

ИНОВАТИВНИ ИНФОРМАЦИОННИ ТЕХНОЛОГИИ ЗА ДИГИТАЛИЗАЦИЯ НА ИКОНОМИКАТА

Сборник с доклади



УНИВЕРСИТЕТ ЗА НАЦИОНАЛНО И СВЕТОВНО СТОПАНСТВО
КАТЕДРА „ИНФОРМАЦИОННИ ТЕХНОЛОГИИ И КОМУНИКАЦИИ“

ИНОВАТИВНИ ИНФОРМАЦИОННИ ТЕХНОЛОГИИ ЗА ДИГИТАЛИЗАЦИЯ НА ИКОНОМИКАТА

Сборник доклади от научна конференция
на катедра „Информационни технологии и комуникации“,
УНСС, 15 октомври 2025 г.

София, 2025

Научна редакция и съставители:

Доц. д-р Ваня Лазарова

Гл. ас. д-р Ивона Велкова

Програмен комитет:

Проф. д.ик.н. Валентин Кисимов

Проф. д-р Любен Боянов

Проф. д-р Камелия Стефанова

Доц. д-р Александрина Мурджева

Доц. д-р Пламен Милев

Доц. д-р Митко Радоев

Доц. д-р Моника Цанева

Гл. ас. д-р Мария Мързованова

Гл. ас. д-р Мариана Ковачева

Организационен комитет:

Председател:

Гл. ас. д-р Ивона Велкова

Членове:

Доц. д-р Ваня Лазарова

Доц. д-р Дорина Кабакчиева

Гл. ас. д-р Станимира Йорданова

Гл. ас. д-р Иван Белев

Гл. ас. д-р Гено Стефанов

Гл. ас. д-р Веска Михова

Гл. ас. д-р Явор Табов

ISSN 3033-0432 (print)

ISSN 3033-0467 (online)

Академично издателство УНСС

Печатница УНСС

UNIVERSITY OF NATIONAL AND WORLD ECONOMY
DEPARTMENT "INFORMATION TECHNOLOGIES AND COMMUNICATIONS"

**INNOVATIVE
INFORMATION
TECHNOLOGIES
FOR ECONOMY
DIGITALIZATION
(I I T E D)**

Conference Proceedings from the Scientific Conference
of the Department "Information Technologies and Communications",
UNWE, 15 October, 2025

Sofia, 2025

СЪДЪРЖАНИЕ

Рисковете за дигитализация, причинени от човешки фактори в академичната среда	9
Любен Боянов	9
The Future of Data: Utopia or Dystopia?	20
Assof Prof. Aleksandrina Murdzheva,	20
Nikolay Vlaychev	20
Съвременни аспекти на приложения на изкуствен интелект в Бизнес Информационните Системи и в частност ERP системите	29
Димитър Димитров	29
Blockchain and Artificial Intelligence for Intelligent and Secure Financial Transactions: An Implementation Perspective with Python	37
Vasil Spasov,	37
AI in Offensive Cybersecurity: Autonomous TTP Emulation and Scalable Continuous Testing	42
Сали Салиев	42
Speech-to-Speech Artificial Intelligence as a New Stage in Automated Communication	54
Radoslav Dodnikov	54
Киберсигурност като инструмент за защита на цифровите активи	63
Андрей Вишневский	63
Exploring Seasonal Trends in Bulgarian Media: Insights from a Data-Driven Analysis of Web Publications	68
Plamen Milev	68
Интернет на обектите и дигитални двойници в обществените поръчки: модел за умен контрол и превенция на корупционни практики	77
Теодор Буров	77
Невидимите бариери във визуализацията на данни: как достъпността оформя потребителското изживяване	86
Лиляна Петрова Сапунджиева,	86
Приложни аспекти на интеграция между системи за управление на съдържание в организацията (ECM) и системи за управление на бизнес процеси (BPM)	95
Иван Белев	95
Теоретични аспекти на статистическия анализ на неструктурирани данни от чатботове за усъвършенстване на клиентското преживяване	100
Биляна Голешова	100

Semantic Search based on Embedding Models	111
Vanya Lazarova	111
Съвременни аналитични подходи за оценка на резултатите от дигитални маркетингови кампании	116
Гергана Тодорова,	116
Интегриране на изкуствен интелект в управлението на човешките ресурси: резултати от пилотно изследване в български организации	123
Милена Миленкова,	123
Приложение на изкуствения интелект в образованието - за и против.....	130
Наталия Маринова,	130
Технологични рискове от дигитализацията в неакадемична среда	137
Мариана Ковачева	137
Technological Features of Modern Video-Sharing Platforms	145
Yavor Tabov,	145
Подобряване на обучението и повишаване на квалификацията чрез симулации с разширена реалност	151
Kaloyan Dimitrov	151
Отворените данни в контекста на устойчивите хранителни системи.....	158
Асен Божилов.....	158
Методи за автоматизация на бизнес анализи чрез интегриране на нееднородни данни ..	166
Теодор Тодоров,	166
Основни аспекти на въвеждането на цифрово евро	173
Аглика Кънева.....	173
Модели и стратегии за дигитална трансформация: новото нормално	182
Светла Ценова	182
Assessment and management of technological risks in the digitalization process of academic institutions	189
Ivona Velkova,	189
Software tools for online analytical processing- from data to strategic decisions in the digital economy	200
Marinela Todorova.....	200
Интелигентни виртуални асистенти като средство за подобряване на клиентското изживяване.....	208
Виктория Габровска.....	208

Алгоритмичен маркетинг, дигитални двойници и манипулация на потребителското поведение: етика и рискове в дигиталната икономика	214
Василена Василева.....	214
Изкуствен интелект в маркетинговото планиране: възможности за подобряване на ефективността	229
Гергана Тодорова,	229
Прилагане на геймификацията в електронното обучение като подход за повишаване на мотивацията.....	236
Кремена Маринова-Костова,	236
Предизвикателства пред професионалното образование и обучение в България и дигитални инструменти за справяне	243
Лора Първанова.....	243
Оптимизиране на автоматизирано тестване на бизнес логика в .NET чрез AI Business logic unit testing optimization in .NET using AI.....	250
Maria Marzovanova,	250
Veska Mihova	250
Типове данни в Управлението на бизнес процеси (BPM) – класификация, сравнение и аналитичен потенциал	257
Гено Стефанов,.....	257
Artificial Intelligence and Tax Administration in Bulgaria	262
Georgi Emilov Hristov.....	262
Използване на машинно обучение за автоматизация при анализ на нехармонизирани фирмени данни	266
Теодор Тодоров,.....	266
Digital Twins in Healthcare: Simulation, Data, and the Future of Personalized Medicine	274
Lyuben Zyumbilski.....	274
Artificial Intelligence in Medical Diagnostics: Algorithms, Data, and Challenges in Practical Implementation	279
Lyuben Zyumbilski.....	279
Методологични подходи в мултимодалния sentiment анализ	284
Станимира Йорданова	284

РИСКОВЕТЕ ЗА ДИГИТАЛИЗАЦИЯ, ПРИЧИНЕНИ ОТ ЧОВЕШКИ ФАКТОРИ В АКАДЕМИЧНАТА СРЕДА

The risks for digitalization by human factors in academic environment

Любен Боянов¹,

e-mail: lboyanov@unwe.bg¹

Абстракт

През последните три десетилетия дигитализацията постепенно, се превърна във първостепенен фактор и основна сила в образователната среда. Тази нова парадигма коренно промени начина, по който знанието се представя, придобива, разпространява и оценява. Нейното значение произтича от възможностите на дигиталния свят да преодолява много от ограниченията на традиционните подходи на аналоговия свят, а в случая с образованието - да подготви хората за непрекъснато развиващия се и усъвършенстващ се нов свят, който в момента е преобладаващо дигитален. Същевременно дигиталната трансформация (ДТ) носи няколко риска. Много от тях са свързани с човешки фактори, които могат или да възпрепятстват успешното внедряване на дигитализацията и качеството на образованието, или да доведат до непредвидени негативни последици. Настоящият доклад разглежда свързаните с човека рискове от дигитализацията, като представя описание на отделните рискове, категоризира ги и дава примери. Материалът завършва с кратко обсъждане и препоръки по разглежданите теми и проблеми.

Abstract

During the last three decades, digitalization has gradually but unstoppable became an indispensable factor and major force in the educational environment. This new paradigm fundamentally reshaped the way how knowledge is presented, acquired, disseminated, and assessed. Its importance stems from the possibilities of the digital world to overcome many of the limitations of traditional approaches of the analogue world and in the case of education - to prepare individuals for the ever progressing and enhancing new world, that is overwhelmingly digital at present. At the same time the Digital transformation (DT) brings a few risks. Many are related to human factors that can either hinder the successful implementation of digitalization and quality of education or lead to unintended negative consequences. The present paper looks at the human-related risks of digitalization, presents a description of each risk, categorizes them and presents examples. At the end, there is a short discussion and recommendations on the topics.

Ключови думи: Дигитална трансформация, образование, рискове, човешки фактори

Keywords: Digital transformation, education, risks, human factors

JEL: I2, I23, I24, O33

Въведение

Дигитализацията на образованието промени и изпрати обучението отвъд традиционните класни стаи чрез модерните технологии. Онлайн платформите (напр. Coursera, Udemy, EdX) предоставят гъвкав достъп до курсове и ресурси, докато системите за управление на обучението (Google Classroom, Moodle, Canvas) помагат на преподавателите да управляват и проследяват напредъка на учениците. Електронните учебници и дигиталните материали намаляват разходите и улесняват актуализациите, а приложения за обучение като Quizlet и Duolingo правят обучението по-ангажиращо чрез игри/геймификация и персонализация. Инструментите за

¹ Професор, д-р, Катедра информационни технологии и комуникации, Факултет по приложна информатика и статистика, Университет за национално и световно стопанство, ORCID 0009-0006-2292-0619.

едновременна работа (Zoom, Teams) насърчават взаимодействието, докато Виртуалната реалност (VR)/Добавената реалност (AR) предлагат визуални и итеративни подходи за представяне и разбиране на по-сложни концепции. Изкуственият интелект (ИИ), в това число и машинното обучение персонализират учебните материали и подобряват качеството на преподаване. Заедно тези иновации създават по-гъвкави, ефективни и интерактивни учебни преживявания.

Дигиталната трансформация (ДТ) в образованието предоставя множество ползи. Тя подобрява ефективността на преподаването и ученето чрез инструменти за електронно обучение, изкуствен интелект и дигитални ресурси, като същевременно създава по-ангажиращи и персонализирани учебни преживявания. ДТ разширява достъпа до образование, като премахва географските и времевите бариери, осигурявайки по-голяма справедливост според различните възможности. Друг важен фактор е, че разходите и ресурсите се оптимизират чрез дигитални материали и автоматизирани системи за управление, докато интерактивните инструменти насърчават взаимодействието и активното участие на преподаващи и учащи се. Освен това, ДТ развива основни дигитални умения за учителите и учениците, укрепва връзките между всички участващите в този процес страни и използва анализ на данни, за да подобри както качеството на преподаване, така и управлението на образованието. Съвременната глобалната дигитална среда с инструменти, достъпни по/от целия свят, позволява съвместна работа, сътрудничество и обучение с колеги и експерти по целия свят, насърчавайки межкултурното разбирателство и различни перспективи.

Въпреки че споменатите по-горе предимства на дигиталните технологии подобряват ефективността и достъпността в образованието, те също така крият рискове. Същите могат да бъдат разделени на две категории - предизвикани от технологиите и предизвикани от човека. Тук ще разгледаме рисковете, свързани с човешкия фактор, които могат да възпрепятстват ДТ и да посочат редица нейни уязвимости. Човешкият фактор се разглежда като критичен фактор за успеха или неуспеха на ДТ. Съществуват няколко подфактори, от които може да се отличат съпротивата, разликите в уменията, опасенията за поверителност, етичните въпроси и неравенството в достъпа до ресурси, които могат да забавят напредъка и да подкопаят ефективността на гореспоменатите ползи. Ако не се управляват по подходящ начин, тези рискове биха могли да поставят под въпрос смисъла на ползите от дигитализацията и нейните предимства относно високата и ефективност, гъвкавост и устойчивост. Проучванията показват, че дигиталната трансформация надхвърля формалното прилагане и приемане на нови технологии. Тя се нуждае от иновативни педагогически подходи, всеобхватна политика, регулаторни рамки и съблюдаването други норми и парадигми [1], [2]. Този доклад разглежда свързаните с хората рискове от дигитализацията в академичните среди, като се фокусира върху съпротивата срещу промяната, по-голямото натоварване, страха от загуба на работа, липсата на стимули и признание, пропуските в дигиталната грамотност, опасенията за поверителност, етичните проблеми и предизвикателствата пред равенството и възможността за достъп от всекиго. В този материал не разглеждаме технологичните фактори, касаещи рисковете за студенти, преподаватели и персонал във висшите учебни заведения.

Рискове за дигитализацията от човешки фактори

Съпротива срещу промяната

Процесът на дигитализация в академичните среди, макар и да предлага огромен потенциал, е изпълнен с рискове, свързани с човешките фактори, които често създават съпротива срещу промяната. Тези рискове произтичат от различни причини, включително дълбоко вкоренени културни норми, липса на институционална подкрепа и индивидуални пропуски в развитието на човешките умения. Преподаватели, студенти и администрация може да се противопоставят на усвояването на новите технологии поради навиците за работа с традиционните методи, страх от загуба на работа или скептицизъм относно ефикасността на дигиталните инструменти. Огромна риск, свързан с хората, е съпротивата им срещу промяната – както от страна на преподавателите, така и от страна на администрацията. Академичните среди, особено тези, които са работили в

продължение на десетилетия с традиционните до дигиталната ера системи, не винаги са склонни да приемат новите технологии. Тази съпротива може да бъде обоснована от няколко фактора:

- **Притеснения с работното натоварване:** Много преподаватели смятат, че новите технологии са просто допълнително бреме върху вече натоварените им графици, а не инструмент за пестене на време. Те се безпокоят от допълнителните усилия, необходими за изучаване на нови системи, конвертиране на материали в дигитална форма и работата с онлайн платформи. Проучване от 2016 г. на Johnson et al. отбелязва, че учителите се съпротивляват на приемането на технологиите, когато те изискват „двойна иновация“, увеличавайки работното си натоварване поради времето, прекарано в изучаване на инструменти и адаптиране на учебните планове [3]. В друго проучване от 2024 г. от WGU Labs, 41% от преподавателите съобщават за „прегаряне“ [4]. Други изследвания показват, че голямото работно натоварване „... е значително предизвикателство във висшето образование, което се отразява на преподавателите и студентите.“ [5].
- **Страх от това да станеш излишен:** Някои преподаватели се страхуват, че ИИ и дигиталните инструменти ще направят уменията им излишни или твърде остарели или ще обезценят човекоцентричния характер на преподаването. Това може да създаде (наред с други неща) чувство за несигурност на работното място и да доведе до съпротива към или даже отхвърляне на технологиите [6].
- **Липса на стимули и признание:** В някои случаи не е имало споразумение за целите и резултатите от ДТ между преподавателите и институциите, а по-скоро текстове за задължението за преподаване. И тъй като дигиталното обучение не се счита за еквивалентно на обучението в класната стая, преподавателите нямат сериозен стимул да го възприемат и прилагат [7].
- **Дигиталното разделение:** Този риск съдържа голяма доза от човешкия фактор. В този случай не става въпрос само за достъпа до технологии, но и за разликата в **дигиталната грамотност и умения** както сред студентите, така и сред преподавателите. Студентите от **по-ниските социално-икономически среди** в редица случаи нямат надежден (или с необходимата бързина) достъп до Интернет или не разполагат с необходимите устройства, докато някои преподаватели може да не притежават умения за ефективно използване на дигиталните инструменти. Това разделение може да влоши съществуващите неравенства в образователната система и висше учебно заведение, оставяйки някои студенти и преподаватели в неизгодно положение. Изследванията показват, че съпротивата произтича от липса на време, дискомфорт, проблеми със зрението и недостатъчни умения в комуникацията, типични за дигиталното разделение. Учителите-ветерани показват по-голяма съпротива към новите технологии поради възрастта и опита си в недигиталните технологии [8], [9]. Дигиталното разделение е налице и в общности със социални неравенства. Последните включват **неравенства между половете, ниски доходи, напреднала възраст и при хора с увреждания**. Всички те засилват рисковете, свързани с изолация, докато **културни фактори като недоверие към дигиталните системи, езикови бариери и ограничена дигитална грамотност** допълнително възпрепятстват пълноценното ангажиране в процеса на трансформация [10].

Характерен пример може да се открие в проучването на EDUCASE от 2023 г. [11], което показва, че 53% от преподавателите са предпочели присъственото обучение пред хибридните модели поради познаваемостта на този модел обучени. Въпреки че по време на пандемията предпочитанията, поради здравословните обстоятелства, бяха за онлайн лекции и преподаване (около 73%), най-предпочитаният начин на преподаване в днешно време присъственият, а най-малко предпочитаният е онлайн. Лесните за използване съвременните инструменти с генеративен ИИ (GenAI) не гарантират непременно тяхното безкритично възприемане. Преподавателите често се въздържат от ефективното интегриране на тези инструменти в своите преподавателски практики [12]. Съпротивата на педагозите към инструментите с ИИ често произтича от страха от загуба на контрол върху учебния процес, опасения относно загуба на работното място и липса на увереност в способността им да използват ефективно технологията. Тази съпротива може да се прояви като нежелание за използване на нови инструменти и методи, притеснение относно това как ИИ ще повлияе на ангажираността и поведението на студентите,

както и етични опасения относно поверителността на данните и потенциалните пристрастия в системите с ИИ [13].

Като цяло, в образователните среди съществува силно културно и академично предпочитание за взаимодействие лице в лице по време на учебния процес. В много случаи липсва доверие в надеждността на технологиите или в ползите и стимулите за използването им. Съществуват също така тревожност, притеснение за по-голямо работно натоварване и страх от да стенеш излишен и да загубиш работата си при използването на нови и зараждащи се дигитални технологии.

Недостатъчна дигитална грамотност

Недостатъчната дигитална грамотност представлява голям риск при дигитализацията на образованието, допринасяйки за съпротива, неефективност и неравнопоставеност при приемането на дигитални инструменти от типа на системи за управление на обучението (LMS), чатботове, управлявани от изкуствен интелект, и други дигитални платформи. Липсата на умения за ефективно използване на дигиталните инструменти, от базовата навигация/ориентиране в LMS до сложни приложения с ИИ, е сериозна причина за съпротива срещу ДТ. Първата и най-важна сред двадесетте бариери, идентифицирани в проучване на проблемите на ДТ във висшите учебни заведения (ВУЗ), е липсата на дигитална грамотност [14].

Липса на обучение и подкрепа. Университетите често внедряват нови системи, без да осигуряват адекватно, непрекъснато обучение или техническа поддръжка. Тази липса на институционална подкрепа може да доведе до сериозно разочарование и да засили убеждението, че технологията носи повече проблеми, отколкото ползи. Различни автори сочат значението на обучението и познаването на новите дигитални технологии от страна на преподавателите, както и на подкрепата от ръководството на учебното заведение [14], [15]. Изследванията посочват като бариери недостатъчните компетенции сред академичния персонал в областта на Информационните Технологии (ИТ), ограничени дигитални умения и знания както сред преподавателите, така и сред административния персонал, а също така и недостатъчна способност на студентите, преподавателите и персонала ефективно да използват съществуващите дигитални услуги [14], [16], [17]. Липсата на обучение е сред важните фактори за работещите, която проявява съпротива под различни форми [18]. Гореспоменатото проучване, проведено през 2024 г. от WGU Labs [4], показва, че 20% от анкетираните студенти се затрудняват в усвояването на новите технологии, докато 8,8% показват липса на увереност в адаптирането си към тях.

Влиянието на недостатъчната дигитална грамотност и липсата на обучение води до неефективно използване на платформите, намалява производителността, увеличава грешките, затруднява комуникацията и ограничава работата в дигиталния свят. Този фактор също така разширява разликата между технологично грамотните/можещи и по-малко квалифицираните и обучени потребители. Изследванията показват, че е необходима по-добра подкрепа както от институциите, така и от квалифицирания в техническо отношение персонал към останалите членове на образователната общност, за да се максимизират ползите от дигиталните платформи, използвани във висшите учебни заведения [19]. Някои изследвания показват, че редица студенти възприемат онлайн обучението и оценяването като по-изискващо, тъй като трябва да инвестират повече време, за да се запознаят с дигиталните платформи за изпити [20]. Преподавателите също се сблъскват с определени предизвикателства при управлението на технологиите във виртуалните курсове, до голяма степен поради недостатъчно предварително обучение в използваните образователните технологии [21]. Същото изследване предполага, че сертифицирането (официалната подготовка) в тази област може значително да намали трудностите, с които се сблъскват преподавателите, когато използват технологии за онлайн обучение.

