



**ТЕХНИЧЕСКИ УНИВЕРСИТЕТ – СОФИЯ**

**Факултет „Компютърни системи и технологии“**

**Катедра „Компютърни системи“**

**маг. инж. Мария Петрова Влахова-Такова**

**МЕТОДИ И АЛГОРИТМИ ЗА ПЕРСОНАЛИЗИРАНИ  
СИСТЕМИ ЗА ПРЕПОРЪКИ**

**А В Т О Р Е Ф Е Р А Т**

на дисертация за придобиване на образователна и научна степен  
**"ДОКТОР"**

Област: 5. Технически науки

Професионално направление: 5.3 Комуникационна и компютърна техника

Научна специалност: „Системи с изкуствен интелект“

**Научен ръководител:**

**проф. д-р инж. Милена Кирилова Лазарова–Мицева**

СОФИЯ, 2025 г.

Дисертационният труд е обсъден и насочен за защита от Катедрения съвет на катедра „Компютърни системи“ към Факултет „Компютърни системи и технологии“ на ТУ–София на редовно заседание, проведено на 10.11.2025 г.

Публичната защита на дисертационния труд ще се състои на 10.02.2025 г. от 15:00 часа в Конферентната зала на БИЦ на Технически университет – София на открито заседание на научното жури, определено със заповед № ОЖ-5.3.-61/21.11.2025 г. на Ректора на ТУ–София в състав:

1. Проф. д-р Румен Трифонов – председател
2. Доц. д-р Ива Николова – научен секретар
3. Проф. д-р Милена Милева-Карова
4. Проф. д-тн Красимира Стоилова
5. Проф. д-р Мариана Горанова

Рецензенти:

1. Проф. д-р Румен Трифонов
2. Проф. д-тн Красимира Стоилова

Материалите по защитата са на разположение на интересуващите се в канцеларията на Факултет „Компютърни системи и технологии“ на ТУ–София, блок №1, кабинет №1443а.

Дисертантът е редовен докторант към катедра „Компютърни системи“ на Факултет „Компютърни системи и технологии“. Изследванията по дисертационната разработка са направени от автора, като някои от тях са подкрепени от научноизследователски проекти.

Автор: маг. инж. Мария Петрова Влахова-Такова

Заглавие: Методи и алгоритми за персонализирани системи за препоръки

Тираж: 30 броя

Отпечатано в ИПК на Технически университет – София

# **I. ОБЩА ХАРАКТЕРИСТИКА НА ДИСЕРТАЦИОННИЯ ТРУД**

---

## **Актуалност на проблема**

В условията на все по-нарастваща дигитализация на обществото, потребителите ежедневно взаимодействат с огромни обеми от информация, съдържание и продукти в онлайн среда. Това води до феномена на информационно претоварване, при който вземането на ефективни и информирани решения се затруднява значително. В този контекст системите за препоръки се явяват ключова технология, способна да филтрира и персонализира достъпа до информация в съответствие с индивидуалните нужди, поведение и предпочитания на всеки потребител.

Същевременно се наблюдава разширяване на обхвата на препоръчителни системи от класическите домейни като електронна търговия и мултимедийни платформи към нови области с висока социална и образователна значимост. Една от тези области е създаването на ефективни визуални презентации – критично умение както в академичния, така и в бизнес контекста. Съществува отчетлива липса на решения, които да предоставят образователна, персонализирана и обяснима обратна връзка относно качеството на презентацията, нейната четливост, структура и визуална ефективност. Традиционните подходи за подобряване на презентации разчитат основно на шаблони или ръчни дизайнерски насоки, което ограничава възможността за динамично и персонализирано обучение на потребителите. Това поражда необходимост от интелигентна препоръчителна система, която не само да идентифицира слабостите на дадена презентация, но и да предлага целенасочени и обосновани препоръки с обучителна стойност. Актуалността на темата се обуславя от сливането на две глобални тенденции: необходимостта от автоматизирани, но адаптивни технологии, които да подобряват дигиталните комуникационни умения на потребителите, и еволюцията на персонализираните препоръки чрез методи на изкуствен интелект. Разработването на система за персонализираните препоръки в сферата на презентациите не само допринася за развитието на интелигентните образователни технологии, но и има потенциално въздействие върху повишаване на визуалната грамотност, ефективността на комуникацията и увереността на презентаторите.

## **Цел на дисертационния труд, основни задачи и методи за изследване**

Целта на научните изследвания в дисертационния труд е да се проектира, разработи и валидира персонализирана препоръчителна система, която използва съвременни подходи от областта на машинно и дълбоко обучение за анализ и подобряване на презентации чрез интелигентна, обяснима и обучителна обратна връзка. Въз основа на мултимодален анализ и комбиниране на визуални и стилови характеристики на слайдове, системата цели не само да подпомогне потребителя в идентифицирането и коригирането на конкретни проблеми, но и да насърчи усвояването на дългосрочни комуникационни и визуални умения. Разработването на такъв инструмент отговаря на съвременните изисквания за адаптивно, технологично и образователно подпомагане на потребителите и представлява значим принос в пресечната точка между изкуствения интелект, визуалната грамотност и дигиталното обучение. Дисертационният труд има за цел да даде методологичен принос в създаването на интелигентни обучителни технологии с висока степен на персонализация и обяснимост.

За постигането на целта са дефинирани следните основни задачи: да се проучат литературни източници в областта на персонализираните системи за препоръки; машинно и дълбоко обучение и тяхното приложение в системи за препоръки; презентационни умения и добри практики в създаване на презентации; да се разработи методология за създаване на персонализирана препоръчителна система, използваща дълбоки невронни мрежи за анализ на визуални и стилови характеристики на слайдове; да се предложат хибридни архитектури на дълбоки невронни мрежи за персонализирана препоръчителна система, които използват различни входове и имат разнообразни възможности за анализ и оценка. да се проектира и имплементира прототип на система за препоръки на базата на предложената методология и хибридни архитектури на дълбоки невронни мрежи, насочена към идентифициране на визуални, структурни и стилови проблеми в презентации и генериране на персонализираните обучителни препоръки; да се проведат експерименти за валидация, оценка и верификация на предложената методология и за анализ на ефективността, точността и приложимостта на системата в реална среда чрез използване на метрики за оценка и сравнение с алтернативни решения.

Научните методи, които са използвани са проучване на литературни източници, теоретично изследване на методи и алгоритми за машинно и дълбоко обучение в системи за препоръки, сравнителен анализ и синтез на хибридни архитектури на дълбоки невронни мрежи за генериране на персонализирани препоръки, проектиране прототип на интелигентна система за препоръки, на експериментална оценка, статистически методи за обработка на експериментални резултати.

## **Научна новост**

Предложени са два подхода за мулти-етикетна класификация на изображения за стилова съгласуваност на презентации: с адаптиране на алгоритъм чрез единен модел и с трансформация на проблема чрез независими бинарни класификатори. Предложена е дву-клонова архитектура на конволюционна невронна мрежа за оценка на стилова съгласуваност на слайдове в презентация. Предложена е концептуална архитектура и е разработен прототип на система за препоръки за интелигентно и персонализирано създаване на презентации, която комбинира поведенчески, съдържателен и визуален анализ на презентации, която интегрира три основни модула за мулти-етикетна класификация на съдържателни и структурни характеристики, за оценка на стилова съгласуваност и хибриден препоръчващ модел с филтриране по съдържание и машинно обучение. Предложена е цялостна MLOps инфраструктура за внедряване на системата за препоръки за интелигентно и персонализирано създаване на презентации. Предложен е подход за формиране и структуриране на комплексен набор от данни с интегриране на различни източници и типове информация. Предложена е методология за събиране, аотиране и предобработка на синтетични и реални данни за структурно-съдържателен анализ и стилова съгласуваност на презентации и са създадени два специализирани набора от данни: ръчно аотиран корпус и синтетичен набор от изображения.

## **Практическа приложимост**

Разработената интелигентна препоръчителна система има практическа стойност и приложимост в образователната и професионалната среда като ефективен инструмент за обучение и самооценка на презентационни умения, приложим в университетски курсове, обучителни центрове и корпоративни обучения. Модулната архитектура и използването на съвременни MLOps подходи позволяват лесно внедряване в реална среда като самостоятелно приложение или като модул интегриран в платформи за електронно обучение и инструменти за създаване на презентации. Проведените експерименти доказват ефективността, надеждността и устойчивостта на системата при реални условия, което потвърждава приложимост и потенциал за бъдещо разширяване в други домейни на визуалната комуникация и дигиталното обучение.

## **Апробация**

Резултатите от дисертационния труд са разглеждани, обсъждани и публикувани в Списание “Computer and Communication Engineering”; Научна конференция EUt+ ELaRa “Technologies and Techniques to Support Sustainable Education in the Academic Sphere”, София, България, 14–15 декември, 2022; Десета международна научна конференция “Computer Science”, София, България, 30 май – 2 юни, 2022; Международно научно списание Journal on Information Technologies & Security; Международно научно списание Journal Engineering, Technology & Applied Science Research; Международно научно списание MDPI Electronics; Научен семинар на докторанти във Факултет по Компютърни системи и технологии, 12 юни 2020, 25 юни 2021, 24 юни 2022, 7 февруари 2025.

## **Публикации**

Основни постижения и резултати от дисертационния труд са публикувани в 6 научни статии, от които 2 самостоятелни. Научните статии са представени и публикувани в национални и международни конференции и международни реферирани и индексирани издания.

## **Структура и обем на дисертационния труд**

Дисертационният труд е в обем от 185 страници, като включва увод, четири глави за решаване на формулираните основни задачи, списък на основните приноси, списък на публикациите по дисертацията и използвана литература. Цитирани са общо 143 литературни източници, всичките на латиница. Работата включва общо 36 фигури и 26 таблици. Номерата на фигурите, таблиците и използваната литература в автореферата съответстват на тези в дисертационния труд.

## **II. СЪДЪРЖАНИЕ НА ДИСЕРТАЦИОННИЯ ТРУД**

---

### **ГЛАВА 1. ПЕРСОНАЛИЗИРАНИ СИСТЕМИ ЗА ПРЕПОРЪКИ – СЪЩНОСТ, ВИДОВЕ, ПРИЛОЖЕНИЯ, ИЗПОЛЗВАНИ МЕТОДИ И АЛГОРИТМИ**

Системите за препоръки са технологии, които подпомагат потребителите в избора на релевантно съдържание, продукти и услуги, като намаляват ефекта от информационното претоварване. Те адресират ключови предизвикателства като разнообразие на съдържанието, бърз ръст на данните, неясни предпочитания и риск от „информационен балон“. Основната им цел е да предоставят персонализирани предложения, които подобряват потребителското изживяване и ангажираност.

#### **1.1. Преглед и еволюция на системите за препоръки**

Съвременните дигитални платформи излагат потребителите на огромно и непрекъснато нарастващо количество съдържание, което затруднява информирания избор и води до информационно претоварване. Персонализираните системи за препоръки възникват като ключов механизъм за филтриране и приоритизиране на релевантни обекти (продукти, услуги, медии) спрямо предпочитания, поведение и контекст, с цел да намалят когнитивното натоварване и да подобрят потребителското изживяване. Еволюцията на системите преминава през ясно разграничими етапи. Първо поколение решения разчитат на евристики и експертни правила – ефективни в тесни домейни, но трудни за адаптация и мащабиране. Следващият етап въвежда колаборативната филтрация и филтрацията, базирана на съдържание, които експлоатират сходства между потребители и елементи и позволяват автоматично обобщаване на предпочитания без експертно кодиране. С нарастване на данните се налагат хибридни комбинации и факторизационни модели (напр. матрична факторизация), които извличат латентни фактори и подобряват точността и мащабируемостта. През последното десетилетие доминират методите с дълбоко обучение – рекурентни мрежи, LSTM, автоенкодери и обучение с подсилване, позволяващи моделиране на последователности, сложни нелинейности и вземане на решения в динамична среда. Текущите тенденции включват контекстно-осъзнати, графови и мултимодални подходи, както и обяснимост на препоръките.

#### **1.2. Примери за успешни препоръчителни системи**

Разгледани са добри практики и водещи решения за препоръчителни системи:

1.2.1 Netflix. Мултиалгоритмична система, комбинираща матрична факторизация, хибридни методи (CBF+CF), дълбоки невронни мрежи и контекстни/мулти-раменни бандити за адаптивни препоръки в реално време; използват се и техники с внимание и вграждания за мащабиране и персонализация.

1.2.2. Amazon. Ядро от item-item колаборативна филтрация, еволюирало към модели със самовнимание, автоенкодери и дори LLM-подкрепени решения; прилага се и DPP за разнообразие и контрол на излишъка при препоръки.

1.2.3. LinkedIn Learning. Невронно колаборативно филтриране с архитектура „две кули“ (learner/course) за вграждания на обучаеми и курсове, плюс модели за предсказване на отговори за фино персонализиране на класациите.

1.2.4. LinkedIn Talent Search. Хибридни модели, графи на знанията и механизми за справедливост/прозрачност; мащабна онлайн/офлайн инфраструктура с A/B тестове за надеждно съвпадение кандидат-позиция.

1.2.5. YouTube. Двуетапна архитектура с дълбоко обучение (кандидатиране + класиране), последователно моделиране на гледанията и корекция на позиционни изкривявания; обучение от неявна обратна връзка (напр. време на гледане) и многозадачни MME компоненти.

1.2.6. Spotify. Хибрид от CF и съдържателен анализ с аудио вграждания, допълнен от GNN и обучение с подсилване, насочено към дългосрочна ангажираност (напр. оптимизация на плейлисти като Discover Weekly).

1.2.7. Duolingo. Хибрид от HLR (за памет и „spaced repetition“) и дълбоко обучение с DRL (policy-gradient/actor-critic) за адаптивно предлагане на упражнения в реално време според успехи, грешки и контекст.

### **1.3. Съществуващи инструменти за презентации**

1.3.1. Slidewise: инструмент за Microsoft PowerPoint, който помага за поддържане на техническа консистентност в презентациите. Той открива и коригира несъответствия в шрифтове, изображения и оформлениа, както и оптимизира размера на файловете. Подходящ е за професионалисти, които искат визуално единни и технически изчистени слайдове, но не предлага обучителна или дизайнерска обратна връзка. Плюсозете му са лесна интеграция с PowerPoint, спестява време и осигурява визуална консистентност, а минусите са ,че няма обучителен характер и е платен (79 \$/година).