Недостатъчната дигитална грамотност може да доведе в някои случаи до:

- Недостатъчно професионално развитие, когато университетите често не успяват да осигурят адекватно обучение на преподавателите по отношение на дигитални инструменти като LMS или чатботове, което засилва съпротивата срещу ДТ и пречи на ефективната интеграция.

- Възрастови технологични диспропорции, когато по-възрастните преподаватели и администратори, с по-малко опит в дигиталните технологии, проявяват по-голяма съпротива в сравнение с по-младите, което допринася за не равномерно възприемане и усвояване.
- Погрешно схващане за дигиталните умения на студентите (че ги имат), на базата на допускането, че последните по подразбиране притежават много високи дигитални умения. Това в много случаи е неправилно, което води до предизвикателства при ефективното използване на технологиите в академичната среда.

Проблеми, свързани с поверителността и сигурността

Друг значителен риск, относно хората и дигитализацията на образованието е свързан с **нарасналият обем на събирани и обработвани лични и академични данни** в дигиталните среди. Това води до значителни проблеми, свързани с поверителността и етиката, за студентите, преподавателите и изследователите. Наличността и потенциалното некоректно използване на лични или академични данни в дигиталните платформи често тревожи участниците относно нарушения на данните, наблюдение и злоупотреба. Редица проучвания показват, че студентите са загрижени за въпросите, свързани с поверителността и сигурността на дигиталните инструменти, като онлайн видеоконферентни и чат инструменти [21] [22]. Изследванията сочат, че неудобството, свързано с практиките за обработка на данни, произтича от рискове, свързани с хора, като неразрешени аудио и видео записи на други студенти, администратори или преподаватели. Страхове водят до нежелание да се използват дигиталните инструменти, което намалява приемането им (например избягване на електронни портфолия), страх от кражба на самоличност или разкриване на чувствителни данни от изследвания и потенциални правни и репутационни щети за институциите (например глоби по GDPR).

Преди няколко години се случи нарушаване на сигурността на данните в платформата Zoom, при който бяха загубени над 500 милиона потребителски имена и пароли от цялата база потребители на платформата [23]. Това нарушение на поверителността от хакери по време на виртуални срещи доведе до **изтичане на програмен код и друга чувствителна информация**, с което се увеличиха страховете и опасенията на потребителите за неприкосновеността на личния живот на академичната общност.

Хората искат правата и свободите им в дигиталния свят да бъдат гарантирани, независимо от ролята им - студенти, учители, членове на семейството или административен персонал. Разглеждането на теми от многобройни научни статии, като защита на личния живот в онлайн среда, право на дигитална сигурност (или киберсигурност) и право на дигитално образование, сочи необходимостта да се положат усилия към **по-обстойно и задълбочено образование на хората, тъй като юридическите уредби на правата и отговорностите**, макар и необходима, сама по себе си не е достатъчна [24].

Проблеми, свързани с поверителността, възникват и при използването на инструменти за откриване на използването на ИИ. Тези системи често се внедряват **без изричното и информирано съгласие на учениците**, което поражда сериозни опасения относно поверителността. Липсата на прозрачност при внедряването на такива инструменти застрашава правата на учениците на защита на данните и информирано участие в образователния процес [25].

Всички тези опасения се дължат главно на ниската осведоменост на потребителите относно практиките за защита на личните данни. Друго притеснение се дължи на факта, че по-голямата част от образователните софтуерни приложения събират лична информация и метаданни от учениците. Освен това тези програми все по-често се управляват не от самите образователни институции, а от външни доставчици [26].

Етични предизвикателства

Масовото въвеждане и интегриране на дигитални инструменти, особено ИИ, поставя под въпрос традиционните представи за академична честност и същността на ученето. Този процес повдига няколко етични въпроса, като академична интегритет, прекомерна зависимост от технологиите, западане на критичното мислене, пристрастност в програмите за ИИ и др. Въпреки че

инструментите за ИИ, като детектори за плагиатство и автоматизирани системи за оценяване, допринасят за укрепване на интегритета и ефективността на академичните изследвания, тяхната способност за генериране на съдържание и потенциалната възможност за злоупотреба пораждаат значителни етични опасения. В редица случаи инструментите, задвижвани от ИИ допринасят за опазването на академичната интегритет чрез механизми като откриване на плагиатство, онлайн наблюдение на изпити и повишена прозрачност на изследванията. В същото време са възможни и злоупотребата с тях –прекомерна зависимост на студентите да пишат есета, генерирани от ИИ, и прекомерна зависимост от автоматизирани системи, което като цяло подкопава основните академични ценности на оригиналност и автентичност [27]. Някои заключения от по-рано споменатото проучване също разкриват опасения относно способността на ИИ да улеснява плагиатството, да подкопава академичната интегритет и да отслабва критичното мислене [28]. Като цяло настроенията на участниците в това проучване са негативни и се отправят силни призови за съблюдаване на етични насоки, строга оценка на инструментите на ИИ и фокус върху запазването на съществените човешки елементи в образованието.

Въздействието от използването на дигиталните инструменти може да доведе до несправедлива оценка в случаите, когато системите за ИИ са обучени на базата на пристрастни данни, което води до несправедливи резултати при оценката, при която е използван ИИ [29]. Проучването разглежда ефектите от пристрастността на ИИ както върху отделните лица, така и върху обществото, и очертава настоящите стратегии за смекчаване на последиците, включително предварителна обработка на данни, избор на модели и техники за последваща обработка. Предизвикателствата, свързани с генеративните модели на ИИ, подчертават необходимостта от целенасочени подходи, специално разработени за справяне с тези сложности [ibid]. Неотдавна проведено проучване относно нарушенията на академичната честност в Обединеното кралство отчита близо 7000 потвърдени случая на измама с използване на AI инструменти през 2023–24 г., което представлява 5,1 инцидента на 1000 студенти — над трикратно увеличение от 1,6 на 1000 през 2022–23 г. [30].

Друг резултат от прекомерното осляняне на използване на средства с ИИ е, че то намалява критичното мислене, като ограничава възможностите за независимо разсъждение и решаване на проблеми. Учениците могат да развият повърхностни знания, тъй като ИИ предоставя готови отговори, което подкопава аналитичното мислене и умения, както и критичното мислене [31]. Събраните данни от различни източници показват, че ИИ в образованието променя академичната интегритет [32]. Зависимостта на учителите от инструменти за откриване на ИИ е нараснала рязко – с 68%, което е с 30 пункта повече, докато действията на учениците по отношение на плагиатството/взимане на готово на информация и данни, свързано с ИИ, се е увеличила от 48% на 64% в рамките на една година [33]. Познаването и използването на ChatGPT е почти всеобщо сред учениците (около 90%), като повечето го прилагат при решаване на домашните си задачи. Обединеното кралство отчита най-високите нива на плагиатство (33% от маркираните документи), въпреки относително ниското съдържание, генерирано от ИИ (10%). Като цяло 89% от студентите използват ChatGPT за домашните си задачи, а 19% от текстовете, генерирани от ChatGPT, съвпадат с предишни текстове и други източници в Интернет [34]. Информация от същия източник показва, че сред студентите, които използват ИИ за училищни задачи, 50% го използват частично - по-голямата част от задачите те ги изпълняват сами; 30% разчитат на ИИ за по-голямата част от заданията, а 17% предават задачи, изпълнени с ИИ, без редакции. В допълнение към споменатото по-рано недостатъчно обучение за отговорно използване на технологиите, може да се посочи, че е необходимо да се доразвият и подобрят етичните насоки за използване на ИИ в академичната среда. Това включва отчитане на важни човешки елементи в образованието, включително културното приемане на съкратени пътища в академичната среда.

Въпроси, свързани с равнопоставеността и приобщаването

Дигитализацията може да задълбочи неравенствата, основани на социално-икономическия статус, пола, възрастта или уврежданията, като ограничава достъпа до технологиите. Едно проучване показва, че осведомеността на учителите по отношение на информационната сигурност в дигитализираното образование е на средно ниво, с различия според пола и

предметната област [35]. Човешките фактори, като небрежност и невежество на потребителите, са основните причини за нарушения на сигурността, като представляват риск за неприкосновеността на личния живот в дигиталното образование. Същото проучване показва, че жените учители проявяват по-ниска осведоменост по темата, което подчертава уязвимостта, свързана с пола [ibid]. Това показва не само важността на човешкия фактор при нарушенията на сигурността, но и важността на по-всеобхватно обучение с акцент върху образованието по въпросите на равенството между половете. Други интересни данни сочат, че макар момчетата като цяло да превъзхождат момчетата по компютърна и информационна грамотност в ранна възраст, те са по-малко склонни да продължат да изучават ИКТ и STEM (Science, Technology, Engineering, and Mathematics) предмети, когато преминават към по-високи нива на образование [36]. Повечето институции предприемат стъпки за осигуряване на достъп до дигитални платформи за студенти с увреждания. Проучванията обаче подчертават, че предоставянето на възможности се простира отвъд базовото предоставяне на технологии и достъп до Интернет. То включва оборудване на студентите с дигитални умения, насърчаване на увереността чрез наставничество, насърчаване на самозащитата и активността на учащите, ангажиране на студентите в процесите на вземане на решения, които отговарят на техните нужди, и осигуряване на достъп до информация в рамките на приобщаваща педагогическа рамка [37]. Учениците в отдалечени райони и от семейства с по-ниски доходи са в неравностойно положение. Доклад на ЮНЕСКО показва, че групите в неравностойно положение (лица, живеещи в отдалечени райони, или тези, които са разселени) се сблъскват с трудности в ученето, имат ограничено време или са пропуснали предишни възможности за образование [38]. Същото проучване показва, че въпреки бързото разпространение на дигиталните технологии, продължават да съществуват значителни различия в достъпа на учащите се до тях. Тези групи обикновено притежават по-малко устройства, имат по-малко надеждна Интернет връзка и разполагат с по-малко ресурси в домакинството. Въпреки че цената на използването и придобиването на дигиталните технологии намалява, за много хора тя остава непосилно висока [ibid]. В резултат на горепосочените проблеми можем да заключим, че социално-икономическите бариери пред достъпа до устройства или Интернет връзка, културните норми, свързани с предразсъдъците по отношение на пола или възрастта при внедряването на технологии, и липсата на достъп до академични платформи представляват сериозно предизвикателство и повод за загриженост по отношение на ДТ в образованието. Може да се твърди също, че по отношение на гореспоменатите подобрения в етичните насоки (в предишния параграф - Етични предизвикателства) за използването на ИИ в академичната среда, институциите трябва да вземат предвид важни културни и социално-икономически теми като социално-икономически статус, пол, възраст, увреждания и ограничен достъп до технологии.

Категоризация на рисковете с примери

Гореспоменатите рискове са описани и категоризирани в таблица 1 с примери за всеки от тях.

Таблица 1: Категоризация на рисковете, свързани с човешкия фактор в академичната среда

Категория	Описание на риска	Примери / Въздействия
Съпротива срещу промяната	Преподаватели, персонал или студенти, които се съпротивляват на въвеждането на нови технологии поради културни норми, страх или скептицизъм.	Предпочитание към преподаване на живо, нежелание да се използват LMS/ИИ, страх от загуба на работа, опасения относно ангажираността на студентите.
Притеснения относно натоварването	Дигиталните инструменти се възприемат като допълнителна задача/бреме, а не като средство за намаляване на усилията.	Необходимо е време за изучаване на системите, преработване на уроците, управление на онлайн платформите; съобщава се за изчерпване на силите на преподавателите (прегаряне).
Страх от отпадане/уволнение	Тревога, че изкуствените интелект или дигиталните инструменти могат да заменят	Несигурност на работното място, подценяване на човешкото преподаване.

	човешката роля в преподаването.	
Липса на стимули/признание	Ограничена институционална мотивация за дигитално преподаване.	Дигиталното преподаване не се третира като равностойно на преподаването в класната стая, по-малко награди за иновации.
Дигитално разделение/противопоставяне	Неравен достъп до технологии и умения сред студентите и преподавателите.	Студенти, които нямат устройства или надежден Интернет; по-възрастни преподаватели с по-малко дигитални умения; неравенства, свързани с доходите, пола, възрастта, уврежданията.
Недостатъчна дигитална грамотност	Ограничена способност за ефективно използване на дигитални инструменти.	Трудности с навигацията в учебни платформи, използването на ИИ; погрешни представи, че всички ученици са дигитално грамотни; възрастови разлики в уменията.
Липса на обучение и подкрепа	Институциите въвеждат нови инструменти без достатъчна подготовка или помощ.	Разочарование на преподавателите, ниска степен на приемане, затруднения на студентите с новите технологии.
Проблеми, свързани с поверителността и сигурността	Рискове от събиране на данни, наблюдение и нарушения.	Страх от кражба на самоличност, нарушения в Zoom, инструменти за откриване на ИИ, използвани без съгласие, обработка на данни от трети страни.
Етични предизвикателства	Прекомерното разчитане на използването на ИИ подкопава академичната честност и критичното мислене.	Есета, създадени с ИИ, плагиатство, пристрастност при оценяването с ИИ, упадък на независимото мислене.
Въпроси, свързани с равнопоставеността и приобщаването	Социални неравенства, засягащи достъпа и участието.	Социално-икономически бариери, неравенства между половете, предизвикателства, свързани с уврежданията, културно недоверие към дигиталните средства.

Източник: Авторът

Заклучения

Успешната дигитална трансформация в образованието изисква институционална ангажираност чрез финансиране и културно привеждане в съответствие, активно участие на заинтересованите страни за намаляване на съпротивата и непрекъснато оценяване на рисковете. Сътрудничеството с доставчиците на технологии гарантира сигурни и достъпни платформи, докато стратегиите, съобразени с културните особености, отговарят на местните нагласи по въпроси като неприкосновеността на личния живот. Някои от най-големите предизвикателства в ДТ за висшите учебни заведения включват липсата на ясни институционални насоки, слабото прилагане на съществуващите дигитални политики и прекалено строгото им изпълнение. Без всеобхватни политики, обхващащи области като поверителност на данните, натовареност и сътрудничество, институциите не разполагат с необходимата структура и инструменти, необходими за ДТ.

Рисковете за ДТ във висшите учебни заведения могат да се разглеждат като предизвикателства, пречки и ограничения, които възпрепятстват отделни лица, групи или организации да постигнат целите на ДТ. В литературата по ДТ тези теми се разглеждат като фактори, които възпрепятстват, забавят или пречат на процеса на трансформация. В допълнение, тези фактори се характеризират като реални или съществуващи характеристики на правни, социални, технологични или институционални рамки, които действат срещу ДТ и ограничават усилията за преконфигуриране на достъпа до информация, хора и организации чрез технологични средства. Ние ги разгледахме по методологичен начин и създадохме категоризация с описание на рисковете пред ДТ. Рискът в дигиталния свят обаче се обновява и изниква във всякакви нови форми постоянно, тъй като динамиката на появата на нови дигитални инструменти и парадигми е много голяма.

Благодарност

Тази работа беше подкрепена от проект, финансиран от програмата за научни изследвания и иновации „Хоризонт Европа“ на Европейския съюз под номер GA N 101131544. Съдържанието отразява единствено мнението на авторите, а REA и ЕК не носят отговорност за използването на информацията, съдържаща се в настоящия документ.

Литература / References

1. O. Joseph, O. Onwuzulike, and K. Shitu, ‘Digital transformation in education: Strategies for effective implementation’, *World J. Adv. Res. Rev.*, vol. 23, Sep. 2024, doi: 10.30574/wjarr.2024.23.2.2668.
2. Murdjeva A. and L. Boyanov, ‘Fulfilled and unfulfilled expectations for digitization’, in *Proceedings of the Innovative Information Technologies for Economy Digitalization Conference*, UNWE, Sofia, Bulgaria: UNWE, Sofia, 2024, pp. 121–130. [Online]. Available: <https://www.unwe.bg/doi/iited/2024/IITED.2024.15.pdf>
3. A. Johnson, M. Jacovina, D. Russell, and C. Soto, ‘Challenges and Solutions when Using Technologies in the Classroom’, 2016, pp. 13–30. doi: 10.4324/9781315647500-2.
4. ‘WGU Labs | CIN Faculty EdTech Survey 2024’. Accessed: Aug. 24, 2025. [Online]. Available: <https://www.wgulabs.org/posts/2024-cin-faculty-edtech-survey-edtech-and-the-evolving-role-of-faculty>
5. P. D. Deep, N. Ghosh, and C. (Yixin) Chen, ‘Faculty Burnout in Higher Education: Effects on Student Engagement, Learning Outcomes, and Artificial Intelligence-Driven Institutional Responses’, *J. Educ. Dev. Psychol.*, vol. 15, pp. 29–29, Apr. 2025, doi: 10.5539/jedp.v15n1p29.
6. Magali Dubosson, Emmanuel Fragnière, Samuele Meier, ‘Early detection of human-related risks in an increasingly digitized work environment’, *Digit. Transform. Soc.*, vol. 1, no. 1, pp. 48–65, Aug. 2022, doi: <https://doi.org/10.1108/DTS-05-2022-0017>.
7. B. Deacon, M. Laufer, M. A. Mende, T. Tschache, and L. O. Schäfer, ‘Resisting digital change at the university: an exploration into triggers and organisational countermeasures’, *Eur. J. High. Educ.*, vol. 0, no. 0, pp. 1–24, doi: 10.1080/21568235.2025.2512735.
8. Richard R. Snyder, ‘Resistance to Change among Veteran Teachers: Providing Voice for More Effective Engagement’, *NCPEA Int. J. Educ. Leadersh. Prep.*, vol. 22, no. 1, 2017.
9. Aimee Thalberg, ‘Technology Resistance: Helping Administration Stop It in Our Schools’. St. Cloud State University, 2019. [Online]. Available: https://repository.stcloudstate.edu/im_etds/24/
10. P. Nirmani, ‘Barriers to digital participation in developing countries: Identifying technological, social, and cultural obstacles to community involvement’, *GSC Adv. Res. Rev.*, vol. 23, pp. 061–071, May 2025, doi: 10.30574/gscarr.2025.23.2.0130.
11. ‘2023 Faculty and Technology Report: A First Look at Teaching Preferences since the Pandemic’, EDUCAUSE. Accessed: Aug. 19, 2025. [Online]. Available: <https://www.educause.edu/ecar/research-publications/2023/faculty-and-technology-report-a-first-look-at-teaching-preferences-since-the-pandemic/introduction-and-key-findings>
12. Z. N. Khlaif *et al.*, ‘University Teachers’ Views on the Adoption and Integration of Generative AI Tools for Student Assessment in Higher Education’, *Educ. Sci.*, vol. 14, no. 10, p. 1090, Oct. 2024, doi: 10.3390/educsci14101090.
13. P. Q. P. Yanique Brown, ‘Overcoming Teacher Resistance to Technology and Artificial Intelligence in the Classroom’, *Click In*, Mar. 2024, Accessed: Aug. 19, 2025. [Online]. Available:

- https://www.academia.edu/116731756/Overcoming_Teacher_Resistance_to_Technology_and_Artificial_Intelligence_in_the_Classroom
14. T. Gkrimpizi, V. Peristeras, and I. Magnisalis, ‘Classification of Barriers to Digital Transformation in Higher Education Institutions: Systematic Literature Review’, *Educ. Sci.*, vol. 13, no. 7, p. 746, Jul. 2023, doi: 10.3390/educsci13070746.
 15. ‘The Role of Digital Transformation in Education in Promoting Sustainable Development’, *Virtual Econ.*, vol. 5, no. 4, pp. 65–86, 2022.
 16. A. M. Syed, S. Ahmad, A. Alaraifi, and W. Rafi, ‘Identification of operational risks impeding the implementation of eLearning in higher education system’, *Educ. Inf. Technol.*, vol. 26, no. 1, pp. 655–671, Jan. 2021, doi: 10.1007/s10639-020-10281-6.
 17. G. Robertsons and I. Lapiņa, ‘Digital transformation as a catalyst for sustainability and open innovation’, *J. Open Innov. Technol. Mark. Complex.*, vol. 9, no. 1, p. 100017, Mar. 2023, doi: 10.1016/j.joitmc.2023.100017.
 18. Nomana Tariq and Nasir Mehmood, ‘Exploring Employees’ Resistance towards Digital Transformation: A Case Study of Higher Education Institution’, *J. Educ. Humanit. Res. Univ. Balochistan*, vol. 18, no. 2, pp. 113–135, 2024.
 19. S. Rafiq, S. Iqbal, and D. Afzal, ‘The Impact of Digital Tools and Online Learning Platforms on Higher Education Learning Outcomes’, May 2024.
 20. H. R. Slack and M. Priestley, ‘Online learning and assessment during the Covid-19 pandemic: exploring the impact on undergraduate student well-being’, *Assess. Eval. High. Educ.*, vol. 48, no. 3, pp. 333–349, Apr. 2023, doi: 10.1080/02602938.2022.2076804.
 21. J. Pacheco and R. Vega-Estrella, ‘Challenges Faced by Teachers When Integrating Technology in the Teaching of Virtual Courses in a Higher Education Institution’, *J. Educ. Dev.*, vol. 9, p. 47, Apr. 2025, doi: 10.20849/jed.v9i2.1503.
 22. B. Almekhled and H. Petrie, *The Role of Privacy and Security Concerns and Trust in Online Teaching: Experiences of Higher Education Students in the Kingdom of Saudi Arabia*. 2024, p. 77. doi: 10.5220/0012616200003693.
 23. ‘An Analysis of the 2020 Zoom Breach | CSA’. Accessed: Aug. 25, 2025. [Online]. Available: <https://cloudsecurityalliance.org/blog/2022/03/13/an-analysis-of-the-2020-zoom-breach>
 24. M.-J. Gallego-Arrufat, I. García-Martínez, M.-A. Romero-López, and N. Torres-Hernández, ‘Digital rights and responsibility in education: A scoping review’, *Educ. Policy Anal. Arch.*, vol. 32, Jan. 2024, doi: 10.14507/epaa.32.7899.
 25. ‘Evaluating the Effectiveness and Ethical Implications of AI Detection Tools in Higher Education’, Sciety. Accessed: Aug. 26, 2025. [Online]. Available: https://sciety-labs.elifesciences.org/articles/by?article_doi=10.20944/preprints202507.2233.v1
 26. M. Forment, M. Guerrero, D. Amo-Filva, C. Severance, and D. Fonseca, ‘Privacy and E-Learning: A Pending Task’, *Sustainability*, vol. 13, p. 9206, Aug. 2021, doi: 10.3390/su13169206.
 27. R. Elaïess, ‘Academic Integrity in the Age of Artificial Intelligence: Exploring Ethical Challenges in A Digital Era’, vol. 8, pp. 12–15, Nov. 2024.
 28. M. Pikhart and L. H. Al-Obaydi, ‘Reporting the potential risk of using AI in higher Education: Subjective perspectives of educators’, *Comput. Hum. Behav. Rep.*, vol. 18, p. 100693, May 2025, doi: 10.1016/j.chbr.2025.100693.
 29. E. Ferrara, ‘Fairness And Bias in Artificial Intelligence: A Brief Survey of Sources, Impacts, And Mitigation Strategies’, *Sci*, vol. 6, no. 1, p. 3, Dec. 2023, doi: 10.3390/sci6010003.
 30. M. Goodier, ‘Revealed: Thousands of UK university students caught cheating using AI’, *The Guardian*, Jun. 15, 2025. Accessed: Aug. 29, 2025. [Online]. Available:

- <https://www.theguardian.com/education/2025/jun/15/thousands-of-uk-university-students-caught-cheating-using-ai-artificial-intelligence-survey>
31. C. Zhai, S. Wibowo, and L. D. Li, ‘The effects of over-reliance on AI dialogue systems on students’ cognitive abilities: a systematic review’, *Smart Learn. Environ.*, vol. 11, no. 1, p. 28, Jun. 2024, doi: 10.1186/s40561-024-00316-7.
 32. ‘AI Plagiarism Statistics 2025: Transforming Academic Integrity’. Accessed: Aug. 29, 2025. [Online]. Available: <https://artsmart.ai/blog/ai-plagiarism-statistics/>
 33. ‘As teacher use of AI detection grows, discipline guidance a mixed bag | K-12 Dive’. Accessed: Aug. 29, 2025. [Online]. Available: <https://www.k12dive.com/news/ai-detection-tools-student-discipline/711835/>
 34. S. Joshi, ‘30+ Plagiarism Statistics and Facts You Need to Know in 2025’, University Edge. Accessed: Aug. 29, 2025. [Online]. Available: <https://learn.g2.com/plagiarism-statistics>
 35. H. Sapanca and S. Kanbul, ‘Risk management in digitalized educational environments: Teachers’ information security awareness levels’, *Front. Psychol.*, vol. 13, p. 986561, Sep. 2022, doi: 10.3389/fpsyg.2022.986561.
 36. Directorate-General for Education, Youth, Sport and Culture (European Commission), *International Computer and Information Literacy Study (ICILS) in Europe, 2023: main findings and educational policy implications*. Publications Office of the European Union, 2024. Accessed: Aug. 29, 2025. [Online]. Available: <https://data.europa.eu/doi/10.2766/5221263>
 37. N. Manase, ‘The Digital inclusion of disabled students in Open and Distance e-Learning: Going beyond access to empowerment’, *Int. J. E-Learn. Distance Educ. Rev. Int. E-Learn. Form. À Distance*, vol. 39, no. 2, 2024, doi: 10.55667/10.55667/ijede.2024.v39.i2.1343.
 38. UNESCO, ‘Technology in education’, 2023 GEM Report. Accessed: Aug. 29, 2025. [Online]. Available: <https://gem-report-2023.unesco.org/technology-in-education/>

THE FUTURE OF DATA: UTOPIA OR DYSTOPIA?

**Assof Prof. Aleksandrina Murdzheva¹,
Nikolay Vlaychev²**

e-mail: amurdjeva@unwe.bg¹,

e-mail: nikolay_vlaychev@scalefocus.com²

Abstract

The journey into the future of data is not merely a technical discussion; it is a profound exploration of conceptual, philosophical, business, and ethical aspects that are rapidly reshaping our world. Is the future of data a utopia or a dystopia, and how do we ensure the choice is ours? The speed of technological development often outpaces its societal adoption, posing significant challenges that demand immediate, critical attention.

Ключови думи: Data, AI, data usage, data ethics, data storage

JEL: O03

The Unstoppable Snowball of Data Generation

Creation vs. Meaningfulness

The modern world is characterized by an incessant, accelerating generation of data. Everything—from our daily interactions to complex global systems—is contributing to this digital deluge. With the adoption of technologies like the Internet of Things (IoT), the concept of blockchain, and the rise of generative AI creating synthetic data, nowadays the world is living in interesting days. The consensus is clear: humanity does not have the luxury of choosing to stop generating data. It is an unstoppable snowball driven by the systems and media we have established.

The volume of data worldwide is growing exponentially, with growth driven primarily by a few key sources in recent years [1] [2] [3].

While precise figures in zettabytes (ZB) for the growth of each individual source (such as Social Media vs. IoT) are not readily available in public reports, the statistics provide a clear picture of the key drivers and the breakdown by data format.

- *Current Volume:* Projections for 2024 place the global volume of data (created, captured, copied, and consumed) at around 149 zettabytes (ZB).
- *Projected Growth:* The volume is expected to grow to 181 ZB by the end of 2025, representing an annual growth rate of over 20%
- *Historical Context:* Approximately 80% of all data in the world has been generated in the last two to four years alone.

Volume growth is driven by the following key categories related to social and corporate behavior:

¹ Assoc Professor, Information technologies and communications, University of National and World Economy, Sofia, Bulgaria

² Software Engineer at ScaleFocus, Bulgaria, <https://www.linkedin.com/in/nvlaychev/>

Source	Description of the Contribution	Key Statistics
User-Generated Data	Data created by end users through personal devices and online activity (photos, videos, posts).	70% of all global data is generated by users.
AI and ML	AI systems are " data vampires " that continuously generate, process, and manipulate vast data sets during training and operation.	AI is a key driver of data search, including synthetic data.
Social Media & Streaming	Platforms like TikTok, YouTube, and Instagram generate billions of daily uploads, especially high-resolution videos (\$4K\$ and \$8K\$), which account for the largest share of internet traffic.	Videos account for more than half of internet traffic.
Internet of Things	Sensors and connected devices in industry and households generate data streams in real-time.	Their growing use has been cited as a major driver of growth.
Cloud Solutions	Moving to the cloud removes storage limitations, allowing organizations to store more data for longer periods and support real-time analytics.	The cloud helps create and retain larger volumes of data.

One of the clearest indicators of the source of data is its structure. The majority of the data generated today comes from non-traditional sources (such as social media and video), leading to the predominance of the unstructured format:

Data Format	Share of Total Volume	Growth Rate
Unstructured Data	80% to 80% of the total volume	They grow approximately three times faster than structured data.
Structured Data	About 20% of the total volume	Include traditional data from databases and transactional systems.

Unstructured data (emails, text messages, blog posts, videos, audio, etc.) is the largest and fastest-growing segment, directly reflecting the dominant contribution of consumer and media sources.

However, the core dilemma is not the volume, but the utility of this torrent. Simply having more data does not automatically translate into wisdom or efficiency. More data does not mean more knowledge. More data does not mean that the business and the society understand our business better. More data does not mean that people are better in our everyday lives.

While the creative segment of humanity needs more data to develop, invent, and understand the world, a crucial question arises for the rest: are these vast quantities of data useful, or do they produce more chaos in our business and everyday life, rather than produce something useful? The immediate challenge is not to stop generation, but how to handle this one that we have now.

The Data Storage Paradox: Capacity vs. Processing

In terms of raw data storage, the industry has achieved remarkable success. A huge progress has been made in the field of storage of ever-increasing volume-the powerful big data technologies and even the advent of quantum computing for data processing. The technology world is successful - powerful technologies how to store a large volume of data are here. Today, what was once a scary amount years ago (terabytes) is now commonplace.

Yet, this success in storage has created a new, deeper challenge in processing and insight extraction. Storing the data is only half the battle. The true difficulty lies in making the data meaningful and valuable. This is captured perfectly with the metaphor of a library: It's like having a library without

bookshelves. So you have just an empty room that physically you can store your books, so your data on top of each other. But if you need to find something, it costs you a lot of efforts and a lot of time to get it.

The storage problem vs. usage problem

The future challenge, therefore, shifts from being a storage problem to a usage problem. A new, powerful ideas on how to process and learn from the data is needed. This need has spurred innovative ideas in data science and computer engineering, such as exploring DNA storage [4] and hologram storage [7], essentially stealing some ideas from the nature to solve social and business cases.

The most innovative research and trends in the field of Big Data analysis currently revolve around the deeper integration of Artificial Intelligence (AI), new data architectures, and a focus on real-time processing and ethical use.

Here are the most significant innovative research areas:

1. Advanced AI and Machine Learning Techniques

The use of AI and ML has moved beyond simple predictive models to highly sophisticated, autonomous systems:

- **Augmented Analytics** [8] [9]: This involves using AI and Machine Learning to automate and enhance the entire data analysis workflow. It helps non-data-scientists generate insights, find patterns, and explain findings in natural language, drastically improving the speed of decision-making.
- **Deep Learning (DL)**: Research continues to push the boundaries of neural networks, particularly in:
 - **Natural Language Processing (NLP)**: Advancements like Large Language Models (LLMs) and Retrieval Augmented Generation (RAG) are transforming how unstructured text data (customer feedback, documents, etc.) is understood, summarized, and queried for analysis.⁶
 - **Computer Vision**: Using Deep Learning models like Convolutional Neural Networks (CNNs) to extract valuable insights from image and video data at scale (e.g., medical imaging diagnostics, monitoring in manufacturing).
- **Agentic AI**: This is an emerging and highly innovative area where AI systems can perform entire workflows autonomously—setting goals, planning tasks, executing actions, and adapting to feedback without continuous human intervention.

Evolving Data Architectures and Real-Time Processing

Innovative research is focused on new data processing models to handle data volume, velocity, and variety more effectively [10] [11]:

- **Data Mesh**: A shift from centralized data warehouses/lakes to a decentralized architecture. It treats data as a product, with domain-specific teams owning and serving their data assets. This architecture is designed for massive scale and speed, improving data accessibility and quality across the organization.
- **Data Fabric**: This is a unified, intelligent platform that uses AI and machine learning to integrate and manage data from multiple disparate sources (on-premises, cloud, edge) without having to physically move all the data. It's focused on making data readily available and governed.
- **Edge Analytics**: With the proliferation of Internet of Things (IoT) devices, this trend focuses on performing analysis directly on the data source (the "edge") instead of sending everything to a central cloud. This is critical for real-time applications like predictive maintenance and autonomous vehicles, where latency must be near zero.

Ethical and Secure Data Usage

As Big Data becomes more influential, research is heavily invested in its responsible application [5]:

- **Responsible AI and Ethics:** A major research focus is on mitigating algorithmic bias, ensuring fairness, and developing mechanisms for transparency (explainability, or XAI) and accountability in AI/ML models that make critical decisions.¹²
- **Data Observability:** This involves applying DevOps principles to data pipelines to ensure data quality, freshness, and reliability. It uses automated monitoring and alerting to detect anomalies in data (schema changes, sudden volume drops) before they corrupt analysis or model training.¹³
- **Privacy-Enhancing Technologies (PETs):** Innovative research into techniques like Homomorphic Encryption and Federated Learning allows data to be analyzed or models to be trained while the underlying data remains encrypted or localized, respectively. This addresses increasingly stringent data privacy regulations.

The Ethical Crossroads: Privacy and Ownership

The Privacy Challenge: Technology vs. Behavior

The conversation naturally transitioned to data privacy, revealing one of the most significant philosophical hurdles of the data age. The panel agreed that privacy is crucial for both businesses (which expend great effort to anonymize and secure data) and individuals, whose personal data is constantly exposed and interconnected with professional and daily lives.

However, the problem is not a lack of protective technology. The technological solutions for protection are numerous and "most of these ideas are successful". The true crisis lies in the fact that social behavior actively undermines legal protection.

The laws and all these compliance documents, they try to protect our data, but our social behavior actually exposes this data without any boundaries.

The regulation system is often too slow and unable to keep up with the fast-changing ways humans behave online. Users often expose sensitive data (like pictures or personal identifiers) on platforms without considering the consequences, sometimes becoming victims to be followed or to be targeted.

The implication is profound: the privacy problem is fundamentally a **non-technological problem**. It is a challenge related to ethics, behavior, and social consciousness. The solution, therefore, requires education and critical thinking, rather than merely passing more laws.

The Ownership Enigma

The privacy dilemma is inseparable from the question of data ownership. The panel raised the perplexing question: when a user interacts with a technology that creates data, who is the true owner—the individual or the company that created the technology?

The example of **IoT (Internet of Things)** devices illustrates this paradox vividly. While a knowledgeable user might be able to change this pattern and transfer the data to a personal server, this highlights a deeply confusing situation where one possesses something but may not be able to fully own or protect it.

In most cases, the user (you) is considered the "data subject" and retains the rights to control and access your personal data, but the company that collects, processes, and stores the data is generally considered the legal "data controller" or "data holder."

1. Data from IoT Devices (Smart Homes, Wearables, etc.)

For data generated by an Internet of Things (IoT) device, ownership is usually determined by a combination of the device's Terms of Service (ToS), the type of data, and legal regulations like GDPR.

- Personal Data (Data Subject Rights): For data that is personally identifiable (like your name, location history, or health data from a fitness tracker), you (the individual user) have significant rights to access, correct, delete, and port that data, especially under comprehensive laws like the EU's General Data Protection Regulation (GDPR).¹ This is often described as "ownership as control" rather than traditional property ownership.
- Non-Personal/Aggregated Data (Company Control): The manufacturer, service provider, or data processing platform owner often claims ownership or an exclusive right to use the non-personal data (or personal data that has been anonymized or aggregated) that is collected. For example, a smart thermostat company may own the aggregated data on average temperature settings to improve its product or sell general market insights.
- Contractual Clauses: The fine print in the Terms of Service you agree to when setting up the device usually grants the company a broad license to use and monetize the collected data.

New Legislation: Laws like the EU Data Act [6] are attempting to give users of connected products (like IoT devices) greater control and access over the data they generate, even for non-personal data.

2. *Data from Your Mobile Phone*

The data generated by your mobile phone is handled by several different entities, none of whom typically have exclusive ownership over all of it.

- You (The User/Data Subject): You own the content you create (photos, contacts, texts). More importantly, you have privacy rights over your personal data.³ You are the one who grants or revokes permissions for apps to access things like your location, microphone, and camera.⁴
- App Developers (Third Parties): Apps (like social media, games, or utilities) collect data based on the permissions you grant and the privacy policies you agree to. They typically own the database they create from your usage patterns and may use or sell this data for targeted advertising.
- Mobile Carrier/Phone Provider (ISP): Your carrier (e.g., T-Mobile, Verizon) can see information about your data consumption, general location (via cell tower connection), the metadata of your calls and texts (who you called, when, and for how long), and sometimes the websites you visit (if you're using their cellular network without a VPN).⁵ They use this for billing and may also use it for marketing purposes, but in many jurisdictions, they are restricted by law from sharing the content of your communications.
- Operating System Providers (Google/Apple): The company that provides your phone's operating system (Android or iOS) collects extensive data on your usage, app activity, and location to provide services, personalize your experience, and target ads. You typically agree to this when setting up the device.

The rapid evolution of technology—specifically cloud solutions and blockchain initiatives—has radically altered the traditional concept of ownership. Because users are often "unable to store them, we are unable to protect them [data], and actually... Or to track what is happening with this data," the opportunity to own one's own things is lost. This is amplified by the fact that owning one's own data has become "very expensive".

Data ownership is often replaced by a structure of complex rights and controls:⁶

- Customers have rights over your personal data (especially in regions with strong privacy laws).⁷
- Companies have contractual rights to use and monetize the data they collect and process.⁸

The challenge to regulatory bodies is significant. Technology has created a world with a new reality, and lawyers and data engineers are not able to respond to this very high speed. Ownership and privacy are intertwined; the absence of one compromises the other. The resolution, according to the discussion,

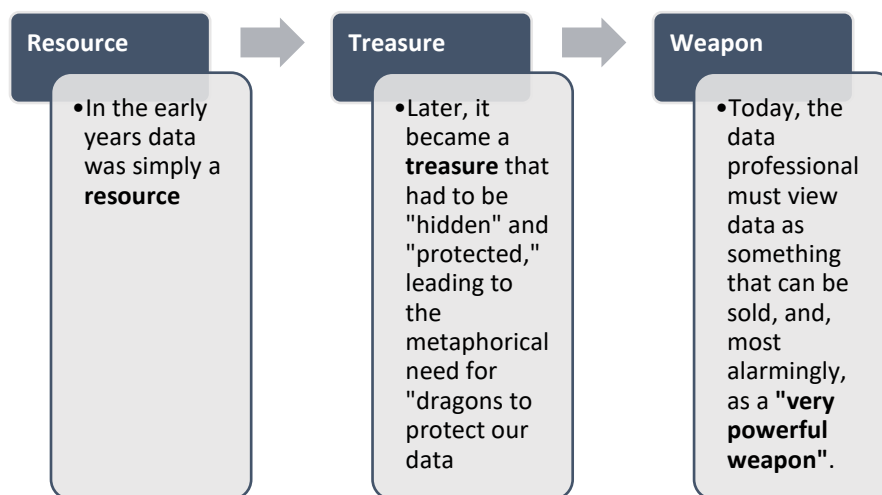
might be defined by one thing: the possibility to choose. Ownership is defined by the freedom to choose whether to share data or to keep it protected.

Data's Transformation: Resource, Treasure, Weapon

Monetization and Risk

The conceptual value of data has undergone a startling transformation, driven by its monetization. The phrase "data is the new oil" encapsulates its immense business importance. In 2006, Clive Humby, a British mathematician and data science entrepreneur, introduced the phrase “Data is the New Oil”. In simple words, what he said is that - “Data is the new oil. It is a valuable source. However unrefined data is not useful. It must be converted into valuable entities that drive profit. Companies now actively monetize their data, leading to the concern that they may be selling not just their own data, but *our* data.

And the data identity evolves from Resource to Weapon



This final identity—data as a weapon—necessitates the invention of more agile technology that can adapt to rapidly changing situations, as today's reality is often quite different from tomorrow's. This volatility and risk make data a gold mine for some, but a potential nightmare when its power is unleashed.

The Need for New Definitions

The current, traditional definition of data is woefully inadequate for the modern age. The complexity introduced by privacy, ownership, monetization, and weaponization demands a conceptual shift.

The current definition of data is not enough. The data is not only bytes, strings. It's something more and to the definition more things should be added so to be able to define the whole thing the whole technology environment creates and chalanges.

The signed professionals in computer science fields should convene to provide another, more conceptual definition of data. Data is no longer merely the raw zeros and ones of information theory [12][13]. It is intrinsically linked to behavior, business, and social dynamics. Redefining the term is crucial for the technology world to be able to adapt and create solutions that cover all these new aspects, particularly concerning storage, privacy, and ownership management.

The AI Game Changer

The Data Vampire and Synthetic Reality

Artificial Intelligence (AI) as the ultimate game changer and a central driver of the data apocalypse or utopia. AI was starkly defined as a data vampire. AI technologies consume data in such volumes that actually AI is one of the reasons that increased the need of data, AI is the reason this noble, this data generation to be unstoppable because AI will continue to develop when there are more data.

AI creates a dilemma regarding the data it consumes—is it real or synthetic? This question of authenticity mirrors dilemmas in other spheres of life (e.g., food or clothing) but holds unique power in the digital realm.