1.3.2. PitchVantage: AI базиран инструмент за трениране на презентационни умения. Той дава обратна връзка в реално време за темпо, тон, пълнежни думи и контакт с публиката, като симулира реална аудитория. Подходящ е за усъвършенстване на публичното говорене, но не анализира PowerPoint слайдове и не оценява съдържанието или дизайна им. Плюсозете му са, че подобрява увереността и презентационните умения на потребителите и осигурява реалистична тренировъчна среда. Минусите са ,че не анализира слайдове, няма оценка на съдържанието и е достъпен главно чрез образователни или корпоративни пакети.

### **1.4. Изводи към първа глава**

В резултат на направения литературен обзор на същността, историческото развитие, класификацията и реалните приложения в различни домейни на персонализираните системи за препоръки могат да бъдат направени следните по-важни изводи:

- Системите за препоръки са критични за ефективно потребителско взаимодействие в дигиталната ера – от намаляване на когнитивното натоварване до повишена ангажираност;
- Развитието на персонализираните системи за препоръки е свързано с еволюционен преход от прости евристични подходи към съвременни интелигентни системи базирани на дълбоко обучение и контекстно-осъзнати алгоритми;
- Еволюцията от прости правила към дълбоки модели ясно показва нарастващата нужда от персонализация, мащабируемост и интелигентно адаптиране;
- Използваните подходи в персонализираните системи за препоръки за колаборативна и базирана на съдържание филтрация, хибридни модели и факторизационни техники демонстрират разнообразие от стратегии за персонализация според спецификата на данните и целите на системата;
- Комбинирането на методи и внедряването на контекстно-зависими и обясними подходи е ключово за бъдещето на препоръчителни системи ;
- Обучението с подсилване и дълбокото обучение се утвърждават като основни инструменти в персонализираните системи за препоръки за откриване на латентни зависимости, по-добра адаптивност спрямо динамично променящите се предпочитания на потребителите и работа в условия на непълни данни;
- Системите за препоръки се прилагат успешно в различни контексти и приложни области – от развлечения и образование до електронна търговия и езиково обучение;
- Системите за препоръки на Netflix, Amazon, YouTube, Spotify, LinkedIn Learning и Duolingo потвърждават ключовата роля на персонализираните препоръки за повишаване на ангажираността, удовлетвореността и задържането на потребителите независимо от приложния домейн;
- Въпреки наличието на презентационни платформи и инструменти като Slidewise и PitchVantage все още съществува потенциал за по-дълбока персонализация и интелигентна подкрепа на потребителите чрез системи за препоръки.

## **ГЛАВА 2. МЕТОДОЛОГИЯ ЗА СЪЗДАВАНЕ НА ПЕРСОНАЛИЗИРАНИ СИСТЕМИ ЗА ПРЕПОРЪКИ С ИЗПОЛЗВАНЕ НА МАШИННО И ДЪЛБОКО ОБУЧЕНИЕ**

### **2.1. Формулиране на проблема за създаване на персонализирани препоръки за подобрене на презентации**

Качествените презентации са критични в бизнес, образование и наука, но мнозина срещат трудности въпреки наличния софтуер. Често срещаните проблеми при създаването на качествени и въздействащи презентации включват прекалено много текст, липса на визуални елементи, неясна структура, несъответстващи стилове, лоша четливост.

За да се преодолеят тези и други проблеми е необходимо разработването на интелигентна система за препоръки, която да помага на потребителите да подобряват своите презентационни умения. Основната цел на такава система е анализиране на презентациите на потребителите, идентифициране на най-често срещаните проблеми и предоставяне на персонализирани препоръки за подобрене на презентациите. За да бъде ефективна, интелигентна система за препоръки трябва да следва добре утвърдени принципи за създаване на успешни презентации, които се основават на проучвания и най-добри практики в областта на дизайна и визуалната комуникация: А) Привличане на вниманието на аудиторията: балансиран цвят, четими шрифтове, ясни заглавия, интерактивност, добра структура и позициониране и като цяло дали публиката биха намерили слайда за интересен; Б) Използване на визуално съдържание: инфографики, графики/диаграми, качествени изображения вместо дълги текстове; В) Яснота и стегнатост на съдържанието: ключови точки, кратки изречения, булети, прост език; правило „гледай–чувай–разбери“ (прочит за ~10 сек); Г) Спазване на „правило 7×7“: ≤7 реда на слайд, ≤7 думи на ред (незадължително, но полезно ограничение); Д) Добро позициониране на съдържанието: подравняване, бяло пространство, важната информация в горната половина; Е) Четливост на текста: ≥24 pt, професионални шрифтове, висок контраст; Ж) Включване на слайд с дневен ред (агенда): начален слайд с дневен ред за ориентация и последователност; З) Последователност и единен стил: еднаква палитра, шрифтове и оформления без резки промени. Целева аудитория на препоръчителната система е непрофесионалисти в презентирането (студенти, младши бизнес/научни специалисти). Препоръчителната система следва да подпомогне потребителите в създаването на по-ефективни и професионални презентации, като осигури необходимите знания за да станат уверени презентатори.

## 2.2. Методология за създаване на препоръчителна система за подобряване на презентационни умения

Предложената методология за създаване на препоръчителна система за подобряване на презентационни умения включва няколко основни стъпки (фиг. 8) и предвижда използването на три модула: модул за проверка на кохерентност на слайдовете: Pstyle Checker; модул за класификация на всеки слайд: Presentation Checker; модул за генериране на персонализирани препоръки: PAdvisor.

Необходимите данни за функционирането на препоръчителната система се събират от различни източници с цел изграждане на цялостно разбиране за нуждите на потребителите и характеристиките на съдържанието на презентациите като част от тези данни се генерират в процеса на реално използване на системата:

А) *Потребителски профилни характеристики*: профилът съдържа предпочитания (проблемни области за фокус) и демографски данни (местоположение, възраст, роля); в продукцията се събират с анкета; за обучение се използват „генерирани профили“; генерирането включва: уникален ID, тип потребител (бизнес, учител, студент, изследовател, мениджър, технически и др.), местоположение, предпочитания (предварително дефиниран списък теми/умения, разширяем).

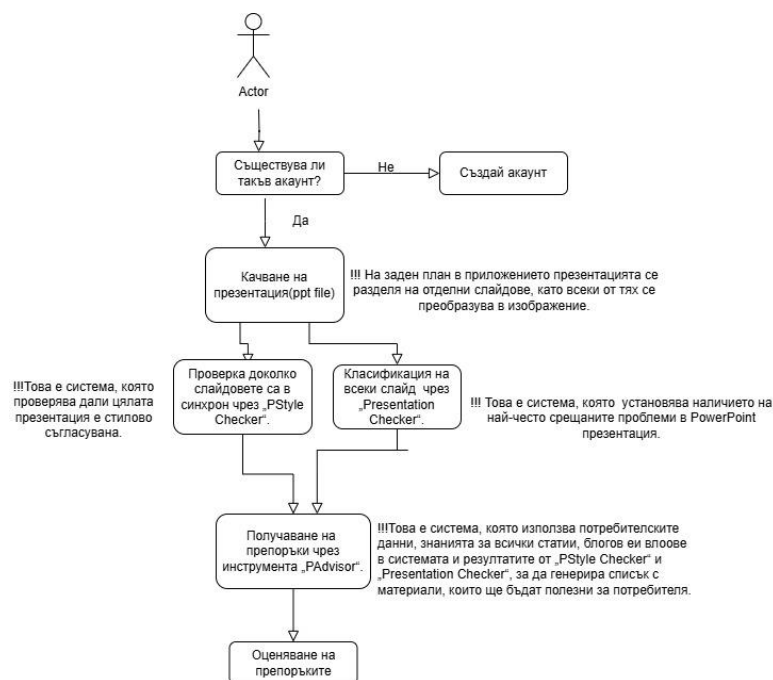
Б) *Презентации*: в реална система: презентациите идват от потребителите; за дисертационното изследване е създадено ръчно анотирано множество от изображения (слайдове на реални презентации) и етикети по категории: интересност; наличие/качество на графики, изображения, таблици, булети, инфографики; четимост; размер на текста; наличие на агенда; предизвикателства при визуални данни: различни гледни точки/мащаби, деформации, закриване, осветление, фон шум/сливане, вътрешнокласово разнообразие, шум/размиване, които затрудняват класификацията и детекцията.

В) *Учебни материали (статии)*: материалите (статии/блогове/видео) се препоръчват според предпочитанията и откритите проблеми; за всеки материал се съхраняват метаданни: за кои проблеми е релевантен, целева аудитория/тип потребители, популярност, гео-ограничения, теми; в реална система популярността се определя от брой отваряния/прочитания; в синтетичните данни е случайна стойност между 0 и 5.

Тъй като няма публично достъпен набор от данни, който да покрива нуждите на системата се използват автогенерирани данни за:

А) Презентации на потребител: заради липса на реални данни се генерират за всеки потребител: брой презентации, проблеми по всяка, случайна дата на подаване.

Б) Матрицата с оценки: създава се персонализирано според минали проблеми, предпочитания, тип презентация, аудитория и популярност на материала.



Фигура 8. Потребителска диаграма на препоръчителна система за подобряване на презентационни умения

Данните се обработват чрез разгъване и бинарно кодиране на категориални полета, като оригиналните списъци се премахват. Резултатът са числови матрици за потребители и материали, готови за изчисляване на сходство и препоръки.

В) Потребителски профили: Профилите се генерират на случаен принцип. Вложените и/или категориални полета (preferences, type, preferred\_presentation\_type) се обработват допълнително чрез разгъване и бинарно кодиране. Оригиналните списъчни полета се премахват, за да няма дублиране на информация.

Г) Данни за препоръки (материали): Характеристиките на материалите се генерират на случаен принцип, след което аналогично с останалите групи се прави бинарно кодиране на списъчните полета presentation\_type и audience\_type и след кодиране оригиналните списъци се премахват.

В резултат се получават съвместими числови матрици за профили и материали, които са удобни за изчисляване на сходство между профили и материали, класификация и други модели с машинно обучение и препоръчване.

### 2.3. Методи за обучение на модели за препоръчителна система за подобряване на презентационни умения

Интелигентната препоръчителна система за подобряване на презентационни умения съдържа три модула, които осигуряват проверка на кохерентност на слайдове, класификация на всеки слайд и генериране на персонализирани препоръки. Целта е да се определят проблеми/грешки в презентациите, да се открие несъответствие в стилового оформлението на презентациите и да се генерират персонализирани препоръки за подобрене на презентационните умения на потребителя и за подобрене на конкретна презентация.

За да бъдат постигнати тези цели, задачата е разделена на следните подзадачи:

1. *Многоетикетна класификация на слайд*:  $g: X_s \rightarrow Y_s, Y_s \subseteq F$ , която определя дали се следват добрите практики;
2. *Оценка на стилова кохерентност*:  $h: (X_{s_i}, X_{s_j}) \rightarrow \kappa_{ij} \in [0, 1]$ , която определя дали презентацията е стилово кохерентна;
3. *Персонализирано препоръчване на обучителни материали*: политика  $\pi(a | u, p, \mathcal{H}_u)$ , чрез която се избират подходящите обучителни материали,

където  $U$  – множество от потребители;  $|U| = N_u$ ;  $P$  – множество от презентации;  $S$  – множество от слайдове; всеки  $p \in P$  е  $p = \{s_1, \dots, s_{mp}\}$ ;  $A$  – множество от обучителни материали (статии/видео);  $F$  – множество от проблемни категории (многоетикетни тагове);  $X_s$  – вектор от характеристики на слайд (визуални/текстови);  $X_u$  – вектор от потребителски характеристики (профил/предпочитания);

$R_i(u, a) \in \mathbb{R}$  – явен/неявен рейтинг на материал  $a$  от потребител  $u$  във време  $t$ ;  $Y_s \subseteq F$  – множество от етикети (проблеми) за слайд  $s$ ;  $\mathcal{D}$  – корпус от данни (изображения на слайдове, метаданни, взаимодействия);  $k_{ij} \in [0,1]$  – мярка за стилова кохерентност (сходство) между слайдове  $s_i$  и  $s_j$ ; ако е оценена от модела  $h$ , бележим  $\hat{k}_{ij}$ ;  $\mathcal{H}_u$  – история на взаимодействията (освен взаимодействия с материали (кликове/оценки), тя включва и данните за проблемите в миналите презентации на потребителя).

В подкрепа на прецизното моделиране е проведен систематичен преглед и оценка на възможните подходи към отделните подзадачи:

#### 1. Многоетикетна класификация на слайд:

Същност на модула: откриване на проблеми/грешки, формулира се като „multi-label“ класификация (един слайд често има няколко проблема).

Подходи: one-vs-one, one-vs-all, „ECOC“, трансформиране (BR, Classifier Chains), адаптирани модели (ML-kNN/NN), ансамбли (RAkEL).

Главни проблеми: дисбаланс на класовете, където възможни решения са undersampling (Random, Tomek/ENN), oversampling („SMOTE“), cost-sensitive learning, ансамбли (EasyEnsemble/BalanceCascade).

#### 2. Силова кохерентност:

Същност на модула: автоматична оценка на сходство по цветове, шрифтове, лога/рамки, подравняване/layout; кохерентността е отчасти субективна.

Подходи: класически похвати: ORB, EMD, SSIM (бързи, но чувствителни към трансформации/шум); дълбоки методи: Triplet/Deep Ranking, Siamese (надеждни при фини различия), Autoencoders (латентни вектори + косинус), CLIP (семантика). Изборът е компромис между точност, данни и ресурси.

Главни проблеми: субективност и липса на единна дефиниция, разделяне „стил“ от „съдържание“, хетерогенност и артефакти на данните, трудно извличане на стилови признаци

#### 3. Персонализирани препоръки:

Сигнали: профил (предпочитания/ниво), текуща презентация (идентифицирани проблеми), минали презентации, доп. ограничения (време/култура/локация).

Модели: базов content + collaborative filtering; DL пергесия (ембединги на user/item + контекст); Autoencoders (по-слаби при оскъдни данни); RL/DQN (онлайн адаптация, баланс exploration–exploitation).

Cold-start: RL/DQN подпомага старта за нови потребители/елементи чрез контролирано изследване.

Обучение/преобучение: предварително обучение, онлайн TD/Q-learning с поетапни Белман актуализации; пълно преобучение само при значим дрифт.

### 2.4. Метрики за оценка на модели

Оценката на ефективността на интелигентната система за препоръки се извършва чрез прилагане на подходящи метрики за всеки от включените модели:

2.4.1. Метрики за оценка на класификационен модел за откриване на проблеми/грешки в презентации: Hamming loss, коефициент на Jaccard, точност, коефициент Капа на Коен, прецизност, пълнота, ROC/AUC, F1 оценка, време за обучение.