The need for AI to generate synthetic data is itself a paradox. According to Gartner the forecast is that 60% of data used for AI will be synthetic by 2024, and 75% of businesses will use generated customer data by 2026 [14]. It arises because either not enough data is produced by real-life systems, or because the existing real data is not accessible—it is hidden due to privacy concerns, or the owner refuses to provide it. Thus, the privacy barriers that society seeks to establish are simultaneously driving AI to create an artificial reality.

The dilemma of real versus synthetic data ultimately comes down to data reliability. Now, with the introduction of AI, the responsibility for quality is shared with a non-human entity, forcing the industry to adapt its entire framework for verifying information.

Hallucination and Critical Thinking

The output of AI is not guaranteed to be reliable - the phenomenon of AI hallucination, a scenario where a machine learning model, if "not learn[ed]... quite enough," becomes prone to error and bias. This risk is amplified by our own incomplete understanding of the parameters driving these systems.

In the new environment, the most vital human skill is **critical thinking** - It's not to know everything, but to ask yourself, is it a valid one, how I can verify that. Without critical thinking, AI will produce a fake reality, unreliable decisions. While technologies are needed to produce the knowledge required to distinguish between real and synthetic data, the human element—curiosity and the willingness to debug the source like code—is essential. The responsibility to detect and check for the synthetic, fake, or biased content—whether in the news, static data, or AI-generated responses—falls upon the users. The future of data with AI is bright only if customers remain critical thinking people. If society chooses merely to consume the results without questioning, the future is not so bright.

The Evolving Role of the Data Professional

Beyond Technology

The complexities of the new data reality have fundamentally altered the role and responsibilities of the data professional—the data engineer, data scientist, or database developer. They are no longer just custodians of technology; they must become strategists balancing ethics and efficiency. The new role of data professionals is defined from the observation that often we have to solve problems with technologies, but actually the problem is not a technology problem.

This requires professionals to adopt their cells because their responsibilities have changed. It is no longer enough to be a "technology person". The modern data professional must consider and find scalable solutions for non-technical aspects like privacy, ownership, and ever-changing social behavior. System architecture decisions, such as choosing between cloud and on-premise solutions, are now driven not just by functionality, but by the need to satisfy the ownership and privacy requirements of the data.

The professional development path has become very interesting because it demands curiosity about many things outside the purely technological domain. Designing systems requires considering the entire spectrum of human behavior and business realities, rather than falling back on traditional ways of using data.

Aligning Technology and Society

The complexity is further compounded by the fundamental shift in the responsibility for data quality. In the past, data quality was largely the responsibility of the data engineers—they were often guilty if the data was unreliable in BI or data warehouse projects.

Today, the game has changed with the introduction of AI as a new player. Now, a technology created by people—AI—will often be responsible for the source used in knowledge extraction systems. This

creates a disorienting feedback loop where it is difficult to distinguish where something ends and something begins in our play.

The new engagement for data professionals is to bridge the gap between technology and society, to align the two technologies and how people think about it and how they use it. A new responsibility is here to teach them how to consume it correctly. Data professionals must act as ambassadors, explaining how systems work so that the public is aware of "possible traps or misleading". The task is not merely to implement, but to teach and guide consumption.

The Mandate for Adaptation and Redefinition

The future's outcome is not set in stone; it is contingent upon human effort and willingness to evolve.

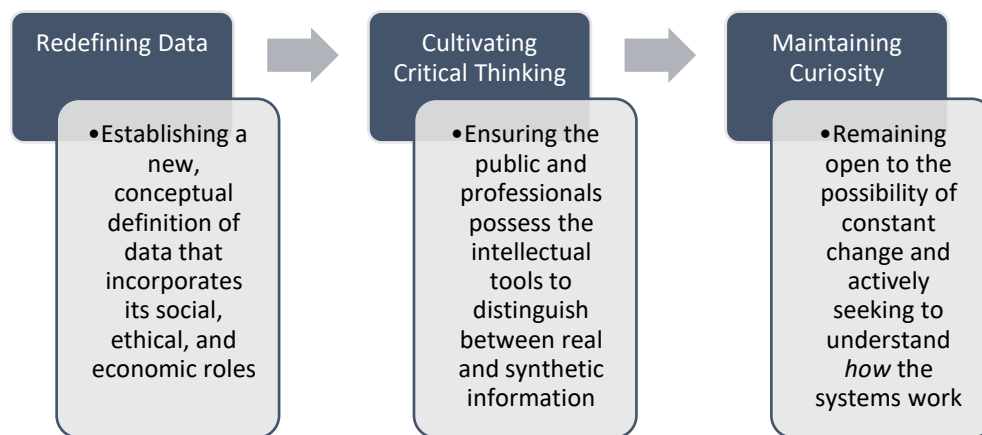
The Necessity of New Problem Solving

Once believed that all technologies in data processing and management solve problems, the current reality suggests the opposite. Nowadays all these technology actually... are creating problems and not always technology problems. These new problems are complex, non-traditional challenges that demand curiosity, critical thinking, and a willingness to adapt. The humorous observation was made that perhaps the industry is bored already and thus uses technology to produce new problems because there don't have enough problems" Yet, beneath the humor is the serious truth that the industry must accept these challenges as its new reality.

The Path to Control

To retain the power of choice—to steer away from dystopia—requires a fundamental commitment to growth and an open mindset.

The path involves three key mandates:



Conclusion: The Choice is Ours

The future of data is not predetermined. It is not an inevitable utopia, nor is it a guaranteed dystopia. The power to shape this future resides entirely with humanity. The journey ahead requires a radical acceptance of the complexity introduced by the digital age. Professionals must embrace a role that transcends pure technology, becoming ethical strategists and educators. Society must accept the necessity of developing its collective critical thinking to combat the risks of synthetic reality.

The ultimate message is one of responsibility: must stay open, remain curious, and never forget that technology often races ahead of our ability to adapt. By actively engaging with the problems our technology creates, by redefining our core concepts, and by choosing to think critically, humanity can ensure that the choice remains **ours**.

References

1. Taylor P. , Amount of data created, consumed, and stored 2010-2023, with forecasts to 2028, Nov 2025, Statista, <https://www.statista.com/statistics/871513/worldwide-data-created/>
2. Gkikas D., Prokopis T, AI in Consumer Behavior, Advances in Artificial Intelligence-based Technologies (pp.147-176), 2022
3. Bartley K., Big data statistics: How much data is there in the world?, 2025, <https://rivery.io/blog/big-data-statistics-how-much-data-is-there-in-the-world/#:~:text=In%202024%2C%20the%20global%20volume,by%20the%20end%20of%202025.>
4. Yiming, Dong & Sun, Fajia & Ping, Zhi & Ouyang, Qi & Qian, Long. (2020). DNA storage: Research landscape and future prospects. National Science Review. 7. 1092-1107. 10.1093/nsr/nwaa007. https://www.researchgate.net/publication/345359466_DNA_storage_Research_landscape_and_future_prospects#:~:text=Abstract%20and%20Figures,and%20sustainable%20data%20storage%20solution.
5. Opeyemi, Kehinde & Samuel, Adebis. (2023). Big Data Analytics: Emerging Trends, Technologies, and Innovations for Future., https://www.researchgate.net/publication/390630564_Big_Data_Analytics_Emerging_Trends_Technologies_and_Innovations_for_Future#:~:text=areas%20like%20drug%20discovery%2C%20financial,more%20accessible%20to%20non%2Dexperts.
6. Data Act explained [2023], <https://eur-lex.europa.eu/eli/reg/2023/2854/oj>, <https://digital-strategy.ec.europa.eu/en/factpages/data-act-explained#:~:text=A%20data%20holder%20must%20have,connected%20product%20or%20related%20service.>
7. Lin, Xiao & Liu, Jinpeng & Hao, Jianying & Wang, Kun & Yuanying, Zhang & Li, Hui & Horimai, Hideyoshi & Tan, Xiaodi. (2020). Collinear holographic data storage technologies. Opto-Electronic Advances. 3. 19000401-19000408. 10.29026/oea.2020.190004.
8. Minu, M S & Ahmad, Zoya. (2020). Augmented Analytics: The Future of Business Intelligence. 5. 7-13. 10.5281/zenodo.3757837.
9. Alghamdi, Noorah & Al-Baity, Heyam. (2022). Augmented Analytics Driven by AI: A Digital Transformation beyond Business Intelligence. Sensors. 22. 8071. 10.3390/s22208071.
10. Lu, Jiaheng & Holubova, Irena & Cautis, Bogdan. (2018). Multi-model Databases and Tightly Integrated Polystores: Current Practices, Comparisons, and Open Challenges. 2301-2302. 10.1145/3269206.3274269.
11. Yasuko Matsubara et al, MicroAdapt: Self-Evolutionary Dynamic Modeling Algorithms for Time-evolving Data Streams, Proceedings of the 31st ACM SIGKDD Conference on Knowledge Discovery and Data Mining V.2 (2025). DOI: 10.1145/3711896.3737048
12. Elmasri, R., & Navathe, S. B. (2016). Fundamentals of Database Systems (7th ed.). Pearson Education.
13. Floridi, L. (2010). Information: A Very Short Introduction. Oxford University Press.
14. Gartner Identifies Top Trends Shaping the Future of Data Science and Machine Learning [2023] <https://www.gartner.com/en/newsroom/press-releases/2023-08-01-gartner-identifies-top-trends-shaping-future-of-data-science-and-machine-learning>

СЪВРЕМЕННИ АСПЕКТИ НА ПРИЛОЖЕНИЯ НА ИЗКУСТВЕН ИНТЕЛЕКТ В БИЗНЕС ИНФОРМАЦИОННИТЕ СИСТЕМИ И В ЧАСТНОСТ ERP СИСТЕМИТЕ

**Modern aspects of artificial intelligence applications in Business Information Systems
and in particular ERP systems**

Димитър Димитров¹
e-mail: dimitar.dimitrov@unwe.bg¹

Абстракт

В съвременния силно свързан свят, софтуерните системи за управление на бизнеса (ERP) са под голям натиск за постоянни подобрения и усъвършенстване, за да отговорят на нуждите на бизнеса. Бумът, който предизвика появата и бурното навлизане на изкуствения интелект (AI) в сферите на маркетинга, потребителското изживяване, кибер- сигурността, не подмина и разработчиците на бизнес софтуер. Усъвършенстването на големите езикови модели (LLM) позволи на водещите ERP производители - SAP, Oracle и Microsoft да стартират разработването на собствени агенти и модели. Интересен е подходът на всеки един от големите, тъй като всеки един от тях показва различен поглед към AI и как той да бъде имплементиран към техните core системи - S/4 Hana, Fusion engine, Dynamics 365.

Abstract

In today's highly connected world, enterprise resource planning (ERP) software systems are under great pressure to constantly improve and evolve to meet the needs of businesses. The boom that has caused the emergence and rapid penetration of artificial intelligence (AI) in the fields of marketing, user experience, and cybersecurity has not passed by the developers of business software. The improvement of large language models (LLM) has allowed the leading ERP manufacturers - SAP, Oracle, and Microsoft to start developing their own agents and models. The approach of each of the major ones is interesting, as each of them showed a different view of AI and how it can be implemented in their core systems - S/4 Hana, Fusion engine, Dynamics 365.

JEL: C88, L86

Увод

Изкуственият интелект (AI) и по-специално генеративният AI (GenAI) променят начина, по който се проектират и използват бизнес информационните системи. Докато през последното десетилетие AI често се интегрира като допълнителни модули, през периода 2023–2025 г. наблюдаваме трансформация към вграждане на AI като първостепенна функционалност в ERP системите. Това налага преосмисляне на архитектурни решения, операционни практики и управление на доверието. Целта на настоящия доклад е да представи съвременните тенденции, да опише подходите на трите основни доставчици (SAP, Oracle, Microsoft) и да обсъди бъдещи перспективи и предизвикателства.

Методология и рамка на анализа

Анализът е направен чрез преглед на официални технически документи, бели книги (white papers), продуктови спецификации на доставчиците и рецензирана литература до 2024 г., допълнен с индустриални доклади и практическата им интерпретация. Рамката за оценка включва: стратегия за вграждане на AI, ключови функции, техническа архитектура и

¹ Докторант в катедра ИТК, УНСС, e-mail: dimitar.dimitrov@unwe.bg

операционна интеграция, възможности за разширяемост, и механизми за управление на регулаторните рамки.

Общи тенденции в приложенията на AI в бизнес софтуера

AI еволюира от отделни аналитични и експериментални решения към вградени транзакционни услуги, които участват пряко в бизнес процесите (напр. автоматично попълване на финансови транзакции, автоматично класифициране на документи и прогнозиране на потребности).

GenAI позволява натуралноезикови интерфейси, генериране на съдържание (например чернови на договори, имейли, резюмета на отчети) и автоматизация на документооборота. Ключов ефект от това е намаляване на времето за ръчни операции и ускоряване на вземането на решения.

Като области на масово приложение може да се определят следните:

- финанси - прогнозиране на парични потоци, откриване на аномалии,
- снабдяване и вериги на доставки - прогнози за търсене, оптимизация на инвентара,
- производство - предвидима поддръжка (predictive maintenance),
- HR - оценка на таланти, автоматизирани процеси по наемане,
- обслужване на клиенти – чат и гласови ботове, автоматично маршрутизиране на случаи.

В същото време растящите регулаторни изисквания налагат повече рестрикции върху поверителността, прозрачността и верността на изходните резултати от AI системите. В отговор доставчиците изграждат механизми за проследяване на произхода на моделите, одитни следи и функционалности за обяснимост.

От гледна точка на архитектурата тенденциите са ориентирани към облачните, API-first и microservices архитектури, заедно с инструменти за управление на жизнения цикъл на моделите (MLOps). Това се превръща в стандарт за интеграция и експлоатация на AI функционалности.

Подходът на SAP

SAP се фокусира върху вграждане на AI в S/4HANA и използване на SAP Business Technology Platform (BTP) като основен слой за интеграция и разработка. Стратегията е да предостави контекстно прецизирани функции, оптимизирани за бизнес процесите по индустрии.



Source: <https://blog.devgenius.io/sap-ai-770a5ba8a710>

Фигура 1: SAP

Ключови възможности

SAP Joule предлага разговорен и контекстно-свързан асистент, интегриращ данни от транзакции и аналитични слоеве. Други възможности включват прогнозна аналитика (predictive analytics), интелигентна обработка на документи (IDP) и RPA/Process Automation за автоматизация на повтарящи се задачи.

Технически аспекти

SAP оптимизира големи езикови модели (LLM) и обработки за HANA in-memory база, използва микросървиси на BTP и партнира с хиперскейлерите (AWS, GCP, Azure) за големи изчислителни натоварвания и LLM хостинг. BTP предоставя инструменти за изграждане и внедряване на модели, API мениджмънт и интеграция с корпоративни данни.

Управление, доверие и съответствие

SAP предлага опции за локализиране на данни (data residency), контрол на достъп чрез ролеви модели, мониторинг на представянето на моделите и управление на техния жизнен цикъл. Фокусът е върху индустриализация на AI с внимание към обяснимост и проследяване.

Подходът на Oracle

Oracle позиционира AI като вграден компонент на Fusion Cloud ERP, разчитащ основно на Oracle Cloud Infrastructure (OCI) и интеграцията с автономни бази данни и инфраструктура.



Source: <https://boxplot.com/oracle-and-ai-innovation>

Фигура 2: Oracle

Ключови възможности

Oracle предоставя дигитални асистенти, интелигентна обработка на документи, адаптивни приложения и вграден анализ. Интеграцията с автономни бази данни осигурява висока степен на автоматизация в управлението на данните и операции.

Технически аспекти

Основната сила е използването на OCI за хостинг на модели и услуги за AI, включително специфични Oracle AI услуги и инфраструктурна автоматизация. OCI дава инструменти за изграждане, обучение и внедряване на персонализирани модели.

Управление и доверие

ОСИ включва вграден набор от механизми за сигурност, одитиране и проследяване на произхода на моделите. Oracle акцентира върху автономните възможности за управление на данните и инфраструктурата, което улеснява съответствието и оперативната надеждност.

Подходът на Microsoft

Microsoft следва AI-first подход, като интегрира Copilot и Azure AI в Dynamics 365 и Power Platform за осигуряване на разговорен UX, low-code автоматизация и аналитика.



Source: <https://boxplot.com/>

Фигура 3: Microsoft AI

Ключови възможности

Copilot в Dynamics 365 предоставя разговорен интерфейс за ERP/CRM сценарии; Power Automate с AI Builder улеснява автоматизацията; Microsoft Fabric и Azure AI предлагат аналитична и ML инфраструктура.

Технически аспекти

Интеграция с Azure OpenAI Service, Dataverse като унифициран модел на данни и Microsoft Fabric за анализ. MLOps възможности в Azure DevOps/ML Ops подпомагат управлението на жизнения цикъл на моделите.

Управление и доверие

Microsoft Purview и наборът от инструменти за отговорен AI предоставят механизми за управление на данни, поверителност и съответствие на ниво tenant и организация. Силната low-code екосистема улеснява бързото внедряване от бизнес потребители.

Сравнителен анализ: прилики и различия

Прилики

- Вграждане на AI директно в ERP работните потоци, а не само като add-on.
- Използване на облачни платформи и партньорства за хостинг и изчисления.
- Предварително изградени възможности: IDP, прогнозиране, детекция на аномалии.
- Фокус върху управление на данни, обяснимост и съответствие.

- Разширяемост чрез API, платформи за разработка и ISV екосистеми.
- Постепенно разгръщане на GenAI копилоти за повишаване на продуктивността

Разлики

- Oracle залага на собствени инфраструктурни решения и автономни функции (OCI).
- SAP акцентира върху индустриални процеси, real-time аналитика и оптимизация чрез HANA.
- Microsoft се отличава с мощни low-code/no-code възможности и широка екосистема за самостоятелно обслужване (self-service).

Таблица 1: Сравнителна таблица на SAP, Oracle и Microsoft по ключови параметри

Параметър	SAP	Oracle	Microsoft
Стратегия за AI	Вграждане в S/4HANA и BTP; фокус върху индустриални процеси и real-time аналитика.	AI навсякъде в Fusion Cloud ERP; опора на OCI и автономни операции.	AI-first интеграция чрез Copilot, Dynamics 365, Azure AI и Power Platform; фокус върху self-service и low-code.
Ключови AI възможности	SAP Copilot, прогнозна аналитика, интелигентна обработка на документи, автоматизация на процеси.	Дигитални асистенти, интелигентна обработка на документи, адаптивни приложения, интеграция с автономни БД.	Copilot в Dynamics, Power Automate + AI Builder, обработка на документи, прогнозни анализи.
Техническа архитектура	HANA-оптимизирани модели; BTP микросървиси; партньорства с хиперскейлъри за LLM/compute.	Модели и услуги хоствани в OCI; автономни бази данни и инфраструктурна автоматизация.	Интеграция със Azure (Azure OpenAI, Fabric); Dataverse като единен модел на данни; силна MLOps поддръжка.
Разширяемост и интеграция	BTP инструменти за разработка и предварително бизнес съдържание; API-first подход; ISV екосистема.	OCI инструменти за персонализирани модели; API и интеграция с Fusion Cloud; ISV/partner marketplace.	Силни low-code/no-code (Power Platform); широк ISV екосистем; богат набор от SDK и API.
Управление на модели и MLOps	Инструменти за мониторинг на модели, управление на жизнения цикъл, локализация на данни.	OCI MLOps услуги; вградено одитиране и проследяване на произхода на моделите.	Azure MLOps, DevOps интеграция, мониторинг и управление на версиите; Power Platform ALM.
Сигурност и съответствие	Опции за data residency; ролев контрол; enterprise security интеграции; акцент върху обяснимост.	Вградени OCI мерки за сигурност; автономни операции за по-малко човешки фактори; одит и проследяване.	Microsoft Purview, tenant-level контрол, сертификати за съответствие и инструменти за отговорен AI.
Подход към GenAI копилоти	Контекстно-осъзнати копилоти (Copilot) в рамките	Digital assistant функционалности;	Copilot фокус с разговорен UX, тесно свързан с Power

	на бизнес процесите; интеграция с транзакционни данни.	интеграция с приложни потоци и данни в OCI.	Platform за автоматизации и потребителски сценарии.
Предпочитани клиенти/сценарии	Големи предприятия с нужда от индустриални процеси и real-time аналитика.	Организации, желаещи автономни операции и силна OCI интеграция; корпоративни среди с акцент върху базите данни.	Организации, които търсят бързо внедряване, self-service и low-code възможности; интеграция с Microsoft 365.
Рискове и ограничения	Зависимост от HANA; сложност при интеграции с legacy системи; усилия за управление на данни.	Vendor lock-in към OCI; управление на чувствителни данни при cloud-first подход; сложност в кастомизации.	Зависимост от Azure и Microsoft екосистема; потенциални ограничения при специфични индустриални изисквания.
Типични случаи на използване	Реално-времева аналитика за производство; интелигентна обработка на поръчки; автоматизация по индустрии.	Автоматизация на финансови операции; интелигентна обработка на счетоводни документи; автономно управление на БД.	Автоматизация на продажби и обслужване; бързи low-code решения за работни потоци; Copilot за служители.