2.4.2. Метрики за оценка на модел за изчисляване на стилова съгласуваност на презентации: както при класификационни модели.

2.4.3. Метрики за оценка на препоръчващ модел: средна абсолютна грешка, средноквадратична грешка,  $R^2$ , средна абсолютна процентна грешка, косинусово подобие, индекс на структурно сходство.

2.4.4. Оценка на цялостната система за препоръки: офлайн методи за оценка (разнообразие, новост, серендипитет, обхват, тестове за използваемост), онлайн методи за оценка (процент на кликове, коефициент на конверсии), оценки след внедряване (пристрастие и справедливост, устойчивост, мащабируемост).

## 2.5. Изводи към втора глава

По-важните изводи и резултати във втора глава могат да бъдат обобщени както следва:

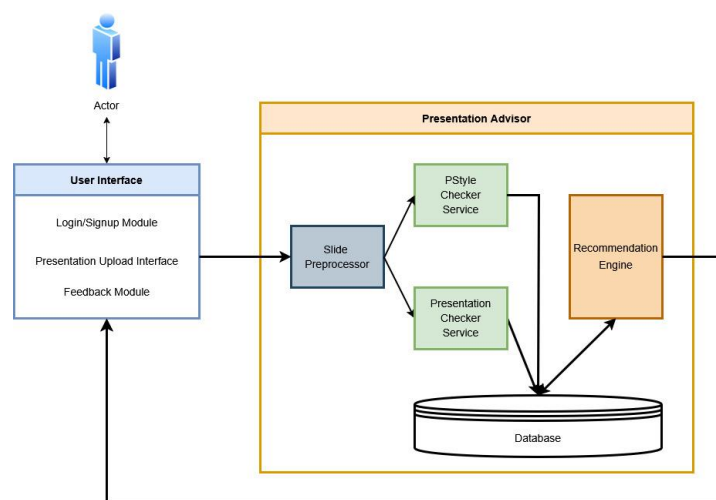
- Много потребители (особено непрофесионалисти) срещат затруднения при създаването на добре структурирани, визуално атрактивни и ангажиращи презентации, като най-често срещаните проблеми са наличие на прекалено много текст, липса на визуални елементи, неясна структура, несъответстващи стилове, лоша четливост;
- Предложена е цялостна методология за персонализирани препоръки за презентации за разработване на персонализирана система за препоръки, ориентирана към подобряване на презентационните умения и визуалната съгласуваност на слайдове чрез методи на машинно и дълбоко обучение;
- Предложена е концепция за интегрирана система, комбинираща поведенчески, съдържателен и визуален анализ на презентации, насочена към дългосрочно подобряване на презентационните умения на потребителите, изградена от три основни модула: за автоматизирана оценка на стилова кохерентност, за откриване на проблеми и грешки в слайдове, за генериране на персонализирани обучителни препоръки;
- Предложен е подход за формиране и структуриране на комплексен набор от данни, който интегрира разнообразни източници и типове информация, необходими за изграждането на система за персонализирани препоръки и включва обединяване на потребителски профили, анотирани презентации, учебни материали и автогенерирани данни;
- Предварителната обработка на данните изисква задаване на анотации, преодоляване на проблема с липсващи стойности и нормализация на визуални и текстови характеристики;
- Разработена е процедура за предварителна обработка и унифициране на данните в числов формат, съвместим с алгоритми за машинно и дълбоко обучение, която включва процедура за предварителна обработка и унифициране на данните, включваща преобразуване на всички категориални и текстови характеристики в числов формат чрез бинарно кодиране; изравняване на вложени структури за постигане на съвместимост с алгоритми за машинно и дълбоко обучение; премахване на несъвместими и дублирани атрибути; обединяване на данни от различни източници в унифицирана структура, позволяваща лесно сравнение и изчисляване на сходство;
- За обучение на модели за всяка от подзадачите на съответните модули на системата (откриване на проблеми в съдържанието на презентации, оценка на стилова съгласуваност между слайдове, генериране на препоръки адаптирани към индивидуалния потребителски контекст) могат да бъдат приложени и адаптирани усъвършенствани алгоритми за машинно и дълбоко обучение;
- Дефинирани са ясни критерии и е определен набор от метрики за класификационните и препоръчителните модули, осигуряващ обективна оценка на ефективността на всеки компонент и на цялостната система, в това число офлайн метрики (точност, F1, MAE), онлайн методи (A/B тестване, потребителска обратна връзка), както и методи за последваща оценка след внедряване на системата в реална среда, което осигурява надежден механизъм за валидиране и адаптиране на системата спрямо нуждите на крайните потребители.

## ГЛАВА 3. ПРЕПОРЪЧИТЕЛНА СИСТЕМА ЗА СЪЗДАВАНЕ НА ПРЕЗЕНТАЦИИ

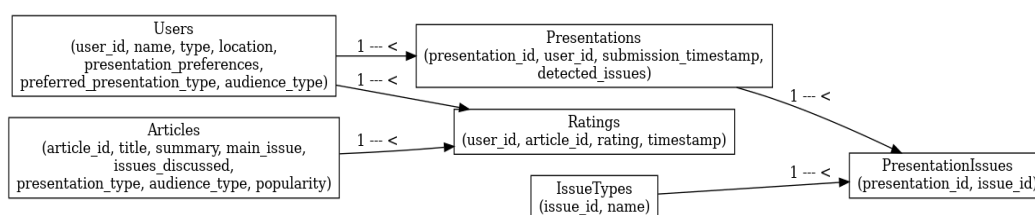
На базата на методологията предложена във втора глава е представено концептуалното разработване на препоръчителна система, чиято основна цел е да подпомага създателите на презентации чрез интелигентни, автоматизирани препоръки.

### 3.1. Концептуална архитектура на препоръчителна система за създаване на презентации

Представени са практическите стъпки по изграждане, обучение и внедряване на препоръчителната система „Presentation Advisor“. На основата на предложената концептуална архитектура са разгледани програмната реализация, моделите, процесът по обучение, внедряване (деплоймънт) и стратегии за поддръжка на системата. Препоръчителната система за презентации е реализирана като модулна микросървисна архитектура, базирана на услуги (service-based), която има за цел да анализира PowerPoint презентации и да предоставя персонализирани, практически приложими препоръки за подобряване както на съдържателното качество, така и на визуалната съгласуваност на слайдовете. Концептуалната архитектура на системата е представена на фиг. 9.



Фигура 9. Концептуална архитектура на препоръчителната система за презентации



Фигура 10. Архитектура на базата данни на препоръчителната система за презентации

След качване на презентация, препроцесорът разделя файла на отделни слайдове и стандартизира входните данни за последващ анализ. Подготвените слайдове се подават паралелно към две независими услуги: „Presentation Checker Service“ – анализира логическата структура, количеството текст, подредбата на елементите и използването на графики, таблици, булети и инфографики; „PStyle Checker Service“ – оценява цвятова палитра, шрифтове, оформление, размери на текстови полета и други характеристики на стиловата съгласуваност чрез Dual-Branch CNN. Получените резултати се съхраняват в централна PostgreSQL база данни (фиг. 10) и се използват от „Recommendation Engine“ за формиране на персонализирани препоръки. Цялостната оркестрация, маршрутизиране и взаимодействие между отделните модули се управлява от Backend API, реализиран с FastAPI. Моделите се разгръщат в изолирани Docker контейнери, а за мониторинг и визуализация се използват Prometheus и Grafana. Деплоймънт архитектурата е показана на фиг. 11. Системата е проектирана за гъвкавост, разширяемост и лесна бъдеща поддръжка.

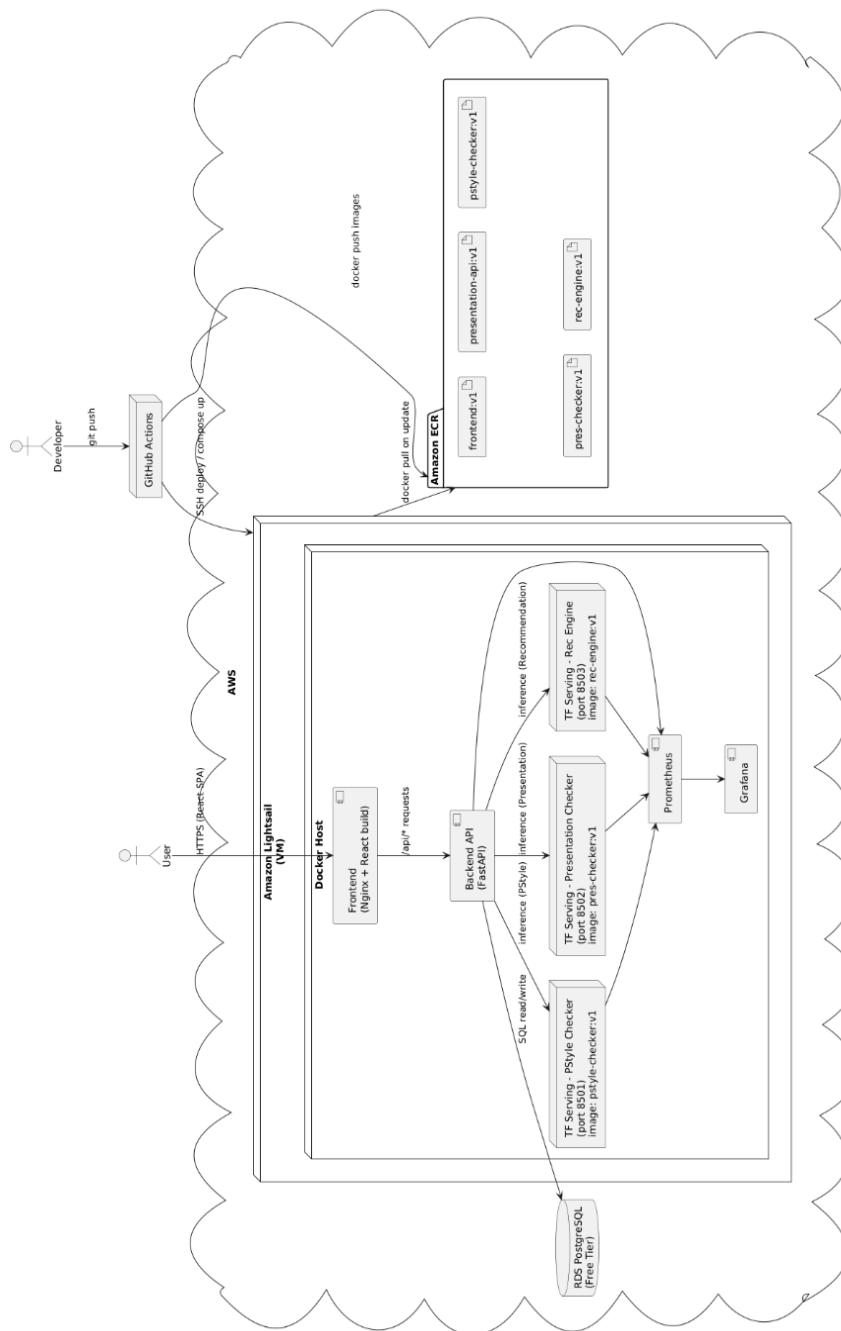
### 3.2. Модул за проверка на съдържание и структура на презентации

Модулът Presentation Checker Service е основен компонент на системата „Presentation Advisor“, изпълняващ функцията на първичен анализатор на отделните слайдове. Той открива и категоризира съдържателни и структурни проблеми, които значително влияят върху качеството и възприемането на презентацията. За изграждането и обучението на модела е използван специално създаден корпус от 912 ръчно анотирани изображения на слайдове, всеки маркиран с десет бинарни етикета (табл. 4), примери от които можете да видите на фигура 12.



Фигура 12. Примерни изображения от наборът от данни за обучение на модел за проверка на съдържание и структура на презентации

Корелационен анализ чрез коефициент на Пийрсън не показва значителни зависимости между отделните класове (табл. 5).



Фигура 11. Деплоймънт диаграма на препоръчителната система за презентации

Таблица 4. Разпределение на етикетите в експерименталния мулти-етикетен набор от изображения

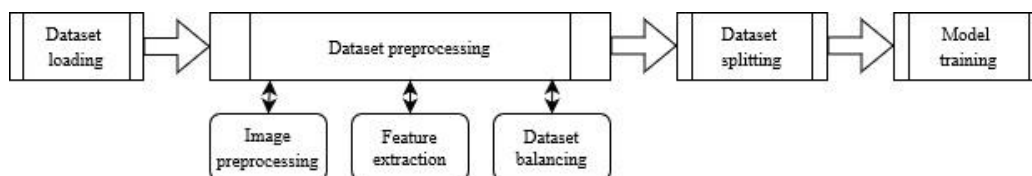
| Етикет          | Брой изображения (%) | Етикет          | Брой изображения (%) |
|-----------------|----------------------|-----------------|----------------------|
| L1: bullets     | 15.8                 | L6: tables      | 4.9                  |
| L2: interesting | 39.7                 | L7: positioning | 79.8                 |
| L3: readable    | 62.2                 | L8: pictures    | 67.5                 |

|                  |      |                  |      |
|------------------|------|------------------|------|
| L4: agenda       | 0.7  | L9: infographics | 22.0 |
| L5: size_of_text | 58.7 | L10: graphics    | 37.5 |

Таблица 5. Корелационна матрица на експерименталния мулти-етикетен набор от изображения

| Етикет     | L1      | L2      | L3      | L4      | L5      | L6      | L7      | L8      | L9      | L10     |
|------------|---------|---------|---------|---------|---------|---------|---------|---------|---------|---------|
| <b>L1</b>  | 1.0000  | 0.0222  | -0.0291 | 0.0027  | -0.1272 | -0.0305 | -0.1237 | -0.0016 | -0.1787 | 0.0742  |
| <b>L2</b>  | 0.0222  | 1.0000  | 0.0403  | -0.0372 | 0.1713  | -0.0536 | 0.1865  | 0.0575  | 0.1983  | 0.1493  |
| <b>L3</b>  | -0.0291 | 0.0403  | 1.0000  | 0.0664  | 0.3946  | -0.0097 | 0.2920  | -0.1582 | 0.1271  | -0.0252 |
| <b>L4</b>  | 0.0027  | -0.0372 | 0.0664  | 1.0000  | -0.0630 | -0.0192 | 0.0422  | 0.0262  | 0.0417  | -0.0335 |
| <b>L5</b>  | -0.1272 | 0.1713  | 0.3948  | -0.0631 | 1.0000  | -0.0274 | 0.0925  | 0.0330  | 0.0227  | -0.1115 |
| <b>L6</b>  | -0.0305 | -0.0535 | -0.0097 | -0.0192 | -0.0274 | 1.0000  | -0.0181 | -0.1162 | -0.0330 | -0.1055 |
| <b>L7</b>  | -0.1237 | 0.1865  | 0.2920  | 0.0422  | 0.0925  | -0.0181 | 1.0000  | 0.0069  | 0.1604  | 0.1351  |
| <b>L8</b>  | -0.0016 | 0.0575  | -0.1582 | 0.0262  | 0.0330  | -0.1162 | 0.0069  | 1.0000  | 0.0232  | -0.1861 |
| <b>L9</b>  | -0.1787 | 0.1983  | -0.1271 | -0.0417 | 0.0227  | -0.0330 | 0.1604  | 0.0232  | 1.0000  | 0.1931  |
| <b>L10</b> | 0.0743  | 0.1493  | -0.0252 | -0.0335 | -0.1115 | -0.1055 | 0.1351  | -0.1861 | 0.1931  | 1.0000  |

Един слайд може да съдържа повече от един проблем, което налага мулти-етикетна класификация. Стъпките на обработка в предложения подход за трансформация на проблема при мулти-етикетна класификация на изображения са показани на фиг. 13.

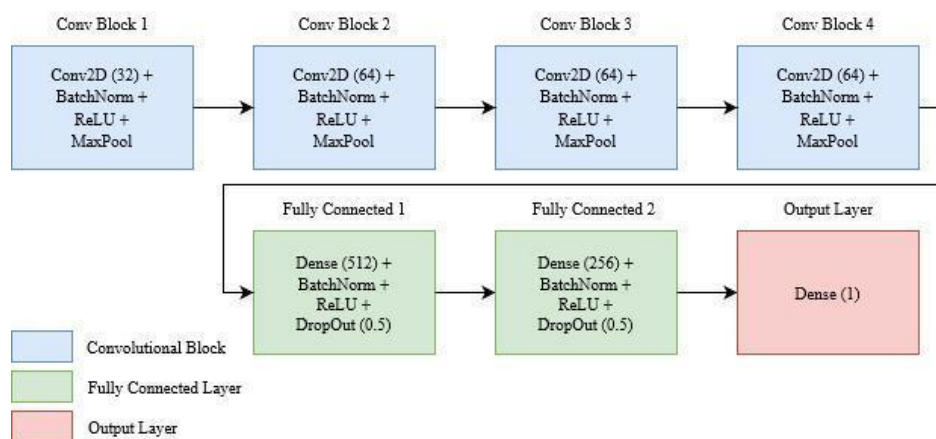


Фигура 13. Стъпки на обработка при подхода за трансформация на проблема за мулти-етикетна класификация

За решаването на задачата са разгледани два основни подхода:

- А) *Трансформиране на проблема* – мулти-етикетната задача се превръща в множество бинарни класификатори (по един за всеки етикет). Това позволява използване на добре познати бинарни алгоритми и независима оптимизация за всяка категория. Стъпки за предварителна обработка са:
1. Преоразмеряване и нормализация на слайдовете до  $224 \times 224$  пиксела и скалиране на стойностите в  $[0,1]$ .
  2. Извличане на характеристики с предварително обучен ResNet50, намалявайки изчислителната сложност.
  3. Премахване на дубликати чрез косинусова прилика.
  4. Балансиране на класовете чрез undersampling за смекчаване на дисбаланса.

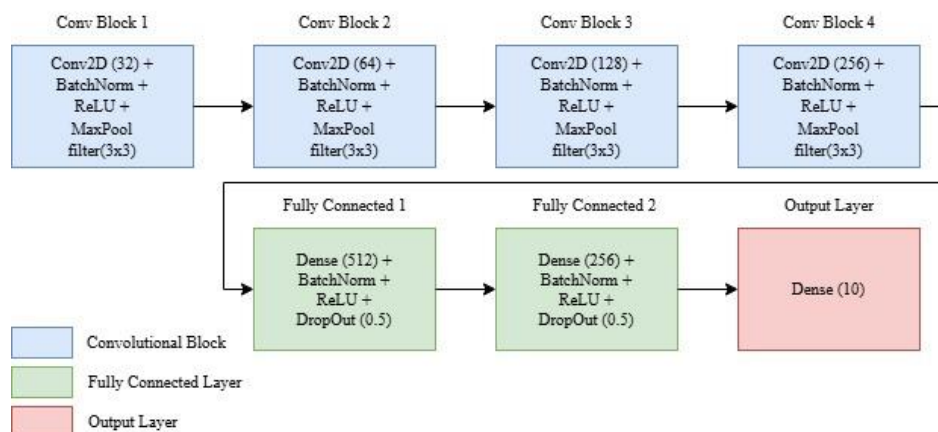
Архитектурата на предложената конволюционна невронна мрежа, представена на фиг. 14, включва: четири конволюционни блока с Convolution  $\rightarrow$  BatchNorm  $\rightarrow$  ReLU  $\rightarrow$  MaxPooling; два плътни слоя (Fully Connected) за интегриране на пространствени характеристики; изходен слой със сигмоидна активация. Обучението и използваните хиперпараметри са както следва: 200 епохи, Batch size = 16, Learning rate = 0.001, оптимизатор Adam; оптимизация чрез GridSearchCV; ранно спиране (EarlyStopping) – за предотвратяване на overfitting; стратегия за разпределение на данните: 80% за обучение (с 3-кратно кръстосано валидиране), 20% за тест. Кръстосаното валидиране осигурява по-надеждна и обективна оценка на модела като намалява риска от преобучение и подобрява генерализационната способност.



Фигура 14. Архитектура на конволюционна невронна мрежа за мулти-етикетна класификация на изображения в препоръчителната система за презентации с подход с трансформиране на проблема

Б) *Адаптиране на алгоритъма* – създаване на единен мулти-етикетен модел с изход, който предсказва вероятности за всички етикети едновременно. Този метод може да улови зависимости между етикетите, ако има такива, но при небалансирани данни се наблюдават предизвикателства като преобучаване, ограничена мащабируемост (при промяна на броя класове е необходимо пълно преобучение), затруднена интерпретация на резултатите, невъзможност да се балансират данните. Архитектурата на предложената дълбока конволюционна невронна мрежа за този подход е показана на фиг. 14.

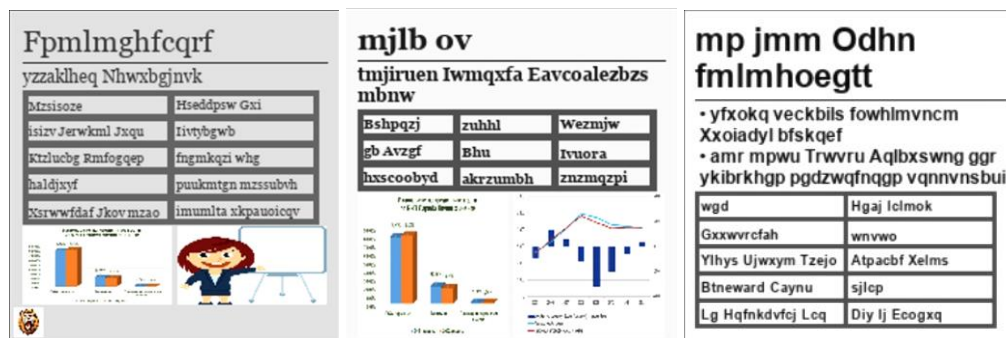
След обучението, моделът идентифицира проблемите с висока точност и подава резултатите към Recommendation Engine, където се формулират персонализирани препоръки за подобрене на презентациите.



Фигура 15. Архитектура на дълбока конволюционна невронна мрежа за мулти-етикетна класификация на изображения в препоръчителната система за презентации с подход с адаптиране на проблема

### 3.3. Модул за проверка на стилова съгласуваност на презентации (PStyle Checker Service)

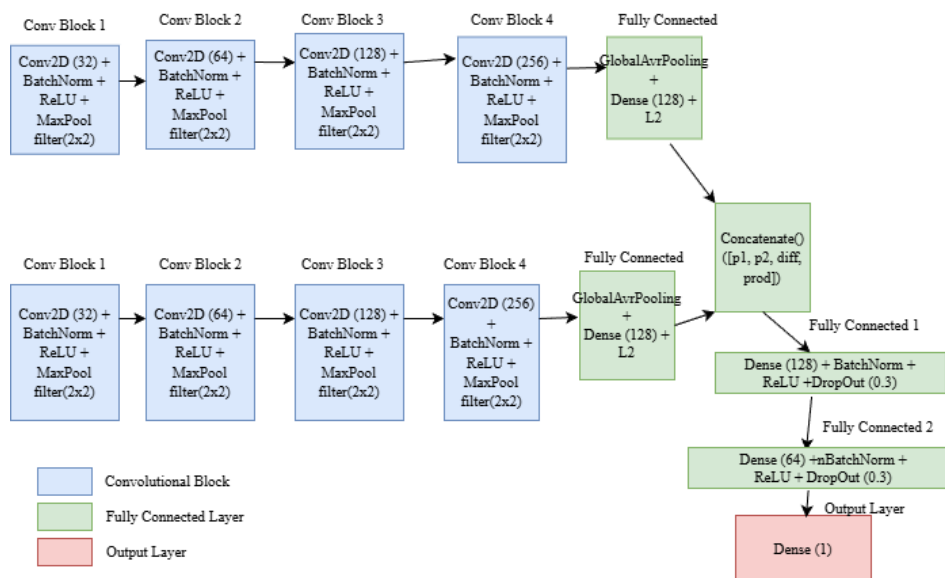
За обучението на модула е създаден синтетичен корпус от 6000 двойки изображения, генерирани с контрол върху ключови визуални параметри: цвят на фона, наличие и позиция на изображение, тип и цвят на шрифт, наличие и позиция на лого, дължина на текстовите области, наличие и избор на изображения и графики, наличие и големина на таблици, както и наличие и стил на рамки. Примерни изображения са показани на фиг. 16. Двойките изображения са равномерно разпределени между категории „сходни“ и „несходни“ стилове, като всички са преоразмерени до 224×224 пиксела.



Фигура 16. Примерни изображения от генерирания набор от данни за обучение на модел за съгласуваност на презентации

Разгледани са няколко подхода за оценка на стилова съгласуваност:

1. *Класически метрики за визуално сходство*: използвани са Structural Similarity Index (SSIM), Oriented FAST and Rotated BRIEF (ORB) и Earth Mover's Distance (EMD). Тези методи са изчислително ефективни и лесни за имплементация, но експериментите показват, че: SSIM не успява да улови фини разлики в шрифтове, цветови нюанси и подравняване; ORB като локален дескриптор е слабо чувствителен към глобални стилови особености; EMD има висока изчислителна сложност при по-големи изображения.
2. *Класически машинно-обучителни модели*: Logistic Regression, Support Vector Machines (SVM) и Random Forest, обучени върху ръчно извлечени визуални характеристики позволяват добра интерпретируемост, но са ограничени при сложни случаи на стилови несъответствия, където ръчните признаци не са достатъчни.
3. *Сиамски невронни мрежи*: архитектури, които сравняват изображения чрез евклидови или косинусни разстояния между embedding вектори, показват по-добра ефективност от класическите подходи, но са чувствителни основно към елементарни визуални прилики и не улавят по-сложни зависимости в дизайна на презентациите.
4. *Дву-клонова конволюционна невронна мрежа (Dual-Branch CNN)*: Архитектурата използва два входни клона със споделени тегла, всеки съдържащ последователни конволюционни блокове (32–256 филтри, BatchNorm, ReLU, MaxPooling). След Global Average Pooling получените embedding-и преминават през плътни слоеве с L2 регуляризация и се конкатенират заедно с допълнителните взаимни характеристики (разлика, произведение). Обединеният вектор се обработва от два плътни слоя с BatchNorm и Dropout, а изходният слой предсказва двоична стойност (0/1). Схематичният ѝ изглед е показан на фиг. 18, а параметрите са обобщени в табл. 8.



Фигура 18. Подобрена архитектура на невронна мрежа с два клона за определяне на сходство между изображения

Таблица 8. Детайли за вътрешната структура на подобрена архитектурата на невронна мрежа с два клона за определяне на сходство между изображения

| Слой           | Компоненти         | Output Shape     | Параметри | Активация | Допълнителна информация     | Слой           |
|----------------|--------------------|------------------|-----------|-----------|-----------------------------|----------------|
| Input 1        | InputLayer         | (None,224,224,3) | 0         | –         | Първи вход                  | Input 1        |
| Input 2        | InputLayer         | (None,224,224,3) | 0         | –         | Втори вход                  | Input 2        |
| Shared Encoder | Functional         | (None,128)       | 423,232   | –         | Споделен между входовете    | Shared Encoder |
| Lambda 1       | Lambda             | (None,128)       | 0         | –         | Обработка на encoder output | Lambda 1       |
| Lambda 2       | Lambda             | (None,128)       | 0         | –         | Обработка на encoder output | Lambda 2       |
| Concatenate    | Concatenate        | (None,512)       | 0         | –         | Комбиниращ векторите        | Concatenate    |
| Dense 1        | Dense              | (None,128)       | 65,664    | –         | Първи FC слой               | Dense 1        |
| BatchNorm      | BatchNormalization | (None,128)       | 512       | –         | Нормализация                | BatchNorm      |
| Dropout 1      | Dropout            | (None,128)       | 0         | –         | Dropout                     | Dropout 1      |
| Dense 2        | Dense              | (None,64)        | 8,256     | –         | Втори FC слой               | Dense 2        |
| Dropout 2      | Dropout            | (None,64)        | 0         | –         | Dropout                     | Dropout 2      |
| Output         | Dense              | (None,1)         | 65        | Sigmoid   | Краен изход                 | Output         |

След сравнителен анализ на представянето на всички изследвани подходи беше установено, че дву-клоновата конволюционна мрежа превъзхожда останалите варианти, като улавя по-сложни взаимовръзки между стилкови елементи и осигурява по-висока точност, особено при трудни гранични случаи. Поради това тя е избрана като финално решение за имплементация в модула.

### 3.4. Модул за препоръки

За проектирането и обучението на модула е създаден специално генериран набор от данни, включващ информация за потребители, статии, взаимодействия и оценки. Данните са структурирани чрез шаблони, отчитащи често срещани проблеми при презентации като четливост, прекомерен текст, липса на консистентност (табл. 9), контекстуални фактори като тип презентация (табл. 10) и тип аудитория (табл. 11), както и времеви характеристики. Това позволява моделиране на реалистични сценарии и персонализирани препоръки за подобряване на презентации.