Влияние върху пазара и доставчиците

Интеграцията на AI създава по-силна зависимост на клиентите към доставчиците (vendor lock-in) чрез дълбока обвързаност с платформените данни, модели и автоматизации. Решението за избор между доставчиците зависи от съществуващите платформи, нуждата от индустриален контекст и предпочитания за low-code възможности.

Бизнес и технически предизвикателства

Управление на данни и качество

AI изисква качествени, структурирани и подходящо етикетирани данни. В ERP среда данните са разпределени между модули и системи, което налага инвестиции в данни, интеграция и управление на метаданни.

Моделна поддръжка и MLOps

Жизненият цикъл на моделите - обучение, валидация, мониторинг, ре-тренинг - изисква MLOps практики, които са сравнително нови за традиционните ERP администрации. Това включва наблюдение на drift, метрики за производителност и стратегии за управление на версии.

Сигурност и поверителност

AI-функции в ERP обработват чувствителни финансови и персонални данни. Необходима е криптография, контрол на достъп, изолация на данни и технически мерки за предотвратяване на изтичане на данни в обучения на външни модели.

Обяснимост и юридическа отговорност

Генерирането на автоматични решения и препоръки налага инструменти за обяснимост и документиране на причините за решенията (explainability). Регулациите изискват проследимост на решенията и възможности за одит.

Управление на промяната и човешки фактор

Интегрирането на AI променя работните процеси и изисква обучение, адаптация на роли и смяна в организационната култура. Съпротивата и липсата на доверие могат да забавят внедряването.

Бъдещи посоки

- Дълбоко вградени GenAI копилоти - Очаква се копилоти, които не само отговарят на заявки, но и проактивно оркестрират процеси през различни модули (например автоматично стартиране на доставки, проверка на кредитни лимити и одобрения).
- Интеграция на LLM с транзакционни системи в реално време - Подобряване на контекстната прецизност и скорост при обработка на големи масиви транзакционни данни в реално време.
- AI за оперативни решения с минимална латентност - В производството и логистиката ще се внедряват решения за вземане на решения в реално време (edge AI, стрийминг аналитика).
- Нови пазари и архитектури - Възникване на корпоративно-адаптирани LLM пазари, специални Data Lakes и синтетични данни за частно обучение и защита на поверителността.
- Засилени регулации и нови професии като CDO – Chief Data Officer
- Още по-строги регулаторни изисквания, трайни одитни следи и методи като частно (federated) обучение, диференциална поверителност и синтетични данни.
- Инвестиции в интеграция на данни, качество и метаданни управление (data governance), включително дефиниране на master data и data lineage.

Заклучение

AI и GenAI трансформират ERP системите от пасивни транзакционни платформи към интелигентни оперативни слоеве, които подпомагат вземането на решения и автоматизацията. SAP, Oracle и Microsoft следват сходни стратегически направления, но се различават по архитектурна реализация, подход към разширяемост и low-code възможности. Бъдещето ще донесе по-дълбоко вграждане на GenAI, по-силна интеграция с транзакционни потоци в реално време и по-строги регулации. Организациите трябва да планират стратегически, да инвестират в данни и MLOps и да изграждат доверие и съответствие, за да извлекат ползи от тези технологии.

References

1. Al-Jarrah, O., Yoo, P., & Smith, R. (2022). AI-Driven ERP Systems: Trends and Challenges. *Journal of Enterprise Information Management*, 35(4), 567–585.
2. Gartner. (2023). Market Guide for AI in ERP. Gartner Research. Retrieved from <https://www.gartner.com>
3. Microsoft. (2024). Microsoft Dynamics 365 and Copilot: Product Overview. Microsoft White Paper. Retrieved from <https://www.microsoft.com>
4. Oracle. (2024). Oracle Fusion Cloud ERP: AI and Autonomous Operations. Oracle White Paper. Retrieved from <https://www.oracle.com>
5. SAP. (2024). SAP S/4HANA and SAP Business Technology Platform: AI Strategy. SAP White Paper. Retrieved from <https://www.sap.com>
6. IBM Institute for Business Value. (2021). The enterprise AI playbook. IBM. Retrieved from <https://www.ibm.com>

7. Kalliamvakou, E., et al. (2021). MLOps: Continuous delivery and automation pipelines in machine learning. *Proceedings of the ACM on Programming Languages*, 4(OOPSLA), 1–29.
8. European Commission. (2021). Proposal for a Regulation laying down harmonised rules on artificial intelligence (AI Act). Retrieved from <https://eur-lex.europa.eu>
9. Bender, E. M., & Gebru, T. (2021). On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? *Proceedings of the ACM Conference on Fairness, Accountability, and Transparency*, 610–623.
10. Zhang, Y., & Lu, H. (2023). Generative AI in enterprise workflows: Opportunities and risks. *International Journal of Information Management*, 68, Article 102509.
11. Redman, T. C. (2018). The Impact of Poor Data Quality on ERP Systems. *Data Quality Journal*, 12(2), 45–58.
12. OECD. (2023). Recommendation on AI and Governance. OECD Publishing.
13. Srivastava, A., & Kumar, P. (2022). Document Intelligence in ERP: Techniques and Implementations. *Enterprise AI Journal*, 2(1), 23–39.
14. Boden, M., & McCulloch, J. (2020). Responsible AI Practices in Enterprise Software. *Journal of Business Ethics*, 162(3), 607–620.

BLOCKCHAIN AND ARTIFICIAL INTELLIGENCE FOR INTELLIGENT AND SECURE FINANCIAL TRANSACTIONS: AN IMPLEMENTATION PERSPECTIVE WITH PYTHON

Vasil Spasov¹,

e-mail: vasil.spasov@unwe.bg¹

Абстракт

Традиционните финансови системи страдат от бавни междубанкови преводи, високи такси и рискове от измами, като централизираните им архитектури създават критични уязвимости. Настоящото изследване предлага архитектура за децентрализирана, интелигентна система за финансови транзакции, базирана на блокчейн и изкуствен интелект, като илюстрира практическата ѝ реализация чрез Python. Целта е система със значително намаляване на транзакционните разходи и автоматизирана детекция на измами. Python е подходящ за тази цел, благодарение на богатата си екосистема от библиотеки.

Abstract

Traditional financial systems suffer from slow interbank transfers, high fees, and fraud risks, with their centralized architectures creating critical vulnerabilities. This study proposes an architecture for a decentralized, intelligent financial transaction system based on blockchain and artificial intelligence, demonstrating its practical implementation using Python. The goal is a system that significantly reduces transaction costs and enables automated fraud detection. Python is suitable for this purpose due to its rich ecosystem of libraries.

Key words: Blockchain, Artificial Intelligence, Python, Financial Transactions, Machine Learning

JEL: G21, C88, O33

1. Introduction

Contemporary financial systems rely on centralized intermediary institutions, resulting in slow and expensive transactions, as well as increased vulnerability to cyber attacks. Traditional methods for combating fraud are no longer adequate against modern threats.

Blockchain technology offers a decentralized and transparent alternative approach, eliminating intermediaries through cryptographically secure mechanisms. Simultaneously, artificial intelligence provides capabilities for big data analysis and real-time risk prediction, complementing blockchain solutions.

Python emerges as a suitable tool for implementing this convergence with its rich ecosystem of libraries: scikit-learn for machine learning, web3.py for Web3 development, Vyper for smart contracts. This report illustrates why Python is an excellent choice for creating hybrid blockchain-AI systems

Research Objective:

To propose a system architecture that combines blockchain and artificial intelligence for secure financial transactions and to demonstrate its practical implementation using Python.

¹ PhD, Faculty of Finance, University of National and World Economy

This research aims to design and develop an integrated system architecture that effectively combines blockchain technology with artificial intelligence to create a secure, intelligent, and efficient framework for financial transactions. The solution that harnesses the complementary strengths of both technologies—leveraging blockchain's inherent advantages of decentralization, immutability, and cryptographic security while integrating AI's advanced capabilities in machine learning, predictive analytics, and adaptive pattern recognition. The primary focus is on building a functional system that addresses critical challenges in modern financial operations, including real-time fraud detection. The goal is to provide a scalable, adaptable model that bridges theoretical innovation with practical application, offering significant improvements in financial security, cost reduction, and system intelligence while maintaining the flexibility to evolve with emerging technological standards and regulatory requirements.

Methodology:

The research employs an analytical approach for designing a system architecture that integrates blockchain networks (Ethereum), smart contracts (Vyper), scikit-learn machine learning algorithms, and decentralized storage (IPFS). The methodology includes a review of relevant literature, conceptual design, and evaluation of practical applicability through Python.

2. Literature Review

Blockchain technology, introduced with the Bitcoin protocol by Nakamoto (2008), represents a decentralized ledger with fundamental characteristics: immutability, transparency, and absence of centralized control. The Ethereum platform by Buterin (2014) expands this concept through smart contracts - self-executing programs that automate contractual agreements without intermediaries.

In the field of artificial intelligence for financial security, West and Bhattacharya (2016) provide a comprehensive review of intelligent fraud detection methods, demonstrating the superiority of algorithms over traditional systems. Modern machine learning algorithms recognize fraudulent transactions with high accuracy.

Liang et al. (2025) extend the research in blockchain context, showing how machine learning detects Ponzi schemes in Ethereum through analysis of behavioral patterns of smart contracts and transactions.

The synergy between blockchain and AI has been analyzed by Salah et al. (2019), who identify how blockchain addresses AI challenges related to model transparency and data trust, while AI enhances the efficiency of blockchain systems through intelligent optimization of consensus mechanisms and automated anomaly detection.

From an implementation perspective, Python dominates machine learning, with Pedregosa et al. (2011) documenting scikit-learn as a powerful library providing a rich set of algorithms. This Python ecosystem extends to blockchain technologies through the Vyper language, which offers Python-like syntax for smart contract development with focus on security and code auditability.

Wang et al. (2021) examine the NFT potential for digitization of unique assets and documents in the financial sector, including securities, contracts, and proof of ownership.

Scalability challenges have been analyzed by Croman et al. (2016), who identify fundamental limitations of blockchain networks in processing high volumes of transactions and propose second-layer solutions. Regulatory aspects are the subject of Finck's (2019) research regarding the tension between blockchain immutability and GDPR requirements, proposing approaches for reconciliation through techniques such as cryptographic erasure and off-chain storage of personal data.

3. Python-Based Architecture for Blockchain-AI Financial Systems

The proposed architecture integrates blockchain and artificial intelligence using Python as a central tool. The system combines decentralized governance with intelligent automation and multi-layered security.

The blockchain layer is implemented through the web3.py library, which serves as a standard interface for Python applications to interact with Ethereum and compatible blockchain networks. The library

provides complete functionality not only for executing transactions and working with smart contracts, but also for real-time event monitoring. For smart contract development, Vyper is used - a Python-like programming language specifically designed for writing secure blockchain applications, which eliminates complex and risky constructs, making smart contracts safer and easier to create.

The machine learning module analyzes transactions in real-time to detect fraud. To assess risk, the module analyzes historical data and multiple parameters such as transaction value and frequency, geographic location, and user behavioral patterns.

For storing large volumes of data that cannot be efficiently recorded directly on the blockchain due to high costs and memory limitations, the system integrates IPFS (InterPlanetary File System). This way, only encrypted data hashes are recorded on the blockchain, while the actual files are stored in the decentralized IPFS network. This hybrid approach ensures data security and immutability while maintaining efficiency and low operational costs.

Before any financial transaction is finally executed, it undergoes an automated risk assessment system based on artificial intelligence. This validation module uses a pre-trained machine learning model created and trained on extensive historical data from past transactions, including both legitimate and fraudulent operations. The model continuously learns from new data to improve its accuracy and adapt to new fraud patterns.

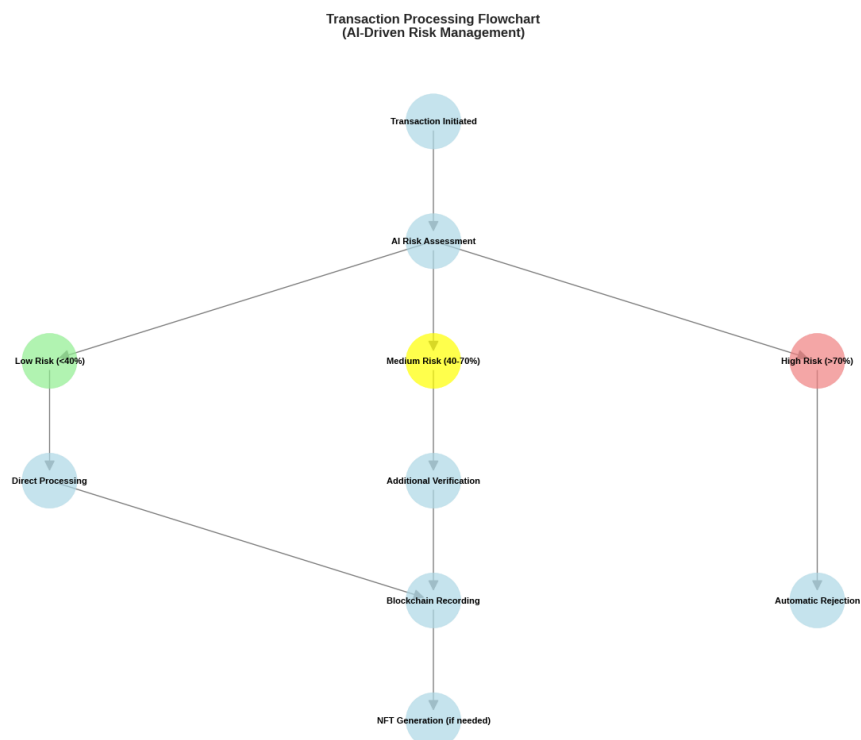
The system analyzes multiple parameters in real-time to generate a probabilistic risk assessment for fraud. Among the analyzed factors are: the transaction amount and its deviation from the user's typical values; the frequency of transactions and their temporal patterns; the geographic location of the operation and its compatibility with previous locations; the user's behavioral patterns and their consistency; the device type and IP address used for the transaction; as well as historical data about counterparties and their risk profiles.

For final decision-making, the system uses a three-tier threshold system that balances security and user convenience. For transactions with extremely high risk (over 70% probability of fraud), the system automatically blocks the operation without human intervention. For medium-risk transactions (between 40% and 70%), the system flags them for additional verification by security specialists or requires additional identity authentication methods. Low-risk transactions (below 40%) are processed directly and immediately without additional delays.

After a transaction successfully passes all checks and is executed, it is recorded on the blockchain as a permanent and immutable record. Each record contains a unique cryptographic hash that serves as a digital fingerprint of the transaction. This hash is generated through cryptographic algorithms and guarantees that even a minimal change in the transaction data will result in a completely different hash. All network participants can verify and authenticate the transaction through this hash, ensuring complete transparency and traceability.

In specific cases where the transaction involves ownership transfer, fulfillment of contractual conditions, or creation of a digital asset, the system automatically generates a Non-Fungible Token (NFT) through a specialized ERC-721 smart contract. This NFT serves as digital proof of ownership or contract fulfillment. The token contains metadata describing the specific asset or condition and is stored in the owner's wallet. The metadata may include a detailed description of the asset, deal terms, creation timestamp, as well as history of ownership transfers. This mechanism provides secure, transparent, and verifiable proof of rights and obligations arising from the financial transaction.

The system implements a continuous learning mechanism for the AI model, where new data from transactions verified as legitimate or fraudulent is periodically collected, and the model is retrained to adapt to evolving fraud patterns. This process can be automated when a certain volume of data is reached.



Source: Author's development

Figure 1: Transaction Processing Flowchart. AI Risk Management

4. Conclusion

The study demonstrates the implementation of a hybrid system for secure financial transactions using Python, integrating blockchain and artificial intelligence. Tools such as web3.py, Vyper, and scikit-learn were employed, providing a comprehensive solution with fundamental advantages: reduced costs and processing time, automated fraud detection, and complete transaction traceability.

The selection of Python for implementing an intelligent financial transaction system is not only technologically feasible but also strategically justified, opening opportunities for innovation, broad accessibility, and easy integration across multiple domains. The combination of blockchain's immutability and transparency with AI's intelligent analytical power, implemented through Python's versatile ecosystem, represents a genuine paradigm shift toward a more efficient financial system.

References

1. Nakamoto, S. (2008). "Bitcoin: A Peer-to-Peer Electronic Cash System." Bitcoin.org.
2. Buterin, V. (2014). "A Next-Generation Smart Contract and Decentralized Application Platform." Ethereum White Paper.
3. West, J., & Bhattacharya, M. (2016). "Intelligent Financial Fraud Detection: A Comprehensive Review." *Computers & Security*, 57, 47-66.
4. Liang, R., et al. "(2025). Towards Effective Detection of Ponzi Schemes on Ethereum with Contract Runtime Behavior Graph. *ACM Transactions on Software Engineering and Methodology*, 34(4), 106, 1-32. <https://doi.org/10.1145/3707458>
5. Salah, K., et al. (2019). "Blockchain for AI: Review and Open Research Challenges." *IEEE Access*, 7, 10127-10149.
6. Croman, K., et al. (2016). "On Scaling Decentralized Blockchains." *Financial Cryptography and Data Security*, 106-125.

7. Finck, M. (2019). "Blockchain and the General Data Protection Regulation." European Parliamentary Research Service. Scientific Foresight Unit (STOA) PE 634.445 – July 2019
8. Wang, Q., et al. (2021). "Non-Fungible Token (NFT): Overview, Evaluation, Opportunities and Challenges." arXiv preprint arXiv:2105.07447.
9. Web3.py Documentation. <https://web3py.readthedocs.io>
10. Vyper Documentation. <https://vyper.readthedocs.io>
11. Pedregosa, F., et al. (2011). "Scikit-learn: Machine Learning in Python." JMLR, Journal of Machine Learning Research, 12, 2825-2830.

AI IN OFFENSIVE CYBERSECURITY: AUTONOMOUS TTP EMULATION AND SCALABLE CONTINUOUS TESTING

**Изкуствен интелект в офанзивната киберсигурност: автономна емуляция на TTP
и мащабируемо непрекъснато тестване**

Сали Салиев¹

e-mail: sali.saliev@unwe.bg¹

Абстракт

Офанзивното тестване постепенно се измества от единични ръчни упражнения към рутинни, автоматизирани проверки. В този доклад представяме система, която използва големи езикови модели за планиране и агенти за контролирано изпълнение на TTP-та върху тестови среди, близки до корпоративни мрежи. Подходът обхваща разузнаване, начален достъп, странично придвижване и въздействие, но работи под ясни политики за риск, ограничения на обхвата и „авариен стоп“. Вграден телеметричен поток превежда техническите находки в управленски метрики за експозиция (критичност на активи, радиус на въздействие, дължина на пътя на атаката). В лабораторни сценарии системата достига съпоставими резултати с човешки red team екипи и намалява времето до откритие с 35–60% според случая, като поддържа нощни „репетиции на атаки“ за проследяване на сигурностния дрейф. Обсъждаме практични похвати за промпт инженерство, синтетични данни за закаляване на моделите и контрамерки срещу халюцинации, излизане извън обхват и OPSEC изтичания. Резултатът е прагматичен модел на „човешко целеполагане + машинен мащаб“ за непрекъсната валидация на контроли и аналитика, готова за решения.

Abstract

Offensive testing is shifting from ad-hoc, manual exercises to routine, automated checks. We present a system that uses large language models for planning and tool-driven agents for controlled TTP execution against corporate-like lab environments. The approach spans reconnaissance, initial access, lateral movement, and impact while operating under explicit risk policies, scope limits, and an emergency stop. A built-in telemetry pipeline turns technical findings into management-ready exposure metrics (asset criticality, blast radius, attack-path length). In lab scenarios, the system delivers results comparable to human red teams and reduces time-to-finding by 35–60% depending on the case, while enabling nightly "attack rehearsals" to track security drift. We detail prompt-engineering tactics, synthetic-data hardening, and countermeasures against hallucinations, scope creep, and OPSEC leaks. The outcome is a pragmatic "human intent + machine scale" model for continuous control validation and decision-ready analytics.

Keywords: offensive security, red team, autonomous agents, TTP emulation, continuous testing, LLM-tools, exposure analytics, safety policies.

JEL: C55, C63, M15, O33

Introduction

Traditional penetration testing and red-team exercises are episodic, expensive, and quickly outdated as environments change. As controls evolve and drift, defenders need continuous, repeatable, and business-aligned validation that reflects real TTPs rather than synthetic checkbox tests.

This paper proposes an AI-assisted offensive testing framework that combines:

¹ Главен експерт, веб системи, Дирекция "Учебна дейност", УНСС, e-mail:sali.saliev@unwe.bg

LLM-based planning to translate high-level testing goals and constraints into actionable steps aligned with MITRE ATT&CK;

Controlled agents that execute those steps via vetted tools within isolated runtimes;

A policy layer that constrains scope, rate, and impact, and provides auditable guardrails; and

Telemetry & exposure analytics that convert technical artifacts into management-ready metrics (asset criticality, blast radius, and attack-path length).