Наборът от данни е изграден на четири нива:

1. **Статии** – всяка статия е свързана с основен проблем (напр. „твърде много текст“, „четимост“, „презентационни умения“, „визуализация на данни“), както и допълнителни свързани затруднения

(Таблица 9). Описани са: име, резюме, тип презентация (Таблица 10), подходяща аудитория (Таблица 11), популярност (предишна ангажираност) и дата на публикуване.

2. **Потребители** – профилите са генерирани така, че да представят различни роли (бизнес професионалист, учител, ученик, изследовател, мениджър), географско местоположение, предпочитания към проблеми, типове презентации и аудитории, което дава възможност за персонализиране на препоръките.

3. **Презентации** – за всеки потребител се симулират презентации, съдържащи множество идентифицирани проблеми. Те имат времеви маркери (между 2005 г. и настоящия момент), като по-новите получават по-голяма тежест чрез функция на времеви спад.

4. **Оценки** – матрица с оценки (от 2 до 5) се генерира въз основа на: съответствие между предпочитания и характеристики на статиите, релевантност към тип презентация и аудитория, популярност и актуалност на проблемите. Всяка оценка включва ID на потребител, ID на статия, стойност и времеви маркер.

Общият набор съдържа 25 000 записа и 66 колони, обхващащи интеракции, времеви отпечатыци и предпочитания, със средна стойност на оценките около 3.09. На фиг. 17 е показано разпределението на оценките.

*Таблица 9. Чести проблеми при презентации*

| Проблем            | Кратко описание                                                                                |
|--------------------|------------------------------------------------------------------------------------------------|
| Скучна презентация | Слайдовете не ангажират аудиторията; липсва енергия или визуален интерес                       |
| Графики            | Графиките са неясни, прекалено сложни или неподходящи за съдържанието                          |
| Четливост          | Текстът е труден за четене поради избор на шрифт, контраст на цветовете или фон                |
| Последователност   | Непоследователни стилове между слайдовете – различни шрифтове, цветове или оформления          |
| Изображения        | Използвани са некачествени, неуместни или твърде много изображения                             |
| Маркери            | Прекомерна употреба на булети, което прави слайдовете монотонни и претрупани                   |
| Размер на текста   | Текстът е твърде малък или твърде голям, нарушавайки визуалната йерархия и четливостта         |
| Твърде много текст | Слайдовете съдържат прекалено много текст, което затруднява възприемането                      |
| Таблицы            | Таблиците са твърде подробни, гъсти или трудни за бързо разбиране                              |
| Дневен ред         | Липсва ясна структура или навигация – аудиторията може да се изгуби по време на презентацията  |
| Инфографики        | Инфографиките са объркващи, прекалено сложни или не подпомагат ефективно основното послание    |
| Подравняване       | Елементите в слайдовете са разместени или зле подредени, което влияе на четливостта и логиката |

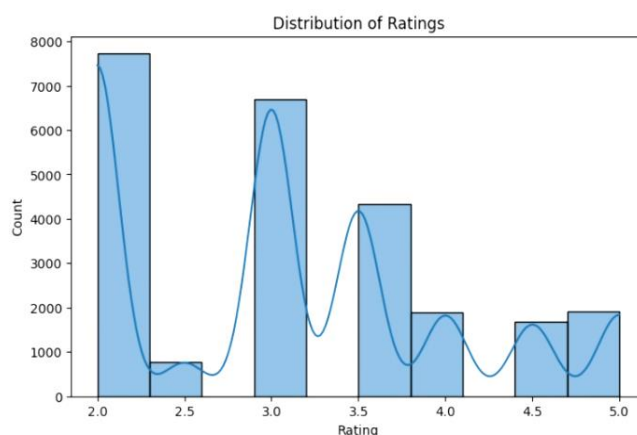
*Таблица 10. Типове презентации*

| Тип презентация | Кратко описание                                                                                |
|-----------------|------------------------------------------------------------------------------------------------|
| Формална        | Структурирана и професионална; използва се често в корпоративна или официална среда            |
| Креативна       | Визуално ангажираща и иновативна; използва разказване на истории, дизайн и нелинейна структура |
| Техническа      | Фокусирана върху подробни данни, процеси или изследвания; предназначена за експертна аудитория |
| Бизнес          | Насочена към вземане на решения, представяне на идеи или отчетност в бизнес контекст           |
| Образователна   | Създадена с цел преподаване или обяснение на тема; използва се в класни стаи                   |

|            |                                                                                |
|------------|--------------------------------------------------------------------------------|
|            | или обучения                                                                   |
| Убедителна | Целта е да убеди или повлияе на аудиторията към определено мнение или действие |

Таблица 11. Типове аудитория

| Тип аудитория         | Кратко описание                                                                                        |
|-----------------------|--------------------------------------------------------------------------------------------------------|
| Академична            | Студенти, изследователи или преподаватели в образователни институции                                   |
| Бизнес                | Професионалисти в корпоративна среда – мениджъри, екипи или вземащи решения лица                       |
| Техническа            | Инженери, разработчици, анализатори или друга аудитория със специализирани технически знания           |
| Деца                  | Малки деца или ученици; изисква просто, ангажиращо и визуално съдържание                               |
| Обща                  | Широка обществена аудитория без специфични предварителни знания; различни нива на подготовка           |
| Изпълнителна          | Висши ръководители или C-level мениджъри; ценят яснота, кратки обобщения и силен визуален ефект        |
| Инвеститорска         | Индивиди или групи, оценяващи възможности – например стартиращи компании или инвестиционни предложения |
| Държавна              | Политици, държавни служители или професионалисти от публичния сектор                                   |
| Уъркшоп/Обучение      | Участници в практически или обучителни сесии, насочени към развиване на умения                         |
| Конферентна аудитория | Смесена публика на събития – може да включва колеги, експерти или представители на медиите             |



Фигура 19. Разпределение на оценките в набора от данни

Разгледани са няколко подхода за изграждане на препоръчващ модел:

1. *Препоръки, базирани на съдържание и колаборативна филтрация*: класическите методи, които използват сходство между потребители и/или елементи въз основа на характеристики или поведение служат като базови модели (baseline). Въпреки простотата и бързата имплементация, тези подходи показват ограничена ефективност в условия на силна разреденост на матрицата и динамични сценарии.
2. *Матрична факторизация и факторизационни модели*: използвани са техники като SVD и ALS, популяризирани от състезанието Netflix Prize, които демонстрират добра мащабируемост и ефективност, но са неподходящи за контекста на „Presentation Advisor“ поради проблеми с нови потребители и нови презентации (cold start).
3. *Автоенкодери*: изследван е вариант на невронни автоенкодери за извличане на латентни представяния на потребители и елементи. Въпреки обещаващи резултати в среди с пълни данни, този подход не е приложим при висока разреденост, каквато е характерна за презентационните сценарии.

4. *Обучение с подсилване*: изследван е модел Deep Q-Network (DQN), който показва потенциал за адаптивност и самообучение в реално време като ефективно адресира проблема с нови потребители и нови елементи. Въпреки това методът изисква значителни изчислителни ресурси и дълга фаза на обучение, което ограничава приложимостта му в продукционни условия.

5. *Дълбоки невронни мрежи*: разгледани са рекурентни архитектури (RNN, LSTM) за моделиране на последователности от потребителски взаимодействия, както и графови невронни мрежи (GNN) и мултимодални подходи, които предлагат висока представителна сила, но страдат от значителна изчислителна сложност и проблеми с интерпретируемостта.

След обстойна сравнителна оценка е предложена хибридна архитектура на невронна мрежа с мулти-колонна структура, която комбинира user-item embeddings (обучени върху взаимодействия), ръчно извлечени характеристики на презентации (поради динамиката на съдържанието) извлечени чрез „Presentation Checker“ и „PStyle Checker“, времеви и поведенчески фактори.

Предложеното решение съчетава:

- *Филтриране, базирано на съдържание* – използва структурните и визуални характеристики на слайдовете за формиране на препоръки;
- *Машинно обучение* – предложената невронна архитектура приема като входни данни латентни векторни представяния на потребителските характеристики и обучителните ресурси, векторизирано представяне на идентифицираните проблеми в презентациите, както и темпорални признаци, описващи контекста на взаимодействието;
- *Утилитарна оценка* – препоръките се преоценяват според дефинирани критерии като популярност, новост и релевантност.

Архитектурата на предложеното решение за Recommendation Engine (фиг. 24) включва няколко „кули“ (towers), всяка специализирана да обработва определен тип данни, аналогично на подход с няколко експерти, всеки от които анализира различен аспект на ситуацията и след това мненията им се обединяват.

#### **1. Кула за потребителите**

- вход: профилни данни (демография, предпочитания, тип потребител, стил на презентации);
- изход: компактно представяне (вектор) на потребителя;
- в модела: по-малко неврони (128), защото потребителите образуват по-ясни групи (фиг. 25).

#### **2. Кула за статиите (обучаващи материали)**

- вход: характеристики на статията (стил, аудитория, дали е текстова или с графики);
- изход: вектор, описващ съдържанието;
- в модела: повече неврони (256), защото статиите са по-разнообразни и трудни за групиране (фиг. 25).

#### **3. Кула за последователности от проблеми**

- вход: към кулата не се подават отделни проблеми от всяка презентация, а един агрегиращ ред, който събира информацията от всички минали презентации на потребителя и е резултат от претегляне на затрудненията – новите проблеми имат по-голяма тежест чрез експоненциална функция на затихване, докато по-старите влияят по-слабо. Така в една компактна форма се съхраняват както основните трудности, така и тяхната динамика във времето;
- изход: вектор, представящ профила на затрудненията на потребителя и показва кои проблеми са най-релевантни в момента и как те ще влияят върху бъдещите препоръки.

#### **4. Кула за времето**

- вход: кога е било последното взаимодействие – ако има взаимодействие се взема се реалното време, ако няма се изчислява „какво ще стане, ако ползвателят кликне сега“.
- данните се кодират циклично (sin и cos трансформации) за да се уловят модели като „по-активен вечер“ или „учене през уикенда“.

#### **5. Обединение на кулите**

- всички вектори от отделните кули се конкатенират;
- следват няколко слоя (Dense + BatchNorm + Dropout), които смесват информацията;
- на финала се изчислява една стойност (оценка) за всяка двойка „потребител – статия“.

#### **6. Утилитарна оценка (Utility score)**

След като се изчисли оценката от обединението на кулите се коригира според съвпадение с аудиторията, свежест и популярност (дълга опашка), разнообразие и новост на съдържанието:

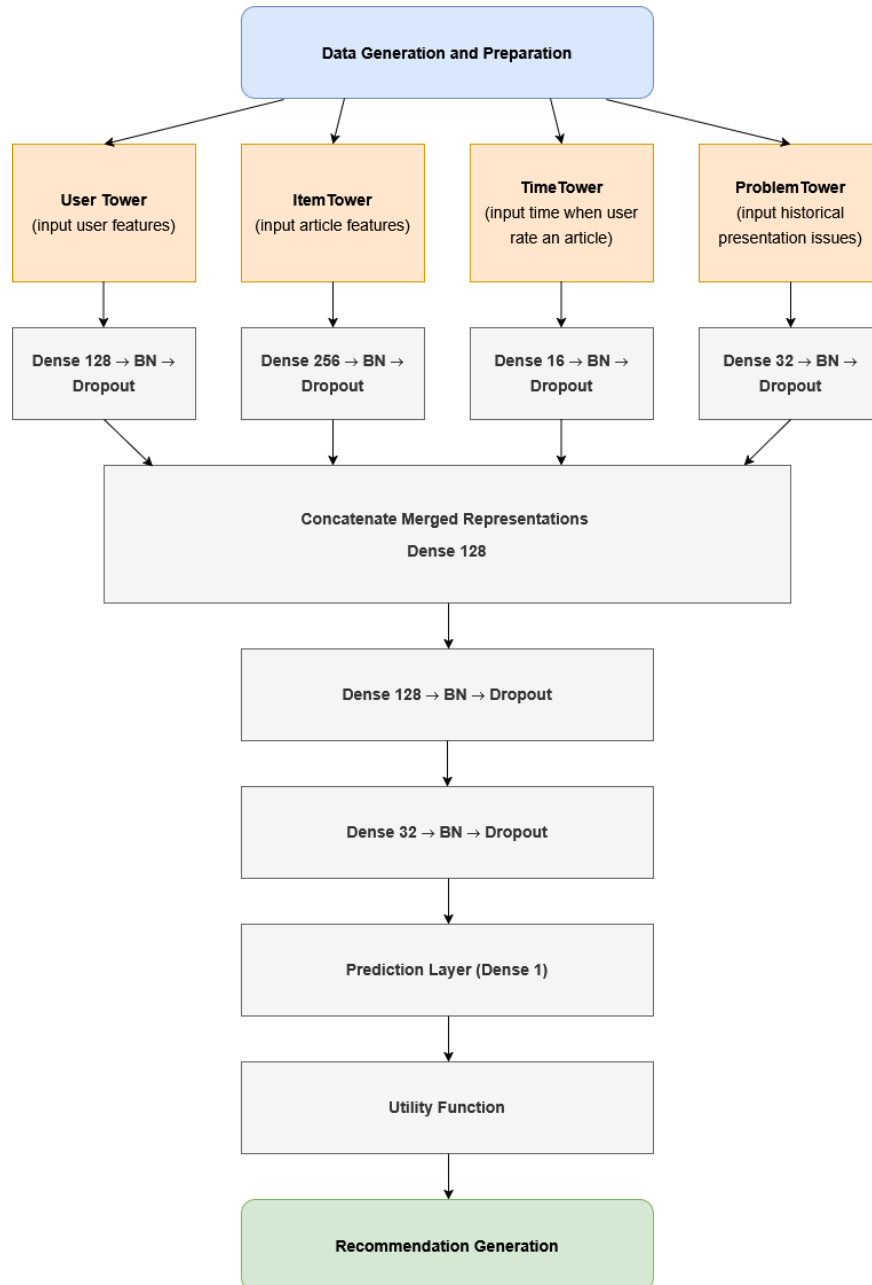
$$utility\_score = predicted\_rating * (1 + \alpha * audience\_boost) * (1 + \beta * long\_tail\_boost)$$

където  $\alpha$  е хиперпараметър, контролиращ усилването на съдържание, съвпадащо с предпочитаната аудитория;  $\beta$  е тегло, регулиращо ефекта на „дългата опашка“ (long-tail boosting); audience\_boost е двоична променлива, отразяваща съвпадение между типа аудитория на потребителя и маркировката на съдържанието; long\_tail\_boost се изчислява по формула:

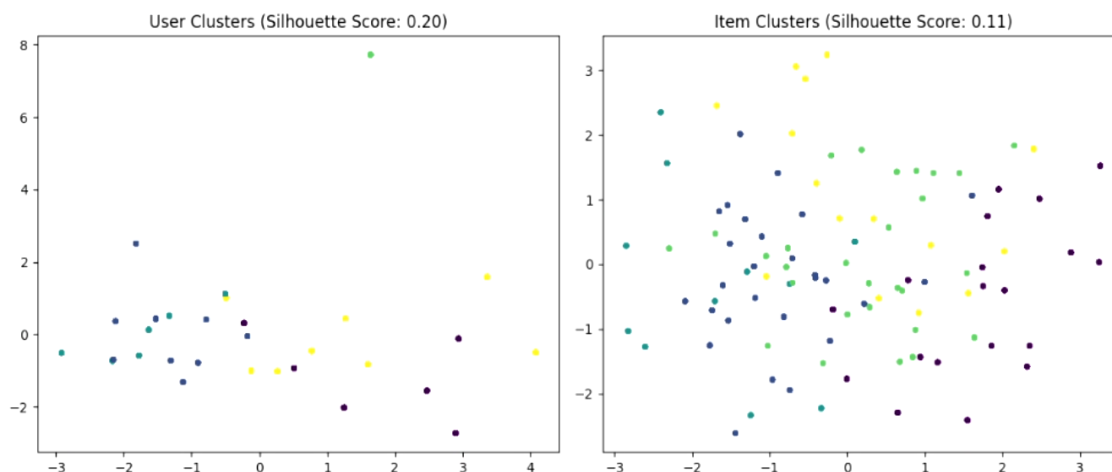
$$\text{long\_tail\_boost} = \gamma * \text{freshness\_score} + (1 - \gamma) * \text{inverse\_popularity}$$

където freshness\_score е нормализирана метрика, отчитаща времето от публикуване на съдържанието; inverse\_popularity е мярка за рядкост/слаба ангажираност с дадено съдържание;  $\gamma$  контролира баланса между актуалност и рядкост.