Contributions.

A modular architecture for safe, controllable TTP emulation at machine scale;

An evaluation methodology and metric suite that compare AI-assisted execution to human baselines;

A telemetry design that yields decision-ready exposure analytics;

Practical safeguards for generative-AI risks in offensive testing (prompt injection, excessive agency, OPSEC).

Related Work and Industrial Context

ATT&CK-aligned emulation. MITRE ATT&CK provides a shared vocabulary for tactics, techniques, and procedures and underpins numerous emulators and test suites used to validate controls. Representative tools include autonomous or scriptable breach-and-attack simulations and lightweight technique-level test packs that map directly to ATT&CK.

Identity/graph analysis. In enterprise domains (Active Directory/Entra/Azure), graph-based analysis of identities and privileges is widely used to reveal attack paths—shortest-path escalations to high-value assets—which correspond closely to business risk.

Continuous exposure management. Recent practice emphasizes continuous threat exposure management: discovering, prioritizing, and validating exposures in context, rather than treating vulnerabilities as flat lists. This aligns with the need for repeatable rehearsals to detect security drift.

LLM tool-use and safety. Emerging patterns (ReAct-style reasoning-and-acting, multi-agent orchestration) enable LLMs to coordinate tools. In parallel, AI risk frameworks and secure-by-design guidance stress mitigations for prompt injection, insecure output handling, excessive autonomy, and data leakage. Our architecture bakes these mitigations into the policy layer.

System Architecture

The framework comprises four layers: an Orchestrator (LLM planner), Execution Agents (tools), a Policy Layer (safety & scope), and Telemetry & Exposure Analytics.

Orchestrator (LLM → goals, constraints, plan)

- Accepts goals and Rules of Engagement (ROE): target scope, no-go lists, risk caps, time/traffic budgets.
- Produces a stepwise plan annotated with ATT&CK tactics/techniques, expected signals, exit/stop criteria, and maximum depth.
- Uses a reason-and-act pattern: only calls tools when required, with explicit input/output schemas; sensitive actions require human-in-the-loop confirmation.

Execution Agents (strictly controlled)

- Discovery: network/service/OS discovery via standard scanners; passive options where possible.
- Exposure checks: misconfiguration and known-issue probes using curated templates and allow-listed targets.
- Identity & paths: directory/graph collection to model privilege relationships and shortest paths to crown jewels.

- ATT&CK emulation: controlled, low-impact executions of ATT&CK-mapped techniques for validation.
- Each agent runs in an isolated container with limited permissions, rate limits, timeouts, and outbound controls; all I/O is logged for audit.

Policy Layer (safety and governance)

- Scope control: allow/deny lists, network segments, identities, data classes.
- Risk caps: execution budgets (time/packets/requests), privilege ceilings, impact thresholds.
- Emergency stop: immediate cancellation if a guardrail triggers (scope breach, anomaly, excessive error rate).
- AI guardrails: prompt filtering, role separation, context isolation between agents, verification gates for tool outputs, and redaction of sensitive data before model exposure.
- Auditability: immutable logs, signed artifacts, and run manifests for review.

Telemetry & Exposure Analytics

- Normalize events to ATT&CK techniques and assets; store in a time-series and graph store.
- Compute management metrics: Asset criticality; Attack-path length; Blast radius; Time-to-Finding (TTF); Repeatability; FP/FN rates.
- Output decision-ready dashboards and diffs across nightly runs to reveal security drift.

Evaluation Methodology

Test Environment

Enterprise-like lab with segmented networks, directory services (AD/Entra/Azure), a small cloud account, representative endpoints/servers, and seeded low-risk misconfigurations and identity pathologies. Each seeded exposure is mapped to an ATT&CK technique and tagged with criticality to support business-aligned scoring.

Metrics

- Time-to-Finding (TTF): median/mean time to detect a seeded exposure.
- TTP coverage: fraction of techniques/sub-techniques within the declared scope that are exercised/detected.
- Repeatability: variance of TTF and success rate across N nightly runs.
- False positives/negatives: precision/recall of exposure identification.
- Operational cost: compute minutes, network budget consumed, and analyst time.
- Path-risk metrics: average attack-path length to crown jewels; expected blast radius given current identity/permission graph.

Baselines

Human baseline: experienced red-teamers (or scripted playbooks) operating under the same ROE. Tool-only baseline: manual execution of emulator/test packs without LLM planning. Comparative analysis: paired scenarios; statistical tests (e.g., paired t-tests) on TTF and repeatability.

Validity, Ethics, and ROE

Clear ROE and legal authorization; no production testing without explicit approval. Safety checks in staging first; impact-minimizing modes for emulation; documented rollback. Reproducibility: fixed seeds where applicable, version-pinned toolchains, and published run manifests.

Tools and Environment

Tooling roles

- Planner/Orchestrator: LLM-driven reasoning-and-acting with structured tool calls; optional multi-agent coordination for decomposition and verification.
- Discovery: standard network/service discovery for inventory and scoping.
- Exposure checks: template-driven probes (curated sets with severity, references, and safe defaults).
- Identity/graph: collectors that build a privilege graph for path analysis and blast-radius estimation.
- ATT&CK emulation: autonomous or scripted technique-level emulation suites to validate detections and response playbooks.

Runtime and isolation

- Containerized agents with least privilege, network egress controls, timeouts, and rate limits; per-tool sandboxes prevent cross-contamination.
- Central logging & audit: structured logs, artifacts, and signed manifests; sensitive data minimized or anonymized.
- Configuration as code: version-controlled ROE, allow/deny lists, and template sets; CI/CD for lab updates.

AI guardrails and data policy

- Prompt safety: input filtering, role separation, and explicit tool-I/O schemas; no implicit code execution.
- Output validation: sanity checks and human approval gates for sensitive or environment-changing actions.
- Data handling: redaction/anonymization before model exposure; prohibition on sending production secrets to external services; retention policies aligned with audit requirements.

Results

Environment & Scope

Run ID: 2025-10-07_2054

Operator: Sali (owner)

Environment: Isolated lab; macOS (Docker) for the app; Kali VM as attacker box

Targets: 192.168.182.108:5173 (Vite dev), 192.168.182.108:3000 (Express dev)

Scope & ROE (lab): owner-controlled assets only; no Internet/external scope; safe/non-destructive techniques; no auth brute-force; no DoS.

Goal: collect reproducible evidence and artifacts for conference documentation.

Executive Summary

- Host reachable; services responsive on 5173 and 3000.
- F1 (Medium): Vite Dev Server @fs Path Traversal / LFI (CVE-2025-31125) — confirmed.
- F2 (Medium): Express app in development mode — stack traces and X-Powered-By disclosed; permissive CORS.
- ffuf (quick): no interesting paths in this pass.
- ZAP Baseline (3-min, passive): 0 FAIL, 7 WARN, 60 PASS (header hygiene on dev server).

Recon

```
nmap -sV -Pn --top-ports 100 192.168.182.108
```

Content discovery

```
ffuf (quick wordlists) → 192.168.182.108:5173
```

Templated checks

```
nuclei (severity: medium,high,critical) → :5173 and :3000
```

Passive DAST

```
ZAP Baseline (3-minute spider) → http://192.168.182.108:5173
```

Table 1. Summary of Identified Vulnerabilities and Recommended Mitigations

ID	Asset	Severity	Description	Evidence (short)	Impact	Recommended Mitigations
F1	http://192.168.182.108:5173 (Vite dev)	Medium	@fs path traversal / LFI (CVE-	Nuclei matched; 200 OK with JS wrapper	Local file disclosure while dev server is	Upgrade Vite; do not expose dev server; serve production dist/

			2025-31125)	containin g base64-embedde d file content via crafted @fs URL	LAN- exposed	behind hardened web server
F2	http://192.168.182.108:3000 (Express)	Medium	Running in developme nt mode	Nuclei node- express- dev-env matched; stack traces; X- Powered- By header; permissiv e CORS	Information disclosure and easier fingerprinti ng if reachable	Set NODE_ENV=pr oduction; disable detailed errors; remove/mask X- Powered-By; restrict CORS

Methodology

Finding Details

F1 — Vite Dev Server @fs Path Traversal (LFI)

Asset: http://192.168.182.108:5173

Severity: Medium (information disclosure / arbitrary file read)

Evidence: Nuclei CVE-2025-31125 matched; server responded 200 OK with JS bundle containing base64-embedded file content from a crafted @fs URL.

Sample matched URL (redacted for lab):
http://192.168.182.108:5173/@fs/etc/passwd?import&?inline=1.wasm?init

Benign PoC to read /etc/hosts (captured artifact):

```
curl -s "http://192.168.182.108:5173/@fs/etc/hosts?import&?inline=1.wasm?init" \
```

```
| python3 -c 'import sys,re,base64;d=sys.stdin.read();m=re.search(r"base64,([^\"]+)",d);print(base64.b64decode(m.group(1)).decode("utf-8","ignore"))'
```

Impact: Local file disclosure while dev server is reachable on LAN.

Mitigations: (1) Upgrade Vite to a patched version; (2) bind dev server to 127.0.0.1 or disable --host; (3) build and serve dist/ with a production server (Nginx/Express static) and strict headers.

F2 — Express App in Development Mode

Asset: <http://192.168.182.108:3000>

Severity: Medium (misconfiguration / information disclosure)

Evidence: Nuclei node-express-dev-env matched; landing page returned stack traces; header X-Powered-By: Express; CORS allowed a dev origin.

Impact: Framework details and error data leaked if endpoint becomes reachable beyond the lab.

Mitigations: (1) NODE_ENV=production for any externally reachable instance; (2) disable detailed errors and remove/mask X-Powered-By (Helmet); (3) restrict CORS to expected origins only.

Recon & Enumeration

Nmap

Host up; 1 IP scanned in ~20s.

Services: 5173/tcp → Vite dev server; 3000/tcp → Node/Express

If banners are needed inline, re-run: `nmap -sV -p 3000,5173 192.168.182.108 -oN artifacts/nmap.txt`

ffuf

Wordlist: quick/common; threads/rate tuned for speed; Result: no interesting paths in this pass.

For deeper coverage: directory-list-2.3-small.txt at a higher rate and manual review.

Passive DAST (ZAP Baseline)

Command (Mac/Docker):

```
docker run --rm \
-v "$HOME/lab-runs/2025-10-07_2054/artifacts:/zap/wrk" \
ghcr.io/zaproxy/zaproxy:stable \
zap-baseline.py -t http://192.168.182.108:5173 -m 3 \
-r zap-baseline.html -x zap-baseline.xml -J zap-baseline.json
```

Automation Framework summary (from console):

FAIL-NEW: 0; WARN-NEW: 7; PASS: 60

Warn items (dev header hygiene):

[10020] Missing Anti-clickjacking Header — 3 URLs (/ , robots.txt, sitemap.xml)
[10021] X-Content-Type-Options Missing — 6 URLs (/ , @vite/client, robots.txt, sitemap.xml, src/logo.png, src/main.jsx)
[10038] Content-Security-Policy Not Set — 3 URLs (/ , robots.txt, sitemap.xml)
[10049] Storable but Non-Cacheable Content — 6 URLs (same set)
[10063] Permissions-Policy Not Set — 5 URLs (/ , @vite/client, robots.txt, sitemap.xml, src/main.jsx)
[10109] Modern Web Application — signal/info
[90004] Insufficient Site Isolation Against Spectre — 11 URLs — signal/info

Production header hardening (Express + Helmet example):

```
import helmet from 'helmet';
app.use(helmet({
  frameguard: { action: 'deny' },
  contentSecurityPolicy: false,
  referrerPolicy: { policy: 'no-referrer' },
  crossOriginOpenerPolicy: { policy: 'same-origin' },
  crossOriginResourcePolicy: { policy: 'same-origin' }
}));
app.disable('x-powered-by');
app.use((_req, res, next) => { res.set('X-Content-Type-Options', 'nosniff'); next(); });
app.use((_req, res, next) => { res.set('Permissions-Policy', 'geolocation=(), camera=(), microphone=()'); next(); });
// Only when serving HTTPS:
app.use(helmet.hsts({ maxAge: 15552000, includeSubDomains: true, preload: false }));
```

Artifacts (relative to ~/lab-runs/2025-10-07_2054/)

Nmap: artifacts/nmap.txt

ffuf: artifacts/ffuf.json

Nuclei (JSONL): artifacts/nuclei.jsonl (matches: CVE-2025-31125 and node-express-dev-env)

Vite LFI PoC (benign): artifacts/vite-lfs-etc-hosts.txt

ZAP Baseline: artifacts/zap-baseline.html, zap-baseline.xml, zap-baseline.json

Metrics: metrics.json (updated with ZAP alert count)

Metrics Snapshot

```
NMAP_BYTES=$(wc -c < "$HOME/lab-runs/2025-10-07_2054/artifacts/nmap.txt" 2>/dev/null || echo 0)
NUCLEI_LINES=$(wc -l < "$HOME/lab-runs/2025-10-07_2054/artifacts/nuclei.jsonl" 2>/dev/null || echo 0)
FFUF_RESULTS=$(jq '.results|length' "$HOME/lab-runs/2025-10-07_2054/artifacts/ffuf.json" 2>/dev/null || echo 0)
ZAP_ALERTS=$(jq '[.site[]|length]' "$HOME/lab-runs/2025-10-07_2054/artifacts/zap-baseline.json" 2>/dev/null || echo 0)
```

```
jq                                     -n                                     \
--arg                                run                                "2025-10-07_2054"                \
--argjson                            nmap                            "$NMAP_BYTES"                    \
--argjson                            nuclei                            "$NUCLEI_LINES"                  \
--argjson                            ffuf                             "$FFUF_RESULTS"                  \
--argjson                            zap                              "$ZAP_ALERTS"                    \
'{run:$run, counts:{nmap_bytes:$nmap, nuclei_lines:$nuclei, ffuf_results:$ffuf, zap_alerts:$zap}}' \
> "$HOME/lab-runs/2025-10-07_2054/metrics.json"
```

Table 2. Security Scan Metrics

Metric	Value
nmap_bytes	(from artifacts)
nuclei_lines	(from artifacts)
ffuf_results	0 (quick profile)
zap_alerts	7 WARN, 0 FAIL

Limitations of This Run

Dev servers by design lack hardened headers; results are not representative of a production build.
ffuf used quick lists only (fast triage).
ZAP: passive baseline, 3-minute spider only.
No destructive techniques; no credential brute-force.

Recommended Remediation & Re-Test Plan

Vite LFI: upgrade to a patched Vite; avoid LAN exposure of the dev server (bind 127.0.0.1);
validate that the @fs path traversal no longer works.
Express hardening: run with NODE_ENV=production; add Helmet; remove/mask X-Powered-By; lock down CORS.
Prod build: serve dist/ via Nginx or Express static with strict headers (CSP, HSTS over HTTPS, X-CTO, Frameguard).
Re-test: re-run Nuclei + ZAP Baseline against the production build to demonstrate remediation; capture deltas in metrics.json.

Discussion & Limitations

Key Learnings

Our lab evaluation shows that an AI-assisted, policy-governed approach can reliably surface development-time exposures without resorting to intrusive techniques. Templated checks (e.g., Nuclei) immediately flagged a dev-server LFI on the Vite stack and a development-mode disclosure on the Express backend. Although both findings are expected in a development environment, they illustrate three important points:

- Fast emergence -> fast detection. Template-driven probes help catch issues (including recently published CVEs) during everyday builds—well before staging/production.
- Guardrails work. Scope enforcement (allow-list only), timeouts, and rate caps prevented out-of-scope scanning and kept traffic bounded, validating the Policy Layer design.
- Decision-ready telemetry. Capturing artifacts (scan outputs, matched templates, ZAP alerts) and counting them into run-level metrics yields repeatable evidence suitable for management review and “security drift” trend analysis.

Where the Results Generalize—and Where They Do Not

- **Generalize:**
 - The orchestration pattern (LLM plan → controlled agents → normalized telemetry) and safe-by-default policies.
 - Early detection of misconfiguration and info-leak classes in web stacks (headers, dev flags, verbose errors).
- **Do not automatically generalize:**
 - Production posture. Dev servers intentionally lack hardened headers and often expose debug endpoints; results here are not representative of a hardened prod build.
 - Coverage depth. We used quick wordlists (ffuf) and a passive ZAP baseline; deeper crawling, authenticated scopes, and custom test packs would find more (or different) issues.
 - Identity/paths. This run did not exercise directory/graph collection (e.g., BloodHound). Attack-path metrics therefore remain unpopulated for this target set.

Threats to Validity

- **Internal validity.** We fixed ROE and used owner-controlled assets, eliminating environmental noise but limiting realism. False positives were minimized by requiring a **clean PoC** (e.g., benign @fs file read) before confirmation.
- **External validity.** Findings depend on stack/version and exposure model. Different frameworks or hardened reverse proxies may change the outcome.

- **Construct validity.** The run emphasizes discoverability and information disclosure; it does not measure exploitability beyond benign proof.
- **Conclusion validity.** No statistical claims (e.g., TTF deltas) are made in this single run. Multi-run series and a human baseline are required for robust comparisons.

Safety, Ethics, and ROE Compliance

All activity was performed on owner-controlled lab assets under explicit Rules of Engagement: allow-listed targets only; no brute force; no destructive payloads; emergency stop on scope violations or error spikes. Artifacts containing potentially sensitive strings were redacted before inclusion in documentation.

Operational Considerations

- **Pipeline fit.** The approach slots naturally into CI for PR-gated checks (templated probes, header hygiene) and into nightly “rehearsals” for drift detection.
- **Cost control.** Containerized tools with time/traffic budgets kept runs predictable; artifact size and log retention can be tuned per project.
- **Human-in-the-loop.** Sensitive actions (mutation, credentialed checks) should require analyst approval; benign discovery can remain fully automated.

Conclusion & Future Work

Conclusion

We presented a pragmatic framework that combines LLM planning, controlled agent execution, and a policy/safety layer to deliver repeatable, auditable offensive testing at machine scale. In a controlled lab, the system surfaced development-time exposures and produced management-ready evidence without breaching ROE. This supports a shift from episodic tests to continuous, business-aligned validation.

Future Work

1. Production hardening loop. Re-run the same playbooks after moving to a production build (Vite patched; Express with Helmet/CSP/HSTS) to demonstrate remediation deltas.
2. Identity & attack-path analytics. Add directory/graph collection (e.g., AD/Entra) to quantify attack-path length and blast radius to crown jewels.
3. Authenticated scopes. Introduce safe, time-boxed authenticated checks (session handling, role scoping) to broaden coverage.
4. Human baseline comparison. Execute the same scenarios with an experienced red team or scripted playbooks to measure TTF, coverage, and repeatability gaps.
5. Cloud techniques. Extend to a test cloud account (AWS/Azure/GCP) using cloud-specific technique packs; report identity/path risks across hybrid assets.

6. Auto-triage & explainability. Generate analyst-ready explanations and fix-it diffs from artifacts; integrate with ticketing for tracked remediation.
7. GRC integration. Map findings and metrics to policy controls and risk registers for decision-grade reporting.
8. Robustness to AI risks. Expand guardrails against prompt injection and excessive agency; add output verification filters and provenance for every tool call.

References

1. MITRE ATT&CK® Enterprise Matrix. MITRE ATT&CK+1
2. NIST SP 800-115: *Technical Guide to Information Security Testing and Assessment*. NIST Computer Security Resource Center+1
3. NIST AI Risk Management Framework (AI RMF 1.0) + Generative AI Profile. NIST+2NIST Publications+2
4. MITRE, (2024). *Caldera – Automated Adversary Emulation Platform*. Достъпно от: caldera.mitre.org.
5. Red Canary, (2024). *Atomic Red Team – Test Library Mapped to ATT&CK*. GitHub.
6. Datadog / Stratus, (2024). *Stratus Red Team – Cloud ATT&CK Technique Emulation*. GitHub; opensource.datadoghq.com.
7. SpecterOps, (2023). *BloodHound – Documentation & Community Edition Quickstart*. BloodHound.
8. ProjectDiscovery, (2024). *Nuclei – Overview and Template Library*. ProjectDiscovery Documentation.
9. FFUF Project, (2024). *ffuf – Fast Web Fuzzer (Docs & Wiki)*. GitHub.
10. OWASP Foundation, (2024). *OWASP ZAP – Baseline Scan (Docker)*. OWASP ZAP Project.
11. OWASP Foundation, (2024). *OWASP Top 10 for Large Language Model Applications*. OWASP.
12. OWASP Foundation, (2024). *AI Security & Privacy Guide / AI Exchange*. owaspai.org.
13. Security Community Sources, (2025). *Vite Dev Server “@fs” Path Traversal / Arbitrary File Read (CVE-2025-30208; CVE-2025-31125)* – advisories и технически анализи.

SPEECH-TO-SPEECH ARTIFICIAL INTELLIGENCE AS A NEW STAGE IN AUTOMATED COMMUNICATION

Radoslav Dodnikov¹

e-mail: radoslav.dodnikov@unwe.bg¹

Abstract

The fast development of AI technologies has led to the emergence of different speech models. One of these models are the speech-to-speech models (S2S) that represent a new stage in automated communication. This paper analyzes its technological foundations, such as encoder-decoder architectures, multilingual acoustic modeling, real-time speech synthesis and prosody transfer. It outlines key applications in education, customer service, healthcare, fintech, and digital accessibility. Big importance is given to the advantages of these models, like better naturalness, reduced latency, and preservation of emotion and speaker characteristics. The paper also discusses current challenges related to model accuracy, language adaptation, data privacy, and ethical use. The analysis concludes that speech-to-speech AI represents a transformative step towards seamless human-machine interaction and opens new opportunities for personalized and context-aware communication technologies.

Ключови думи: speech-to-speech models, ai communication, encoder-decoder architecture, real-time speech synthesis, multilingual acoustic modeling

JEL: C88, O33, L86

Understanding Speech-to-Speech AI

We have all seen sci-fi movies where characters meet other characters that speak a completely different language and their voice is instantly translated into the other language while keeping their emotion, style and personality. This is what we get with speech-to-speech AI models, an out-of-the-future technology that instantly transforms voice without first converting it to text.

Traditional approaches work with three points in their pipeline. First, we get the voice input translated into text (Speech-to-Text or STT). Second, it is processed, for example translated. Finally, it is converted back to speech (Text-to-Speech or TTS). However, each of these processes adds an additional delay and loses important information along the way.

S2S AI, by contrast, works like a direct pipeline. Your voice goes in one end, and transformed voice comes out the other, preserving all the subtle nuances that make human speech rich and expressive: the pauses, the emotional inflections, the rhythm, and even your unique vocal characteristics.

Why is this important

Moving from the traditional speech processing pipelines to S2S is really a big change in the way machines process and generate human speech. This technology enables more natural conversations with AI assistants and agents that sound human, real-time translation that also preserves your voice and emotion, makes communication accessible for people with speech impairments, and makes possible the personalization of voice assistants that can understand and respond in appropriate emotional context.

We have all seen how in recent years and even months, voice AI assistants and customer service bots increase their communication quality and emotional correctness. This is due to the S2S technologies that make our communication with computers more human-like. As we interact more and more with AI in our daily lives, the quality of these interactions becomes crucial. S2S AI represents the next step in

¹ Докторант към катедра ИТК, УНСС, email: radoslav.dodnikov@unwe.bg

making these interactions feel less like talking to a machine and more like conversing with another person.

How S2S AI Works

Understanding S2S AI means looking at four main technologies that work together to make direct voice transformation possible. Think of these as four essential parts of a machine that all need to work together perfectly.

Encoder-Decoder Architectures

At the heart of S2S AI is something called the encoder-decoder architecture. This is a type of neural network that works similar to how our brain processes and reproduces speech.

Imagine you are listening to a song and trying to hum it back. Your brain does not memorize every single sound wave. Instead, it captures the essence of the melody, rhythm, and emotion. This is the encoding process. Then, when you hum it back, your brain reconstructs the song from this compressed memory. This is the decoding process.

S2S AI works in a similar way. The encoder listens to the input speech and compresses it into what we call a context vector. This vector captures everything important about the speech: what is being said, how it sounds, and the emotional tone. Then the decoder takes this context vector and generates the output speech step by step, reconstructing the transformed voice while keeping all the characteristics that were captured.

One good example of this is Microsoft's SpeechT5 framework. It uses this encoder-decoder architecture and can handle multiple tasks like speech recognition and voice conversion using the same core system. Modern S2S systems use something called Transformer architectures with attention mechanisms. These allow the system to focus on the most relevant parts of the input when generating each part of the output, similar to how you might pay special attention to certain words when trying to understand someone's emotional state.

Multilingual Acoustic Modeling

Multilingual acoustic modeling is what allows S2S AI to work across different languages. It trains models on data from many languages at the same time, discovering patterns that are common across all languages while also preserving what makes each language unique.

Different languages have very different sounds and rhythms. Mandarin Chinese uses tones where the pitch changes the meaning of words. English relies more on stress patterns. Arabic has sounds that do not exist in European languages. Traditional systems would need separate models for each language, which is inefficient and hard to maintain.

Multilingual modeling solves this by creating shared representations that capture universal speech patterns. The model learns that certain features are common across languages, like the sound of laughter or emotional expressions, while others are specific to each language, like the rolled R in Spanish or the clicks in some African languages.

This approach has three main benefits. First, it enables cross-lingual speech synthesis, where the system can generate natural speech in languages it has not been extensively trained on by using knowledge from other languages. Second, it allows real-time translation where speech is translated from one language to another while preserving the speaker's voice and emotional tone. Third, it supports low-resource languages by transferring knowledge from languages with lots of training data to languages with limited data.

Imagine attending an international conference where everyone hears the speaker in their native language, but with the original speaker's voice and emotional delivery intact. This is what multilingual S2S AI makes possible.

Real-time Speech Synthesis

Real-time speech synthesis means generating natural-sounding speech fast enough for normal conversation. In human conversation, we expect responses within about 200 to 300 milliseconds. Delays longer than this make conversations feel unnatural and frustrating.

Traditional speech processing pipelines often introduce delays of 500 to 1000 milliseconds or more because the speech has to pass through multiple stages. First it converts speech to text, then processes the text, then converts it back to speech. Each step adds delay.

S2S AI achieves low latency by processing everything in one go. Instead of waiting to process the entire input before generating output, the system can begin producing speech as soon as it has enough information. This is like a simultaneous interpreter who starts translating while the speaker is still talking, rather than waiting for complete sentences.

Modern S2S systems also use optimized neural network designs that balance quality with speed. They process speech in small chunks, typically 10 to 50 milliseconds, which allows them to maintain low latency even for long utterances. This makes S2S AI suitable for live applications like simultaneous translation, interactive voice assistants, and real-time dubbing of live broadcasts.

Prosody Transfer

Prosody refers to the rhythm, stress, and intonation patterns of speech. It is essentially the music of language. Prosody conveys emotion, emphasis, and meaning beyond the words themselves. Prosody transfer is the ability to preserve and transfer these characteristics in synthesized speech.

Consider the sentence "I did not say he stole the money." Depending on which word you emphasize, the meaning changes completely. If you emphasize "I," it means someone else said it. If you emphasize "say," it means you implied it. If you emphasize "he," it means someone else stole it. Prosody also conveys emotion. The same words can sound happy, sad, angry, or sarcastic depending on how they are spoken.

Traditional text-based systems lose all this information because text does not capture pitch, rhythm, or emotional tone. S2S systems maintain prosodic information throughout the transformation process. They encode not just what is being said, but how it is being said: the pitch contours, speaking rate, voice quality, and emotional characteristics.

Advanced S2S models can even separate the linguistic content from the prosodic style. This allows the system to translate words to a different language while preserving the original speaker's emotional delivery. The system can also extract and preserve speaker-specific characteristics like voice timbre, speaking habits, and accent, creating a unique fingerprint for each speaker's voice.

This technology has important applications. It enables emotional speech synthesis for AI assistants that can express appropriate emotions. It allows voice cloning for people who have lost their ability to speak. It creates expressive audiobook narrators that sound engaging and natural with appropriate emotional variation.

Where S2S AI is Used

S2S AI is not just a theoretical technology. It has practical applications that are already improving people's lives in many different areas.

Education

In education, S2S AI is transforming how people learn languages and interact with educational content. Imagine a language learning app that does not just tell you if your pronunciation is correct, but actually demonstrates the correct pronunciation in your own voice, showing you exactly how to adjust. This is now possible with S2S AI.

Educational AI assistants can engage in natural conversations with students and adjust their speaking style based on the student's emotional state. They can speak more encouragingly when a student is frustrated, or more enthusiastically when they are making progress. Students with speech impairments can also participate in oral presentations using S2S systems that generate natural-sounding speech while preserving their intended emotional expression.

A student learning Mandarin Chinese can practice tonal pronunciation with an AI tutor that provides immediate, natural feedback in conversational form. The system can demonstrate correct tones using the student's own voice, making it easier to understand and remember the corrections.

Customer Service

Customer service is another area where S2S AI is making a big difference. Customer service bots powered by S2S AI can detect frustration in a caller's voice and respond with appropriate empathy. They can adjust their speaking style to match the customer's emotional state, being calm and reassuring for anxious customers, or efficient and direct for those who prefer quick resolution.

A single AI agent can also seamlessly switch between languages while maintaining consistent personality and service quality. A customer in Spain can speak Spanish, while a customer in Japan speaks Japanese, both interacting with what feels like the same helpful agent. Companies report that customers rate interactions with S2S-powered agents as significantly more satisfying than traditional text-to-speech bots.

Healthcare

In healthcare, S2S AI is improving both efficiency and accessibility. Doctors can dictate notes, order tests, and access patient information through natural conversation while maintaining sterile conditions or keeping their hands free during procedures. The system understands medical terminology and context, reducing documentation time.

Patients with speech disorders can communicate with healthcare providers using S2S systems that generate clear, natural speech while preserving their intended meaning and emotional context. This is particularly valuable for patients with conditions like ALS, stroke, or vocal cord damage. AI systems can also conduct check-in calls with patients, asking about symptoms and medication compliance in a natural, conversational manner.

A patient with Parkinson's disease, which affects speech clarity, can use an S2S system to communicate with their doctor. The system clarifies their speech while preserving their emotional expression, allowing the doctor to understand not just what the patient is saying, but how they are feeling about their condition.

Fintech

In the financial services industry, S2S AI enables secure voice-based authentication that is difficult to spoof. The system can verify identity through natural conversation, analyzing not just voice characteristics but also speaking patterns and prosody. AI financial advisors can provide guidance through natural conversation, adjusting their communication style based on the customer's financial literacy and emotional state.

S2S systems can also detect anomalies in voice patterns that might indicate stress or deception, helping identify potential fraud attempts during phone banking or transaction authorization. Visually impaired customers can manage their finances through natural voice interactions without needing to navigate visual interfaces.

A customer calling their bank to authorize a large transaction can be verified through voice biometrics during natural conversation. The system detects that they sound confident and calm, suggesting the transaction is legitimate, and processes the authorization without asking security questions or requiring PIN codes.

Digital Accessibility

Perhaps one of the most important applications of S2S AI is in digital accessibility. People with speech impairments can use S2S systems to generate natural-sounding speech that reflects their personality and emotional state. Unlike traditional text-to-speech systems that sound robotic, S2S can create personalized voices that feel authentic.

S2S AI can provide live audio descriptions of visual content for blind users, or generate captions with prosodic information indicating tone and emotion for deaf users. People with cognitive disabilities can interact with technology through simplified, natural voice interfaces that adapt to their communication style and pace.

A person with ALS who has lost the ability to speak can use an S2S system trained on recordings of their voice from before the disease progressed. They can communicate with family members using their own voice, maintaining their identity and emotional connection. This is transformative for people who are losing or have lost their ability to speak.

Global Communication

S2S AI is also breaking down language barriers in global communication. International conferences and meetings can use S2S AI for real-time interpretation that preserves speakers' voices and emotional delivery across languages. Business professionals can negotiate and collaborate across language barriers while maintaining the nuances of tone and emotion that are crucial for building trust.

Media content like podcasts, videos, and audiobooks can be translated to multiple languages while preserving the original speaker's voice and emotional delivery, creating more engaging localized content. Imagine a CEO giving a keynote speech at a global company meeting. Employees around the world hear the speech in their native language, but with the CEO's actual voice and emotional delivery intact. The passion and conviction in the CEO's voice comes through regardless of the listener's language.

Why S2S AI is Better

S2S AI offers several important advantages over traditional speech processing approaches. Understanding these benefits helps explain why this technology represents a real advancement.

Better Naturalness

When speech is converted to text and back to speech in traditional systems, crucial information is lost. Text does not capture the subtle variations in pitch, the natural pauses, the breathing patterns, or the small expressions in voice that make human speech rich and expressive. The result is speech that sounds robotic and unnatural.

S2S AI preserves all these acoustic nuances by processing speech directly without text conversion. It maintains breathing patterns, micro-pauses that indicate thinking or emphasis, pitch variations that convey emotion, voice quality changes that express feeling, and the way sounds blend together naturally in connected speech.

Studies show that listeners consistently rate S2S-generated speech as more natural and human-like compared to traditional TTS output. In some blind tests, S2S speech is indistinguishable from human speech, while traditional TTS is almost always identifiable as synthetic.

Compare these two scenarios. With traditional TTS, "I am so happy to help you today" is spoken in a monotone, robotic voice with even spacing between words. With S2S AI, the same sentence has natural

enthusiasm, slight emphasis on "happy," and a warm, genuine tone. The words are identical, but the S2S version conveys genuine emotion and creates a more positive interaction.

Lower Latency

Traditional pipelines introduce significant delays because speech must pass through multiple stages. Speech to text takes 200 to 500 milliseconds, text processing takes 50 to 200 milliseconds, and text to speech takes another 200 to 500 milliseconds. The total delay is often 450 to 1200 milliseconds or more. These delays disrupt natural conversation flow and create frustrating user experiences. S2S AI achieves much lower latency by eliminating intermediate stages. Direct transformation reduces total latency to 100 to 300 milliseconds, which is within the range of natural human conversation. This is achieved through single-pass processing where one model handles the entire transformation, streaming architecture where output generation begins before input is complete, and optimized neural networks that balance quality and speed.

Low latency enables applications that were previously impossible, like simultaneous interpretation for real-time translation during live conversations, interactive voice assistants with natural back-and-forth dialogue, and live dubbing for real-time translation of broadcasts. S2S AI brings machine conversation into the range of natural human interaction.

Keeping Emotion and Voice Identity

Traditional text-based systems face a fundamental problem. Text cannot encode emotional tone, speaker identity, or prosodic features. When speech is converted to text, all of this information disappears. When text is converted back to speech, the system must guess at appropriate emotion and prosody, often getting it wrong.

S2S systems maintain speaker characteristics throughout the transformation. Voice timbre, accent, speaking habits, and other unique characteristics that make each person's voice recognizable remain intact. The emotional state conveyed in the input speech is preserved in the output. Rhythm, intonation patterns, and emphasis are maintained. Individual speaking habits, like speaking quickly when excited or slowly when explaining something complex, are preserved.

This has important real-world applications. A person who is losing their voice due to illness can have their voice preserved and used to generate speech, maintaining their identity and emotional expressiveness. Actors' voices can be preserved when dubbing films into different languages, maintaining their emotional performance. Virtual assistants can maintain consistent personality and emotional appropriateness across interactions.

Imagine a motivational speaker giving a passionate speech. With S2S translation, the speaker's energy and passion come through in every language, their unique voice remains recognizable, emphasis on key points is preserved, and the emotional arc of the speech remains intact. With traditional translation, all of this would be lost, replaced by a generic voice reading translated text.

Challenges We Still Face

While S2S AI offers tremendous potential, it also faces significant challenges. Understanding these issues is important for the responsible development and use of the technology.

Technical Challenges

One major challenge is model accuracy. Human speech is incredibly diverse. People speak with different accents, dialects, speaking rates, voice qualities, and in various environments like quiet rooms, noisy streets, or over phone lines. Training S2S models that work well across all this diversity is extremely difficult.

A model trained primarily on American English might struggle with Scottish, Indian, or Australian English accents. This creates unfair performance where the technology works better for some groups than others. People also speak differently in different contexts, like formal presentations versus casual

conversation. Models must handle all these variations. Background noise, poor microphone quality, and room acoustics all affect speech quality, and S2S systems must be robust to these real-world conditions.

Another challenge is language adaptation. While S2S AI can theoretically work across many languages, most training data exists for a small number of high-resource languages like English, Mandarin, and Spanish. The world's 7000 plus languages are vastly underrepresented. Many languages have little to no digital speech data available for training. Collecting such data is expensive and requires cooperation from native speaker communities.

Languages also differ dramatically in their sound systems, prosodic patterns, and grammatical structures. A model trained on English might struggle with tonal languages like Mandarin, click languages like Xhosa, or languages with complex consonant clusters like Georgian. Some languages lack standardized writing systems, making it difficult to create training datasets.

Data privacy is another serious concern. Training S2S models requires vast amounts of voice data. Voice is highly personal. It reveals not just what you say, but who you are, your emotional state, your health status, and potentially sensitive information about your identity and background. Voice recordings can be used to identify individuals and track their activities. Voice can reveal health conditions, emotional states, age, gender, ethnicity, and socioeconomic background.

Large databases of voice recordings are attractive targets for hackers and could be used for identity theft, fraud, or blackmail. Current protections like anonymization, encryption, and consent frameworks have limitations. As S2S AI becomes more common, we need clearer regulations specifically addressing voice data, technical innovations for privacy-preserving AI training, greater transparency about how voice data is used, and user-friendly controls for managing voice data.

Ethical Concerns

S2S AI can be used to create highly convincing fake audio, or deepfakes, that sound like someone saying things they never said. This technology can be weaponized for fraud, political manipulation, harassment, and misinformation.

Criminals can impersonate executives, family members, or authority figures to trick people into transferring money or revealing sensitive information. There have been documented cases of fraudsters using voice cloning to impersonate CEOs and authorize fraudulent wire transfers. Fake audio of politicians or public figures can be used to spread misinformation or manipulate elections. Individuals can be targeted with fake audio recordings that damage their reputation or relationships.

The existence of convincing fake audio also undermines the reliability of audio evidence in legal proceedings. "I did not say that, it is a deepfake" becomes a plausible defense even for genuine recordings. Current countermeasures like detection technologies, watermarking, authentication systems, and legal frameworks have limitations. We need industry-wide standards for responsible S2S development, built-in safeguards that prevent misuse, public education about synthetic audio, and international cooperation on legal frameworks.

Another ethical concern is bias in training data. S2S AI systems learn from training data, and if that data reflects societal biases, the systems will perpetuate and potentially amplify those biases. Models may perform better for majority groups than for underrepresented groups, creating unfair access to the technology. Certain accents or dialects may be treated as standard while others are marked as non-standard, reinforcing linguistic discrimination.

Prosodic norms and emotional expression vary across cultures. A model trained primarily on Western data might misinterpret or inappropriately modify emotional expression from other cultures. Training data often overrepresents educated, affluent speakers, potentially making systems less effective for working-class or economically disadvantaged users. People whose voices do not match the training data distribution may find S2S systems do not work well for them, effectively excluding them from services and opportunities.

Questions about consent and ownership are also important. S2S technology raises complex questions about voice ownership and control. Do you own your voice? Can you control how it is used? What rights do you have if someone creates a synthetic version of your voice? Should companies be allowed to use voice recordings to train S2S models without explicit consent? If an S2S system can generate speech in your voice, should it require your permission each time?

Voice rights are poorly defined in most jurisdictions. Voice actors whose livelihoods depend on their unique voices are concerned about being replaced by S2S technology. Some people might want to donate their voices for beneficial uses, like giving voice to people with speech impairments, but there are no clear frameworks for doing this ethically and legally. We need clear legal and ethical frameworks that define voice rights and ownership, establish consent requirements, provide compensation mechanisms, and balance innovation with protection of individual rights.

Finally, there are concerns about transparency and accountability. S2S AI systems are complex black boxes. Even their developers often cannot fully explain why they produce specific outputs. This lack of transparency creates accountability challenges. When an S2S system makes an error or produces problematic output, it is often difficult to understand why, making it hard to fix problems or assign responsibility.

People interacting with S2S systems may not realize they are talking to AI, especially as the technology becomes more convincing. This raises concerns about deception and informed consent. When S2S systems cause harm, it is unclear who is responsible: the developers, the deployers, the users, or the systems themselves. The complexity of S2S systems makes them difficult to audit for bias, errors, or malicious behavior.

We need clear standards for transparency in S2S development and deployment, mandatory disclosure when people are interacting with S2S systems, liability frameworks that assign responsibility for S2S-related harms, independent oversight and auditing mechanisms, and public registries of deployed S2S systems and their capabilities.

Looking to the Future

S2S AI represents a fundamental shift in how machines process and generate human speech. By eliminating the intermediate text conversion step, S2S systems achieve better naturalness, lower latency, and better preservation of emotional and speaker characteristics compared to traditional approaches. This is not just an incremental improvement. It is a qualitative change that makes human-machine interaction feel genuinely natural for the first time.

The technology's impact extends across many areas. Education becomes more personalized and accessible with AI tutors that can engage in natural, emotionally-aware conversations. Healthcare benefits from more efficient documentation, better patient communication, and accessible interfaces for people with disabilities. Customer service becomes more empathetic and satisfying with AI agents that can understand and respond to emotional context. Global communication becomes more seamless with real-time translation that preserves voice and emotion across language barriers. Accessibility improves dramatically, giving voice to people with speech impairments and enabling more inclusive technology.