В конкретната имплементация са използвани следните стойности, установени чрез емпирично валидиране:  $\alpha = 0.2$ ,  $\beta = 0.3$ ,  $\gamma = 0.5$ .



Фигура 24. Архитектура на мулти-колонна невронна мрежа за препоръчителния модел



Фигура 25. Оценка чрез силуетен коефициент на потребителските и предметните (артикулни) клъстери.

С цел да се оцени устойчивостта и генерализиращата способност на модела, е приложена 5-кратна крос-валидация. При този подход обучаващото множество е разделено на пет равни части, като всяка от тях последователно служи като тестово множество, а останалите четири – като обучаващо. Получените резултати са усреднени, което позволява да се минимизира зависимостта от конкретното разпределение на данните и да се постигне по-надеждна оценка на представянето на модела. Моделът осигурява надежден механизъм за откриване на стилови разминавания, които често остават незабелязани от потребителя, но оказват значително влияние върху възприемането на презентацията.

### 3.5 Внедряване на моделите

Всички модели на системата „Presentation Advisor“ са внедрени чрез TensorFlow Serving, опаковани в Docker контейнери за лесно управление и преносимост.

- Архитектура на внедряване (фиг. 11):
  - контейнерите се хостват върху Amazon Lightsail VM, което осигурява надеждна и мащабируема инфраструктура;
  - контейнерните образи се съхраняват и версионират в Amazon ECR, като всяка нова версия се публикува под уникален номер (1/, 2/, ...);
  - REST API интерфейс позволява достъп до препоръчителния модул, като поддържа както единични заявки, така и batch inference за оптимално използване на ресурсите.
- CI/CD процес:
  - GitHub Actions автоматизира целия цикъл: изграждане, тестване и внедряване на моделите;
  - Поддържат се A/B тестване и механизъм за rollback към предишна версия при проблеми.
- Мониторинг и наблюдение:
  - Prometheus събира метрики за работата на системата като точност на препоръките, грешки, латентност на заявките, натоварване на процесора/паметта, брой обслужени заявки и други;
  - Grafana визуализира тези метрики в реално време чрез интерактивни табла (dashboards);
  - включени са аларми и известия при отклонения от нормалната работа (например спад в точността или повишена латентност);
  - наблюдението е интегрирано с процеса по версионизиране, което позволява автоматично връщане към по-стабилна версия на модела при критични проблеми.
- Поддръжка и актуализация:
  - инкрементално обучение – моделите се адаптират към нови данни без пълно преобучение;
  - активно обучение – трудни или неясни случаи се маркират за ръчна анотация;
  - обратна връзка от потребители се използва директно за валидиране и усъвършенстване на препоръките;
  - версионизиране и автоматизирано наблюдение гарантират надеждност, проследимост и бързо възстановяване при грешки.

Така системата за внедряване осигурява не само автоматизация и мащабируемост, но и постоянен мониторинг на качеството и ефективността на препоръките в реална среда.

### 3.6. Изводи към трета глава

В трета глава е представена цялостна концептуална архитектура и компонентите на интелигентна препоръчителната система предназначена да подпомага създаването на висококачествени и визуално съгласувани презентации „Presentation Advisor“. Системата е изградена върху модулна, микросървисна архитектура с акцент върху гъвкавост, мащабируемост и устойчивост, като следва най-добрите практики на MLOps и CI/CD. По-важните изводи и резултати могат да бъдат обобщени както следва:

- Предложена е оригинална концептуална архитектура на препоръчителната система за презентации съчетаваща три основни модула: Presentation Checker Service за мулти-етикетна класификация на съдържателни и структурни характеристики, PStyle Checker Service за оценка на стилова съгласуваност чрез иновативна дву-клонова конволюционна невронна мрежа и Recommendation Engine с хибриден препоръчващ модел, комбиниращ филтриране по съдържание и машинно обучение.
- Предложени са и сравнени два подхода за мулти-етикетна класификация на изображения за стилова съгласуваност на презентации, базирани на трансформация на проблема чрез независими бинарни класификатори и адаптиране на алгоритъма чрез единен модел. Анализирани са ограниченията на подхода с адаптиране на алгоритъма по отношение на преобучаване, мащабируемост и интерпретируемост. Изведени са предимствата на подхода с трансформация на проблема, който е по-ефективен в условия на небалансирани класове.
- Предложена е дву-клонова архитектура на конволюционна невронна мрежа за оценка на стилова съгласуваност на слайдове в презентация, която извлича паралелно характеристики от два слайда, използва конкатенация и допълнителни плътни слоеве за по-прецизно моделиране на взаимовръзките и постига по-висока точност спрямо класически сямски невронни мрежи и традиционни метрики като SSIM или ORB.
- Предложена е методология за събиране, аотиране и предобработка на синтетични и реални данни за структурно-съдържателен анализ и стилова съгласуваност на презентации в условия на ограничена наличност на публични множества данни;
- Създадени са два специализирани набора от данни: ръчно аотиран корпус от 912 реални слайдове за структурно-съдържателен анализ и синтетичен набор от 6000 изображения за стилова съгласуваност с контрол върху визуални параметри.
- Предложена е цялостна MLOps инфраструктура за внедряване, която включва контейнеризация, версионизиране, автоматично деплойване и мониторинг, стратегии за инкрементално и активно обучение с включени механизми срещу катастрофално забравяне и събиране на потребителска обратна връзка за затворен цикъл на подобрене;
- Разработеният прототип на система за препоръки за интелигентно и персонализирано създаване на презентации „Presentation Advisor“ демонстрира ефективна интеграция между модулна архитектура, съвременни модели за машинно обучение и реални потребителски нужди и предлага персонализирани препоръки за подобрене на съдържанието и дизайна на презентации, реално-времеви анализ с обратна връзка от потребителите, възможност за адаптация към нови визуални стандарти, езици и типове презентации.

## ГЛАВА 4. ПЛАНИРАНЕ, ПРОВЕЖДАНЕ НА ЕКСПЕРИМЕНТИ И АНАЛИЗ НА ЕКСПЕРИМЕНТАЛНИ РЕЗУЛТАТИ

Експерименталната проверка има за цел да валидира резултатите от проектирането на препоръчителната система „Presentation Advisor“. Основните задачи са утвърждаване на коректността на поведението на системата и оценка на точността и ефективността на препоръките. Цялостната експериментална платформа е изградена с помощта на програмния език Python и фреймуърка TensorFlow, а хардуерната конфигурация включва процесор Intel(R) Core(TM) i7-8700K @ 3.70GHz и 64 GB RAM. Експериментите се провеждат на ниво отделни модели и цялостна архитектура.

### 4.1. Експериментална оценка на модул за проверка на съдържанието и структурата на презентации (Presentation Checker Service)

Модулет Presentation Checker Service е проектиран да извършва автоматична проверка на основни стилови и структурни характеристики в слайдове от презентации.

Експерименталната оценка има за цел да валидира ефективността на алгоритмите и правилата за верификация на презентациите, както и тяхното въздействие върху крайните препоръки на системата.

Използвани са следните метрики за оценка:

- Hamming Loss – за измерване на грешките при мулти-етикетни предсказания.
- Accuracy – обща точност на класификацията.
- Precision, Recall и F1-Score – стандартни метрики за оценка на качеството на класификацията.
- Jaccard similarity – за оценка на сходството между предсказаните и реалните етикети.
- Kappa coefficient (Cohen's Kappa) – за измерване на степента на съгласие отвъд случайността.
- ROC AUC – за оценка на дискриминационната способност на модела.
- Training time – време за обучение на модела.

Метрики за оценка на модел за мулти-етикетна класификация на изображения в препоръчителната система за презентации с подход за трансформиране на проблема заедно с измерените стойности са представени в табл. 13, а в табл. 14 е дадена точността по класове при подход с трансформиране на проблема. Както се вижда от получените експериментални резултати, модулът демонстрира висока ефективност. Постигната е средна точност е 0.7627, F1 оценката е 0.7889, а стойността на ROC AUC е 0.7622, което доказва, че системата може да предсказва с висока надеждност. Резултатът за Hamming Loss е 0.2373, което показва ниско ниво на грешки в предсказаните етикети. Коефициентът на Капа е 0.5253, което означава, че моделът не взема решения случайно, а демонстрира съществена степен на съгласуваност с реалността.

Анализът на получените експериментални резултати за използваните показатели за оценка на модела ясно демонстрира, че избраният подход е както приложим, така и ефективен в контекста на разглежданата задача. Един от основните недостатъци на метода за трансформиране на проблема е дълготното време за обучение. За използваната експериментална постановка времето за обучение е приблизително 37 минути ако всички модели се обучават последователно (един след друг). Дълготното време за обучение на модела не представлява сериозен проблем за системата тъй като:

- всички бинарни модели могат да се обучават паралелно, което значително съкращава общото време за обучение;
- в реална (онлайн) среда системата използва предварително обучени модели, така че времето за обучение не засяга потребителското изживяване;
- използването на предложените стратегии за частично преобучение и актуализация на моделите осигурява устойчивост и адаптивност на системата във времето при поява на нови типове презентационни елементи.

Резултатите от подхода за адаптиране на проблема са представени в табл. 15 и табл. 16. Получените експериментални резултати показват, че при адаптиране на проблема се наблюдават по-ниски стойности в почти всички ключови метрики за ефективност. Hamming Loss е по-висок – 0.2918 при подхода с адаптиране спрямо 0.2373 при подхода с трансформиране, което означава, че моделът прави повече грешки при предсказване на етикетите. Точността също е по-ниска – 0.7306 при подхода с адаптиране спрямо 0.7627 при трансформирането. Освен това прецизността, пълнотата и F1-оценките както в микро, така и в макро измерване са в полза на метода на трансформиране. Например, микро F1-оценката при адаптиране е 0.6755, докато при трансформиране достига 0.7889, а ROC AUC е съответно 0.6118 при адаптиране спрямо 0.7622 при трансформиране на проблема. Визуално сравнение на резултатите между моделите, които решават квалификационната задача чрез трансформиране на проблема и адаптиране на проблема е показано на фиг. 26.

Таблица 13. Метрики за оценка на модел за мулти-етикетна класификация на изображения в препоръчителната система за презентации с подход за трансформиране на проблема

| <b>Метрика (средно)</b> | <b>Стойност</b> | <b>Метрика (средно)</b> | <b>Стойност</b> |
|-------------------------|-----------------|-------------------------|-----------------|
| Hamming Loss            | 0.2373          | Jaccard similarity      | 0.5790          |
| Accuracy                | 0.7627          | Kappa coefficient       | 0.5253          |
| Precision               | 0.8137          | ROC AUC                 | 0.7622          |
| Recall                  | 0.7164          | Training time           | 37 min          |
| F1-Score                | 0.7889          |                         |                 |

Таблица 13. Точност по класове за мулти-етикетна класификация на изображения в препоръчителната система за презентации с подход за трансформиране на проблема

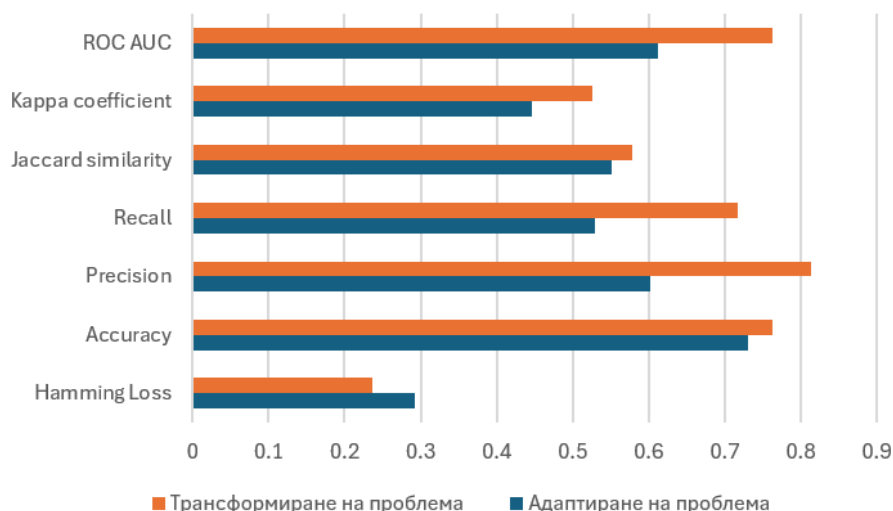
| <b>Етикет</b>    | <b>Точност</b> | <b>Етикет</b>    | <b>Точност</b> |
|------------------|----------------|------------------|----------------|
| L1: bullets      | 0.7213         | L6: tables       | 0.7557         |
| L2: interesting  | 0.8379         | L7: positioning  | 0.8556         |
| L3: readable     | 0.7378         | L8: pictures     | 0.7874         |
| L4: agenda       | 0.8261         | L9: infographics | 0.7007         |
| L5: size_of_text | 0.7935         | L10: graphics    | 0.8493         |

Таблица 14. Метрики за оценка на модел за мулти-етикетна класификация на изображения в препоръчителната система за презентации с подход за адаптиране на проблема

| <b>Метрика (средно)</b> | <b>Стойност</b> | <b>Метрика (средно)</b> | <b>Стойност</b> |
|-------------------------|-----------------|-------------------------|-----------------|
| Hamming Loss            | 0.2918          | Jaccard similarity      | 0.5511          |
| Accuracy                | 0.7306          | Kappa coefficient       | 0.4459          |
| Precision (Macro)       | 0.6013          | Precision (Micro)       | 0.6404          |
| Recall (Macro)          | 0.5297          | Recall (Micro)          | 0.7146          |
| F1-Score (Macro)        | 0.4914          | F1-Score (Micro)        | 0.6755          |
| Training time           | 21.87 seconds   | ROC AUC                 | 0.6118          |

Таблица 15. Точност по класове за мулти-етикетна класификация на изображения в препоръчителната система за презентации с подход за трансформиране на проблема

| <b>Label</b>     | <b>Accuracy</b> | <b>Label</b>     | <b>Accuracy</b> |
|------------------|-----------------|------------------|-----------------|
| L1: bullets      | 0.8092          | L6: tables       | 0.7557          |
| L2: interesting  | 0.4656          | L7: positioning  | 0.6718          |
| L3: readable     | 0.5573          | L8: pictures     | 0.8092          |
| L4: agenda       | 0.6318          | L9: infographics | 0.7328          |
| L5: size_of_text | 0.9542          | L10: graphics    | 0.6183          |



Фигура 26. Визуално сравнение на резултатите между моделите, които решават квалификацията задача чрез трансформиране на проблема и адаптиране на проблема

#### 4.2. Експериментална оценка на модул за проверка на стилова съгласуваност на презентации (PStyle Checker Service)

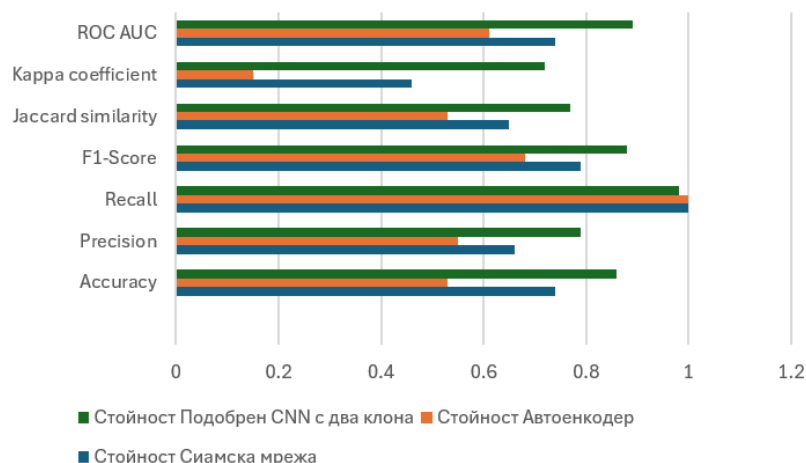
Процесът по проектиране и валидиране на модула за проверка на стилова съгласуваност (PStyle Checker Service) е реализиран в рамките на експериментална методология, насочена към разработването на архитектурно и алгоритмично решение, базирано на дълбоки конволюционни невронни мрежи, което да оценява стиловата съгласуваност между двойки слайдове, представени като изображения.

За целите на експерименталната оценка е изградена система за обработка на данни и обучение на модела, обучена върху синтетично генериран набор от изображения, разделен в съотношение 80:20 за обучение и тестване (4800 двойки за обучение и 1200 за тестване). Обучението е проведено за 100 епохи при размер на пакета 16 и начална стойност на коефициента на обучение  $3 \times 10^{-4}$ , като е приложена петкратна кръстосана валидация и механизъм за ранно спиране за ограничаване на пренастройването (overfitting). Оценката на представянето на предложената архитектура е извършена чрез сравнение със сиамска невронна мрежа и автоенкодерна архитектура. За обективна оценка са използвани множество метрики: Accuracy, Precision, Recall, F1-Score, Jaccard similarity, Kappa coefficient и “ROC AUC”, както и измерване на времето за обучение.

Експерименталните резултати, представени в табл. 18 и фиг 33, показват превъзходство на двузаклонената CNN архитектура спрямо базовите модели. при която се достига Accuracy = 0.86, превишавайки резултатите на сиамската мрежа (0.74) и автоенкодера (0.59). По отношение на Precision (0.79) и F1-Score (0.66) моделът също демонстрира най-добри стойности, като постига оптимален баланс между чувствителност (Recall = 0.98) и контрол над фалшивите положителни резултати. Допълнителните показатели – Jaccard similarity = 0.77, Kappa = 0.73 и ROC AUC = 0.89 – потвърждават надеждността и устойчивостта на предложената архитектура.

Таблица 18. Метрики за оценка на моделите обучени със сиамска мрежа и подобрен модел с конволюционна невронна мрежа с два клона върху усложненото множество данни

| Метрика                  | Dataset size | Accuracy | Precision | Recall | F1-Score | Jaccard similarity | Kappa coefficient | ROC AUC | Training time |
|--------------------------|--------------|----------|-----------|--------|----------|--------------------|-------------------|---------|---------------|
| Сиамска мрежа            | 6000         | 0.74     | 0.66      | 1      | 0.79     | 0.65               | 0.46              | 0.74    | 3 часа        |
| Подобрен CNN с два клона | 6000         | 0.86     | 0.79      | 0.98   | 0.88     | 0.77               | 0.72              | 0.89    | 50 min        |



Фигура 33. Графично представяне на резултатите на подобрения модел

Избраният модел е тестван върху 60 реални анотирани слайда и постига точност 0.8750. Грешките са обясними: пропуска 4 от 16 случая при силно сходни по стил слайдове поради липса на тренировъчни примери с различни цветови варианти, огледални или „рисувани“ презентации и слайдове с номерация.

#### 4.3. Експериментална оценка на модул за препоръки (Recommendation Engine)

Валидирането на модула е извършено чрез синтетично генерирани данни, включващи проблеми в презентации, контекстуални фактори и потребителски предпочитания.

Набора от данни използвани за оценката е изграден от 25 000 взаимодействия между потребители и ресурси, с богата дескрипция на всеки обект. Информацията включва: ID на потребителя, тип аудитория, тип презентация, предпочитания относно стил и визуализация, времеви момент на взаимодействие, както и рейтинг, даден на конкретна препоръка. Данните са структурирани с оглед на улесняване на входа към мулти-входните невронни мрежи. Разделянето на множеството данни е 70% за обучение, 15% за валидация и 15% за тест, като се осигурява стратификация по ключови характеристики. Това разделение осигурява достатъчно примери за обучение, валидация и крайна оценка на генерализацията на модела.

Таблица 19. Оценка на ефективността на моделите за препоръчване

| Модел                                        | MAE  | MSE   | RMSE | време за обучение |
|----------------------------------------------|------|-------|------|-------------------|
| CBF+CF                                       | 1.13 | 1.96  | 1.40 | 0.29 sec          |
| Autoencoders                                 | 3.05 | 10.46 | 3.23 | ~1 hour           |
| RL (DQN)                                     | 3.22 | 12.22 | 3.49 | ~15 min           |
| Hybrid model                                 | 0.11 | 0.033 | 0.18 | ~25 min           |
| Hybrid model + custom pre-trained embeddings | 0.49 | 0.36  | 0.60 | ~1 hour           |

Резултатите от експерименталното сравнение на различните модели за препоръки по отношение на точността на прогнозиране, измерена чрез изчисляване на метриците MAE, MSE и RMSE, както и времето за обучение, са представени в табл. 19. Сравнението е направено за модели за персонализирани препоръки с използване на препоръки базирана на съдържание и система с колаборативна филтрация (CBF+CF), модел с автоенкодерна архитектура, модел с обучение с подсилване (DQN), предложеният хибриден модел с архитектура на мулти-колонна невронна мрежа, както и използване на хибридният модел с предварително обучени embedding представяния.

Хибридният модел показва най-добри резултати спрямо всички оценени варианти – осигурява най-висока точност на персонализираните препоръки, най-ниски грешки и добра ефективност по време на обучение, особено при използване на предварително обучени embeddings. За разлика от него, моделите CBF+CF, автоенкодери и обучение с подсилване не постигат висока ефективност,

тъй като са или прекалено опростени, или прекалено сложни за спецификата на използвания набор от данни. Допълнително са направени и анализи на системата преди и след „utility scoring“ в табл. 20 и табл. 21. Основния извлечен извод е, че оценяването на полезност пренаreja препоръките, като вместо само по предсказан рейтинг ги подрежда според баланс между точност и актуалност, давайки приоритет на по-нови и релевантни взаимодействия.

Таблица 20. Най-добрите 10 препоръки преди прилагане на оценка на полезност („utility scoring“)

| Ранг | ID на статия | Реален рейтинг | Предсказан резултат | Дни след оценка от потребителя |
|------|--------------|----------------|---------------------|--------------------------------|
| 1    | 9            | 4.5            | 4.77493             | 3413.88                        |
| 2    | 8            | 5.0            | 4.717092            | 3587.75                        |
| 3    | 6            | 4.5            | 4.642761            | 1771.29                        |
| 4    | 35           | 4.5            | 4.627602            | 2197.96                        |
| 5    | 17           | 4.5            | 4.605495            | 3002.75                        |
| 6    | 3            | 4.5            | 4.59516             | 7124.75                        |
| 7    | 41           | 4.5            | 4.577928            | 560.04                         |
| 8    | 26           | 4.5            | 4.568518            | 7069.38                        |
| 9    | 46           | 4.5            | 4.562939            | 6663.63                        |
| 10   | 25           | 5.0            | 4.555358            | 7126.96                        |

Таблица 21. Най-добрите 10 препоръки след прилагане на оценка на полезност („utility scoring“)

| Ранг | ID на статия | Реален рейтинг | Предсказан резултат | Дни след оценка от потребителя |
|------|--------------|----------------|---------------------|--------------------------------|
| 1    | 9            | 4.5            | 4.77493             | 3413.88                        |
| 2    | 8            | 5.0            | 4.717092            | 3587.75                        |
| 3    | 6            | 4.5            | 4.642761            | 1771.29                        |
| 4    | 17           | 4.5            | 4.605495            | 3002.75                        |
| 5    | 25           | 5.0            | 4.555358            | 7126.96                        |
| 6    | 35           | 4.5            | 4.627602            | 2197.96                        |
| 7    | 31           | 4.5            | 4.423069            | 5007.63                        |
| 8    | 3            | 4.5            | 4.59516             | 7124.75                        |
| 9    | 26           | 4.5            | 4.568518            | 7069.38                        |
| 10   | 63           | 4.5            | 4.27723             | 3257.38                        |

#### 4.4. Оценка на ефективността на системата за препоръки „Presentation Advisor“

За оценка на ефективността на цялостната системата за препоръки са използвани следните офлайн методи: ръчно тестване със симулирани потребители и стандартни метрики (MAE, MSE,

RMSE, Precision@K, Recall@K, F1-score, Novelty, Diversity, Serendipity, Coverage). Експерименталните резултати са представени в табл. 23.

Получените резултати показват, че разработената система за препоръки постига стабилно и балансирано представяне спрямо ключовите показатели за ефективност, като значително превъзхожда базовия филтрационен подход. По отношение на Precision@10, Recall@10 и F1@10, системата достига стойности от 0.34, спрямо 0.20 при филтрацията. Това означава, че на всеки потребител се предлагат средно 3.4 релевантни артикула в топ 10 препоръки, като се поддържа добра пропорция между точността и пълнотата на препоръките. Подобреното спрямо базовия метод е значимо, което е индикатор за по-ефективно елиминиране на нерелевантно съдържание.

Таблица 23. Офлайн оценката на ефективността на системата за препоръки

| Метрика         | Стойност на препоръчителна система | Стойност без препоръчителна система (филтрация) |
|-----------------|------------------------------------|-------------------------------------------------|
| Precision@10    | 0.3400                             | 0.2                                             |
| Recall@10       | 0.3400                             | 0.2                                             |
| F1@10           | 0.3400                             | 0.2                                             |
| nDCG@10         | 0.4428                             | 0.2553                                          |
| Novelty@10      | 6.6439                             | 6.6438                                          |
| Diversity@10    | 0.5557                             | 0.5815                                          |
| Item Coverage   | 0.2700                             | 0.1                                             |
| User Coverage   | 1.0000                             | 1.0000                                          |
| Users Evaluated | 5                                  | 5                                               |

Особено показателен е резултатът за nDCG@10 (0.4428 при системата срещу 0.2553 при филтрацията), който свидетелства за значително по-добро ранжиране на релевантните артикули в горната част на списъка. Това увеличава вероятността за взаимодействие от страна на потребителя и демонстрира напредък в оптимизацията на алгоритъма за ранжиране. Метриците за новост показват, че и двата подхода успяват да предоставят съдържание, което не е прекалено популярно или добре познато на потребителите. По отношение на разнообразието, филтрацията постига малко по-висока стойност (0.5815) спрямо 0.5557 при системата, което показва, че базовият метод разпределя предложенията по-широко в тематично отношение, но за сметка на релевантността. Покритието на артикулите (Item Coverage) е значително по-високо при системата за препоръки – 27% от каталога срещу 10% при филтрацията, което показва по-широко използване на наличното съдържание и потенциал за намаляване на ефекта „filter bubble“.

След внедряване на системата в реална среда като част от следващия етап на проучването се предвижда „онлайн валидиране“ чрез прилагане на следните подходи:

1. *A/B тестване*: група А: базова версия без персонализации, група В: пълна версия с utility scoring и времеви фактори, метрики: потребителска оценка, честота на взаимодействия, реална употреба в презентации, подобрения в автоматичния анализ.
2. *Дългосрочно проследяване на умения (Longitudinal Study) на потребителите*: събиране на данни за приложени препоръки и еволюция на умения.
3. *Интерактивна обратна връзка (Feedback Loop)*: вграден интерфейс за оценка полезността и приложимостта на препоръките.
4. *Анализ на поведение на потребителите*: проследяване на продължителност на сесии, честота на използване, кликания, типове съдържание което е предпочитано и други.

Получените експериментални резултати показват, че системата за препоръки „Presentation Advisor“ демонстрира стабилни резултати в офлайн среда и е готова за онлайн оценка, чрез която да се потвърди ефекта ѝ върху реалната употреба, удовлетвореността на потребителите и дългосрочното развитие на презентационните умения.

#### 4.5. Изводи към четвърта глава

В рамките на проведените експериментални изследвания са оценени трите ключови модула на разработената интегрирана интелигентна система за препоръки „Presentation Advisor“, насочена към

подобряване на съдържанието, структурата и стиловата съгласуваност на презентации, както и към персонализирано подпомагане на потребителите при тяхното създаване и оптимизация. Експерименталните изследвания обхващат трите основни компонента на системата – модул за проверка на съдържание и структура, модул за стилова съгласуваност и препоръчващ модул, всеки от които е оценен с използване на специализирани метрики и реални данни.

Резултатите от модула за проверка на съдържанието и структурата на презентациите показват висока степен на точност при идентифициране на често срещани грешки в презентациите с използване на предложеният в дисертационния труд модел. Постигнати са високи стойности за метриците за оценка на класификация (точност, прецизност, пълнота), което демонстрира, че моделът е способен да се адаптират към различни типове презентационно съдържание.

При модула за стилова съгласуваност, предложеният модел успява да разграничи стилово несъвместими двойки, което е особено важно за визуалната цялост на професионални презентации.

Препоръчителният модул демонстрира висока персонализация, особено след включване на контекстуални и времеви фактори. Прилагането на оценка на полезност и времево-чувствителен механизъм за внимание значително подобрява ранга на препоръките и повишава тяхната уместност спрямо нуждите на потребителя.

Цялостната оценка на системата „Presentation Advisor“, извършена чрез комбинация от офлайн метрики показва, че системата е не само функционално пълноценна, но и потенциално приложима в реална среда. В допълнение е разработена схема за поэтапна онлайн оценка (включително A/B тестове, дългосрочно проследяване на умения интерактивна обратна връзка, анализ на поведение), осигуряваща надеждно измерване на дългосрочния ефект върху презентационните умения на потребителите при използване на системата за препоръки.

## ЗАКЛЮЧЕНИЕ

В рамките на дисертационния труд е проектирана и емпирично валидирана персонализирана препоръчителна система, основана на многомодулна архитектура, която обединява няколко специализирани дълбоки модела – всеки със строго дефинирана аналитична роля. Интеграцията на тези компоненти позволява системата да функционира едновременно на три взаимодопълващи се равнища:

- „Виждане“ на слайда в детайл – конволюционен клон, обучен върху 6 000 анотирани изображения, който локализира визуални несъответствия (четливост, размер и контраст на шрифтове, неподходящи графики и др.);
- „Усещане“ на глобалната стилова кохерентност – сиамски модул, сравняващ латентните представяния на последователни слайдове и сигнализира за отклонения в цветови палитри, визуална йерархия и оформление;
- „Разбиране“ на поведенческия контекст на потребителя – хибриден препоръчващ слой, който в реално време интегрира личната история на работа, моментното ниво на ангажираност и времеви контекст, за да предложи оптимално полезни и образователни препоръки.

Многомодулният дизайн осигурява висока интерпретируемост и лесно надграждане: всеки клон може да се актуализира или заменя независимо, без да се нарушава работата на останалите. Това улеснява внедряването на допълнителни аналитични модули (напр. за съответствие между текст и визуално съдържание) с минимални промени в ядрото на системата. Комбинирането на отделните модели подпомага персонализацията в реално време. Докато визуалните клонове оценяват статичните характеристики на слайдовете, поведенческият модул динамично отчита индивидуалните особености и цели на презентатора. Така генерираните препоръки са не просто „коректни“ по някаква обща метрика, а максимално релевантни и обучаващи за конкретния потребител в конкретния момент. В резултат Presentation Advisor се превръща от пасивен инструмент за детекция на грешки в активен дигитален ментор, който едновременно диагностицира, обяснява и обучава. Доказано е, че многомоделната, задачно-ориентирана интеграция е ключът към следващото поколение интелигентни препоръчителни системи – системи, които не само филтрират информация, а добавят реална стойност за крайния потребител.

Системата Presentation Advisor постигна убедителни резултати в контролирана тестова среда и е готова за внедряване в учебни и корпоративни контексти. В образователна среда тя може да подпомага студентите при самостоятелната подготовка на презентации, като автоматично открива проблеми с четливостта, размера на текста и визуалната консистентност, и предлага точни, визуално обяснени корекции. В корпоративния сектор системата може да служи като цифров

наставник, който не просто маркира пропуски, а обучава потребителите в добри практики на презентационния дизайн, насърчавайки повишена визуална грамотност и по-висока ефективност на предадената чрез презентацията информация.

Предстоящата работа по системата очертава няколко ключови предизвикателства и направления за развитие. На първо място стои мащабното ѝ изпитване с реални потребители в продукционна среда, което ще позволи дългосрочно да се проследи как препоръките влияят върху презентационните умения на автори от различни професионални домейни. Паралелно с това ще се разширява самият модел – ще се добавят специализирани подмодули за семантично съответствие между текст и слайд, както и за автоматична оценка на сложността в графики и инфографики, допълнени от нови класификатори за специфични дизайнерски проблеми. Важна стъпка е и интегрирането на големи езикови модели: планира се използването им за генериране на обяснения на естествен език, съобразени с експертното ниво на потребителя, за препоръки при преформулиране на заглавия и списъчни елементи с оглед на краткост и яснота, както и за автоматично предлагане на по-добра подредба на елементите в слайдовете. Цялата архитектура ще бъде подготвена за непрекъснато и подобряващо обучение чрез техники от reinforcement learning и continual learning, така че се очаква системата да се адаптира динамично към развиващия се стил и опит на всеки отделен презентатор. Обединяването на този многомоделен аналитичен подход с утвърдените практики за индустриално внедряване показва, че персонализираните препоръки могат да надскочат ролята на коректив и да се превърнат в активен обучителен инструмент – решаваща стъпка към ново поколение интелигентни асистенти за създаване на презентации.

## **ПРИНОСИ ПО ДИСЕРТАЦИОННИЯ ТРУД**

### ***Научни приноси:***

- Предложени са два подхода за мулти-етикетна класификация на изображения за стилова съгласуваност на презентации с адаптиране на алгоритъм чрез единен модел и с трансформация на проблема чрез независими бинарни класификатори, който осигурява по-добро преобучаване, мащабируемост и интерпретируемост и са по-ефективни в условия на небалансирани класове;
- Предложена е дву-клонова архитектура на конволюционна невронна мрежа за оценка на стилова съгласуваност на слайдове в презентация, която извлича паралелно характеристики от два слайда, използва конкатенация и допълнителни плътни слоеве за по-прецизно моделиране на взаимовръзките и постига по-висока точност спрямо класически сямски невронни мрежи и традиционни метрики като SSIM (Structural Similarity Index Measure) или ORB (Oriented FAST and Rotated BRIEF);

### ***Научно-приложни приноси:***

- Предложена е методология за персонализирани препоръки за презентации за разработване на персонализирана система за препоръки, ориентирана към подобряване на презентационните умения и визуалната съгласуваност на слайдове чрез методи на машинно и дълбоко обучение;
- Предложен е подход за формиране и структуриране на комплексен набор от данни, който интегрира различни източници и типове информация, необходими за изграждане на система за персонализирани препоръки и включва обединяване на потребителски профили, анотирани презентации, учебни материали и автогенерирани данни;
- Предложена е концептуална рамка на интегрирана система, комбинираща поведенчески, съдържателен и визуален анализ на презентации, насочена към дългосрочно подобряване на презентационните умения на потребители, изградена от три основни модула: Pstyle Checker за автоматизирана оценка на стилова кохерентност, Presentation Checker за откриване на проблеми и грешки в слайдове, PAdvisor за генериране на персонализирани обучителни препоръки;

- Предложена е концептуална архитектура на препоръчителна система за презентации съчетаваща три основни модула: Presentation Checker Service за мулти-етикетна класификация на съдържателни и структурни характеристики, PStyle Checker Service за оценка на стилова съгласуваност чрез иновативна дву-клонова конволюционна невронна мрежа и Recommendation Engine с хибриден препоръчващ модел, комбиниращ филтриране по съдържание и машинно обучение.
- Предложена е методология за събиране, аотиране и предобработка на синтетични и реални данни за структурно-съдържателен анализ и стилова съгласуваност на презентации в условия на ограничена наличност на публични множества данни;

■ **Приложни приноси:**

- Създадени са два специализирани набора от данни: ръчно аотиран корпус от 912 реални слайдове за структурно-съдържателен анализ и синтетичен набор от 6000 изображения за стилова съгласуваност с контрол върху визуални параметри;
- Разработен е прототип на система за препоръки за интелигентно и персонализирано създаване на презентации „Presentation Advisor“ с ефективна интеграция между модулна архитектура, съвременни модели за машинно обучение и реални потребителски нужди, която предлага персонализирани препоръки за подобрене на съдържанието и дизайна на презентации, реално-времеви анализ с обратна връзка от потребителите, възможност за адаптация към нови визуални стандарти, езици и типове презентации;
- Предложена е цялостна MLOps инфраструктура за внедряване на системата за препоръки за интелигентно и персонализирано създаване на презентации, която включва контейнеризация, версионизиране, автоматично деплойване и мониторинг, стратегии за инкрементално и активно обучение с включени механизми срещу катастрофално забравяне и събиране на потребителска обратна връзка за затворен цикъл на подобрене

## СПИСЪК НА ПУБЛИКАЦИИ ПО ДИСЕРТАЦИОННИЯ ТРУД

1. **Vlahova, M.**, “Recommender Systems: Types, Advantages, Disadvantages and Evaluation Techniques,” *Journal Computer and Communications Engineering*, vol. 15, no. 2, 2021, pp. 3-11, ISSN 1314-2291.
2. **Vlahova, M.**, “A Scientific Approach to Problem Solving”, Proc. of 1<sup>st</sup> EUt+ ELaRa Conference “Technologies and Techniques to Support Sustainable Education in the Academic Sphere”, 14-15 December 2022, TU-Sofia
3. **Vlahova, M.**, Lazarova, M., “Collecting a Custom Database for Image Classification in Recommender Systems,” Proc. of 10<sup>th</sup> IEEE International Scientific Conference on Computer Science (COMSCI), Sofia, Bulgaria, 30 May 2022 – 2 June 2022, <https://doi.org/10.1109/COMSCI55378.2022.9912591> (индексирана в Scopus)
4. **Vlahova-Takova, M.**, Lazarova, M., “CNN Based Multi-Label Image Classification for Presentation Recommender System,” *International Journal on Information Technologies & Security*, vol. 16, no. 4, pp. 73-84, 2024, <https://doi.org/10.59035/PUYE7368> (индексирана в Scopus и WoS, IF 1.0, Q4)
5. **Vlahova-Takova, M.**, Lazarova, M., “Dual-Branch Convolutional Neural Network for Image Comparison in Presentation Style Coherence,” *Journal Engineering, Technology & Applied Science Research*, vol. 15, no. 2, pp. 21719–21727, 2025, <https://doi.org/10.48084/etasr.9571> (индексирана в Scopus)
6. **Vlahova-Takova, M.**, Lazarova, M., “A Recommender System Model for Presentation Advisor Application Based on Multi-Tower Neural Network and Utility-Based Scoring,” *Journal MPDI Electronics*, vol. 14, 2025, <https://doi.org/10.3390/electronics14132528> (индексирана в Scopus и WoS, IF 2.6, Q1)

## SUMMARY

### METHODS AND ALGORITHMS FOR PERSONALIZED RECOMMENDATION SYSTEMS

**Maria Petrova Vlahova-Takova**

The research in the dissertation is aimed to design, develop, and validate a personalized recommendation system that uses modern approaches from the field of machine and deep learning to analyze and improve presentations through intelligent, explainable, and educational feedback.

A methodology for personalized presentation recommendations is proposed for the development of a personalized recommendation system aimed at improving presentation skills and visual consistency of slides through machine learning and deep learning methods. A conceptual framework for an integrated system is proposed, combining behavioral, content, and visual analysis of presentations, aimed at long-term improvement of users' presentation skills, consisting of three main modules: Pstyle Checker for automated assessment of stylistic coherence, Presentation Checker for detecting problems and errors in slides, PAdvisor for generating personalized training recommendations.

Two approaches are proposed for multi-label image classification for stylistic consistency in presentations: (1) an approach with algorithm adaptation through a unified model and (2) an approach with problem transformation through independent binary classifiers, which provides better retraining, scalability, and interpretability and is more effective in conditions of unbalanced classes. A two-branch convolutional neural network architecture is proposed for assessing the stylistic consistency of slides in a presentation, which extracts characteristics from two slides in parallel, uses concatenation and additional dense layers for more precise modeling of interrelationships, and achieves higher accuracy compared to classical Siamese neural networks and traditional metrics such as SSIM or ORB.

An approach is proposed for forming and structuring a complex data set that integrates various sources and types of information necessary for building a personalized recommendation system and includes the merging of user profiles, annotated presentations, training materials, and auto-generated data. Two specialized datasets are created: a manually annotated corpus of 912 real slides for structural and content analysis and a synthetic dataset of 6 000 images for stylistic consistency with control over visual parameters.

A prototype recommendation system for intelligent and personalized presentation creation “Presentation Advisor” is developed with effective integration between modular architecture, modern machine learning models, and real user needs, which offers personalized recommendations for improving the content and design of presentations, real-time analysis with user feedback, and the ability to adapt to new visual standards, languages, and presentation types. A comprehensive MLOps infrastructure is proposed for implementing the recommendation system for intelligent and personalized presentation creation, which includes containerization, versioning, automatic deployment and monitoring, strategies for incremental and active learning with built-in mechanisms against catastrophic forgetting, and collection of user feedback for a closed-loop improvement cycle.

Based on multimodal analysis and combining visual and stylistic characteristics of slides, the prototype system aims not only to assist the user in identifying and correcting specific problems, but also to encourage the acquisition of long-term communication and visual skills. The development of such a tool meets contemporary requirements for adaptive, technological, and educational support for users and represents a significant contribution at the intersection of artificial intelligence, visual literacy, and digital learning. The dissertation also makes a methodological contribution to the creation of intelligent learning technologies with a high degree of personalization and explainability.

The experimental results demonstrate convincing results for usage of the “Presentation Advisor” in a controlled test environment that also reveal its readiness for implementation in educational and corporate contexts. In an educational setting, the system can assist students in preparing presentations by automatically detecting problems with readability, text size, and visual consistency, and offering precise, visually explained corrections. In the corporate sector, the system can serve as a digital mentor that not only flags gaps but also trains users in good presentation design practices, promoting increased visual literacy and greater effectiveness of the information conveyed through the presentation.