However, realizing the full potential of S2S AI while avoiding its risks requires collaboration among many groups. Researchers must continue advancing the technology while prioritizing fairness, transparency, and safety. Developers must implement responsible AI practices, including bias testing, privacy protection, and misuse prevention. Policymakers must create legal frameworks that protect individual rights while enabling beneficial innovation. Civil society must advocate for fair access and hold developers accountable. Users must be educated about S2S capabilities and limitations and empowered to control their voice data.

Future Research Directions

The field of S2S AI is rapidly evolving with several promising research directions. Enhanced multilingual capabilities are needed to develop models that work seamlessly across all the world's languages, including low-resource and endangered languages. This requires new approaches to transfer learning and few-shot adaptation.

Improved prosody control will create systems with finer-grained control over emotional expression, speaking style, and prosodic features, enabling more nuanced and context-appropriate speech generation. Personalization and adaptation will build S2S systems that can adapt to individual users' preferences, communication styles, and needs, learning from user feedback and adjusting behavior over time.

Robustness and reliability improvements will enhance performance in challenging conditions like noisy environments, poor audio quality, and diverse accents, while reducing errors that could cause miscommunication or harm. Efficiency and sustainability work will develop more computationally efficient models that can run on mobile devices and edge hardware, reducing energy consumption and environmental impact.

Privacy-preserving techniques will create methods for training and deploying S2S systems that protect user privacy, such as federated learning, differential privacy, and secure multi-party computation. Ethical frameworks will establish industry standards, legal frameworks, and technical safeguards to ensure responsible development and deployment of S2S technology.

Final Thoughts

S2S AI is not just a technological advancement. It is a step toward more natural, empathetic, and inclusive human-machine interaction. As the technology matures, it has the potential to break down communication barriers, enhance accessibility, and create new forms of expression and connection. However, this potential can only be realized if we address the significant technical and ethical challenges the technology presents. We must ensure that S2S AI works fairly for all people, regardless of their language, accent, or background. We must protect against misuse while enabling beneficial applications. We must respect individual rights to privacy and voice ownership while advancing the technology.

The future of S2S AI is not predetermined. It will be shaped by the choices we make today. By approaching this powerful technology with both enthusiasm and caution, technical innovation and ethical reflection, we can create a future where machines understand and speak with us in ways that feel genuinely human, opening new possibilities for communication, creativity, and connection.

References

1. Ao, J., Wang, R., Zhou, L., et al. (2022). "SpeechT5: Unified-Modal Encoder-Decoder Pre-Training for Spoken Language Processing." arXiv:2110.07205. Available at: <https://arxiv.org/abs/2110.07205>
2. Vaswani, A., Shazeer, N., Parmar, N., et al. (2017). "Attention Is All You Need." Advances in Neural Information Processing Systems, 30.
3. Microsoft SpeechT5 GitHub Repository. Available at: <https://github.com/microsoft/SpeechT5>
4. Partnership on AI. Guidelines for responsible development of synthetic media.
5. European Union. "EU AI Act: Regulatory framework for high-risk AI systems including speech technologies."

КИБЕРСИГУРНОСТ КАТО ИНСТРУМЕНТ ЗА ЗАЩИТА НА ЦИФРОВИТЕ АКТИВИ

Cybersecurity as a tool for protecting digital assets

Андрей Вишневский¹

e-mail: avishnevskiy_2520829@unwe.bg¹

Резюме

В настоящия доклад ще бъде разгледана киберсигурността като инструмент за защита на цифровите активи. Каква е ролята за потребителя, която оказват съвременните технологии, какви цели се постигат, каква потенциална заплаха представляват за потребителя и как киберсигурността помага за съхраняването на ценни данни и ресурси. В допълнение към това ще се представят възможностите и предизвикателствата, срещани при киберсигурността в различни направления.

Ключови думи: киберсигурност, защита, изкуствен интелект, цифрови активи

JEL: L86,O33, D83,E42,G18

Abstract

This report will examine cybersecurity as a tool for protecting digital assets. What role does modern technology play for the user, what goals are achieved, what potential threats do they present to the user, and how cybersecurity helps to safeguard valuable data and resources. In addition, the opportunities and challenges encountered in cybersecurity in various areas will be presented.

Въведение

С напредването на технологиите, процесът на дигитализацията става все по-усилен и по-често потребителите използват своите “умни” устройства, такива като телефони, компютри, часовници и други популярни устройства в своето всекидневие. Тези устройства се използват както за развлечения, така и с учебна или работна цел, в допълнение към това, те са много удобни за използване, имат голямо количество опции и възможности и могат значително да улеснят всекидневните задачи, с тях потребителят може да осъществява комуникация с други хора, да се разпорежда със своите финанси, да съхранява важна информация и т.н. Във всяко от устройствата по един или друг начин се съхраняват и създават различни данни, които могат да представляват ценност не само за потребителя или компанията която разполага с тях, но и за различни престъпници които могат да ги използват в личните си цели, обикновено свързани с финансов интерес, такива престъпници извършват киберпрестъпления, а хората извършващи тях се наричат киберпрестъпници и често още ги наричат “хакери”. Практика създадена с цел на защита от такива престъпления се нарича “Киберсигурност” и тя ще стане предмет на обсъждане в този доклад.

Същност на киберсигурността

Киберсигурност представлява практика за защита на компютри, компютърни мрежи, приложения и различни видове софтуер от различни видове заплахи. Киберсигурността е набор от процеси, практики и технологични решения които имат цел да предпазят, данни и важни за работа и не само, системи и устройства, от злонамерен софтуер и други видове дигитални или цифрови заплахи които са способни да донесат различни финансови, репутационни и други вреди. Киберсигурност включва, не само добри практики за защита на ценна информация, но и важен софтуер и други

¹ Студент-магистър, катедра “Информационни технологии и комуникации”, УНСС

технологични решения, които също оказват голямо влияние относно защита на потребителя от опасни действия с лоши последици.

Ролята на киберсигурността и задачи за които тя помага

Киберсигурността служи за защита на потребителите на смарт устройства от различни видове опасности, създавани от недобросъвестни специалисти с различни цели. Такива цели могат да бъдат кражба или неконтролиран достъп до ценна информация, финансови активи, репутационни щети, които могат да бъдат нанесени върху дадена компания или частно лице, в допълнение към това различни киберпрестъпници могат временно или в някои случаи за постоянно да прекратят работа на дадена система, която те атакуват.

Задачи свързани с киберсигурност

- **Защита на конфиденциална информация:** (Лична информация, фирмени тайни, финансова информация)
Такава информация могат да са данни на потребителите в една система, например в онлайн магазин може да има ценна информация свързана със служителите и с клиенти, тази информация могат да бъдат банкови данни, адреси, имена, телефонни номера и сведения за лични документи. Фирмени тайни могат да са свързани, не само с информация за клиенти и служители, но и свързани с фирмени стратегии, технологични патенти, производствени тайни и друга важна информация. Това което се касае финансова информация, обикновено е свързана с ценни активи или сведения за тях, които също трябва да са строго пазени.
- **Предотвратяване на финансови загуби:** Кибератаки могат да доведат до значителни финансови загуби, особено ако това се случва в различни финансови организации, застрахователни дружества, банки и други организации работещи с големи парични суми.
- **Съхранение на репутация:** Загуба на данни или успешна кибератака може силно да навреди репутацията на компания. С това да понесе както и репутационни, така и финансови загуби по различни причини, тези причини могат да бъдат прекратяване на отношения с партньори, загуба на клиенти, съдебни разходи, ако се стига до такива и други видове загуби.
- **Спазване на законови норми:** Повечето държави имат определени норми свързани със съхраняването и обработването на лична информация, заради което за фирмата е важно да има начини как това да се реализира и да се предпази от злонамерени действия причинявани от киберпрестъпници.

Какво предпазва киберсигурността?

Практиките свързани с киберсигурност целят да предпазват от:

- **Fishing (Фишинг) атаки:** Това са подвеждащи фалшиви съобщения обикновено по имейл с цел на получаване на конфиденциална информация. Съществуват още Smishing (Смишинг) и Vishing (Вишинг) това са подобни видове атаки само че, с помощта на съобщение и глас съответно, тоест по телефон и други гласови канали за комуникация.
- **Злонамерен софтуер:** Това е всякакъв софтуер нанасящ технически или финансови щети с някаква зла умишъл, такъв може да бъдат вируси, различен софтуер за кражба, изтриване или достъп до лична информация.

- **DDoS-атаки:** Атаки от много различни реални или виртуални устройства към един сървър в едновременно с цел частично или пълно изключване на даден сървър за неопределен срок от време.
- **Атаки с използване на уязвимости:** Такива могат да бъдат различни ситуации, в които атакуваното устройство е с остарял софтуер, когато то не е обезопасено според необходимите потребности или е открит някакъв труднозабележим проблем, в такъв момент се извършва атака с цел максимална резултативност.
- **Атаки на социална инженерия:** Това са атаки където се използват, не само технически средства, но и психологически техники и умения от престъпника с цел на получаване на ценна информация или финансови активи, като пари, акции, криптовалута и други.
- **Ransomware:** Това са програми криптиращи информация на устройствата на потребител или компания с цел изнудване за пари, унищожаване на информация или загуба на достъпа до нея.
- **Какви са практики за защита на данни и ценни финансови активи?**

Следните практики могат в значителна степен да предпазят потребителите и компаниите от злонамерени действия и сериозни проблеми.

- **Нулево доверие (Zero trust)**
 - Нулево доверие- принцип на киберсигурността, в който се подразбира липса на доверие в такива случаи всеки отдел или служител има достъп само до тази част, до тези ресурси или данни, до които им е необходимо за изпълняваните им задачи, в случая те нямат достъп до други данни или ресурси за които отговарят други ресурси. Допълнително към това се използват средства за достъп и аутентификация с цел, за да се удостовери, че точно това е служител, който трябва да извърши някаква задача и че точно той има нужда от достъп до тези ценни ресурси, а не някой друг. В случай, ако служител си променя своята длъжност или напуска компанията, то и достъпите му се променят или прекратяват.
- **Поведенческа аналитичност**
 - Поведенческа аналитика позволява да се следи предаването на данни и да следи за подозрителни достъпи, например ако извън работно време на компанията има опити за достъп до някакви фирмени системи или някакви лични данни, които не трябва да се разкриват, то това е подозрително поведение, за което може да бъде уведомено по различни начини. Също така при опит за достъп от място различно от обикновенните, например извън сградата на компанията или за достъп до устройство от друг град или държава, това също е подозрително поведение, което ще бъде забелязано и ще бъдат взети мерки.
- **Система за откриване на проникване**
 - Компании и организации използват системи за откриване на проникване, за да идентифицират кибератаки и да реагират бързо на тях. Съвременните решения за сигурност използват изкуствен интелект, машинно обучение и анализ на данни, за да идентифицират скрити заплахи в компютърните системи на компаниите и организациите.

Използват се механизми за предотвратяване на проникване, които събират данни за инциденти, за достъп където се вижда време и място, от което е направен такъв опит, след което, екипите по сигурността могат да идентифицират източниците им.

➤ **Облачно шифроване (Криптиране).**

- Това е процес по криптиране на данни преди попадането им в облака, където те ще се съхраняват, в случая на непредвиден достъп от киберпрестъпници данните ще бъдат криптирани, в неизползваем вид, за тези които не трябва да имат достъп до тях и нямат специален ключ или алгоритъм с който да ги разшифроват. Това също е важно и полезно решение, в допълнение на това, фирмите обикновено разполагат с копие на данните си и на други носители.

Принципите на работа на киберсигурността са следните:

- Удостоверяването и оторизацията проверяват самоличността на потребителите и контролират достъпа до ресурси.
- Криптирането на данните преобразува информацията в криптиран формат, което я прави недостъпна за престъпниците.
- Наблюдение и откриването на заплахи осигуряват непрекъснат контрол на мрежовата активност, за да се идентифицира подозрителна дейност.
- Обучението на служителите повишава осведомеността на потребителите за киберзаплахите и сигурното поведение.
- Копирането на данни позволява възстановяване в случай на атака и загубата им от основният носител.
- Управлението на събитията включва набор от процедури и инструменти за реагиране на кибератаки и намаляването на тяхното въздействие, разполагайки с данните от миналото е възможно да се предотвратят някои от сценариите.

Технологии помагачи за киберсигурност.

За защита се използват различни средства, те постоянно се подобряват спрямо нови потенциални проблеми, за такива средства могат да се представят следните съвременни технологии:

- Изкуствен интелект (AI), той се използва за анализ на поведението, препоръки към потребителите и намирането на аномалии.
- Машинно обучение помага да предсказва вероятните сценарии и да предотвратява заплахи.
- Блокчейн осигурява съхранение на данни, повишавайки тяхната сигурност.
- Многофакторното удостоверяване (MFA) използва множество методи за проверка на самоличността за достъп до системи.
- Платформите за управление на информацията за сигурност и събития (SIEM) събират, анализират и съпоставят данни за събития за сигурност.
- Криптографията използва методи за криптиране на данните, за да осигури максимално тяхната защитеност.

Заклучение

Съвременният свят непрекъснато се развива и технологиите все повече са свързани със съвременният човек, всеки ден обработваме големи обеми от информация и работим с големи обеми от данни, също така все повече задачи се изпълняват от устройствата ни, различни покупки-продажби, поръчки, подписване на договори, работа с документи и други важни действия. Всяко едно от тези важни за нас действия може да бъде по някакъв начин атакувано от киберпрестъпници или така наречените “хакери”. Техните атаки могат да бъдат с различни цели и върху различни дейности извършвани от обикновен потребител или някаква компания

или организация, което поставя в опасност данни и финансови активи, които могат да бъдат засегнати. В резултат на такива кибератаки могат да бъдат откраднати, загубени и получени ценни данни, а също така и финансови активи от различни видове. За да се избегнат такива атаки и да се минимизират последствията от потенциални такива, съществува Киберсигурност, това е набор от практики и технологии създадени с цел на обезопасяване в технологичните пространства и спазването на всички препоръки и използването на съществуващите средства за защита, значително помага за сигурността на един потребител или цяла компания.

Литература:

1. Microsoft, (н.д.). Какво представлява киберсигурността? Достъпно от: <https://www.microsoft.com/bg-bg/security/business/security-101/what-is-cybersecurity>
2. Microsoft, (н.д.). Какво е фишинг? Достъпно от:
3. <https://www.microsoft.com/bg-bg/security/business/security-101/what-is-phishing>
4. Amazon Web Services, (2025). What is Cybersecurity? Cybersecurity Explained. Достъпно от: <https://aws.amazon.com/what-is/cybersecurity>
5. Kaspersky, (н.д.). Cyber Security Solutions for Home & Business. Достъпно от: <https://www.kaspersky.ca>
6. Economic.bg, (2016). Yahoo призна за най-голямата кибератака в историята. Достъпно от: <https://www.economic.bg/bg/a/view/yahoo-prizna-za-naj-golyamata-kiberataka-v-istoriyata-64311>
7. РБК Тренды, (2021). 10 самых громких кибератак XXI века. Достъпно от:
8. <https://trends.rbc.ru/trends/industry/600702d49a79473ad25c5b3e>
9. bTV Новините, (н.д.). 5 кибератаки, по-страшни от тази срещу НАП. Достъпно от: <https://btvnovinite.bg/svetut/5-kiberataki-po-loshi-ot-tazi-sreshtu-nap.html>
10. BBC News, (2016). Petya ransomware encryption system cracked. Достъпно от: <https://www.bbc.com/news/technology-36014810>
11. BBC News, (2016). The ransomware that knows where you live. Достъпно от: <https://www.bbc.com/news/technology-35996408>
12. Goodreads, (н.д.). Quotes Library. Достъпно от: <https://www.goodreads.com/quotes>
13. CrowdStrike, (н.д.). What Is a Ransomware Attack? Достъпно от: <https://www.crowdstrike.com/en-us/cybersecurity-101/ransomware>
14. Khan Academy, (н.д.). Phishing Attacks – Cyber Crime Section. Достъпно от: <https://www.khanacademy.org/computing/.../phishing-attacks>
15. PayPal, (н.д.). PayPal Bulgaria – Home. Достъпно от: <https://www.paypal.com/bg/home>
16. DTU.kz, (н.д.). Что такое кибербезопасность и ее основные характеристики. Достъпно от: <https://www.dtu.kz/blog-post/chto-takoe-kiberbezopasnost-i-ee-osnovnye-harakteristiki>

EXPLORING SEASONAL TRENDS IN BULGARIAN MEDIA: INSIGHTS FROM A DATA-DRIVEN ANALYSIS OF WEB PUBLICATIONS

Plamen Milev¹

e-mail: pmilev@unwe.bg¹

Abstract

This paper presents a data-driven analysis of Bulgarian online media, aimed at identifying seasonal trends and event-driven fluctuations in news coverage. The study quantitatively analyzes the total monthly publication volume and the frequency of ten most popular keywords over a one-year period. The findings reveal two primary influences on the editorial agenda: predictable seasonal patterns, with a notable decrease in media output during the summer months, and the strong impact of specific high-profile events. For instance, coverage related to political terms like "Trump" and "Elections" peaked in direct alignment with major electoral cycles, while geopolitical topics such as "Ukraine" and "Israel" showed sharp fluctuations corresponding to key international developments. The paper discusses these dynamics, providing empirical insights into how the Bulgarian media agenda is shaped by the interplay of cyclical trends and event-driven news.

Key words: Web Scraping, Media Intelligence, Data Analysis, Web Publications, Editorial Focus

JEL: C88, L86.

Introduction

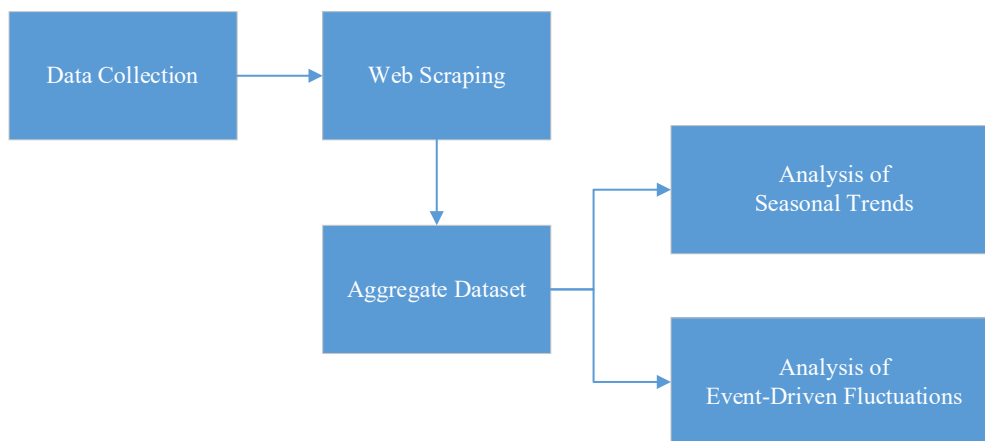
In the modern digital age, online media plays a crucial role in shaping public opinion and defining the public agenda. The constant 24-hour flow of information creates a dynamic and highly competitive environment where editorial priorities shift rapidly. Understanding the factors that influence these priorities is essential for analysts, researchers, and media professionals alike. While intuitive assumptions are often made about "calmer" or "busier" news periods, there is a lack of specific, data-driven research to outline these trends in the Bulgarian media context. This paper aims to fill this gap by conducting a quantitative analysis of web publications from leading Bulgarian media outlets. The primary objective of the research is to identify and analyze both recurring seasonal trends and specific event-driven fluctuations in media coverage. By measuring the total volume of publications and tracking the frequency of key topics, we aim to provide an empirical basis for understanding when and why certain topics dominate the news flow. For the purpose of this analysis, a full one-year period was examined – from April 2024 to March 2025. The methodology involves web scraping and processing of publication data from eight prominent media platforms: Bulgaria ON AIR, Bloomberg, BNR, BNT, bTV, DW, Eurocom, and Nova. The analysis includes two main components: first, an aggregated analysis of the total number of monthly publications to establish baseline levels of media activity and their seasonal fluctuations; and second, a detailed frequency analysis of the ten most frequent keywords identified through this scraping process. These prominent keywords ("Bulgaria," "Sofia," "Ukraine," "Trump," "USA," "Elections," "Russia," "Israel," "NATO," and "Europe") inherently represent a mix of domestic, geopolitical, and international political topics. As will be shown in the subsequent sections, the analysis reveals clear and predictable seasonal patterns, most notably a significant decline in media activity during the summer months. However, sharp, event-driven peaks stand out against this baseline. For example, coverage of topics related to the "USA" and "Trump" reaches its zenith in line with the US election cycle (November 2024 – January 2025), while topics like "Ukraine" and "Israel" show fluctuations directly related to the escalation of events. This paper proceeds by detailing the data collection methodology, followed by an analysis of the aggregated publication volumes to establish

¹ Associate Professor, PhD. Department of Information Technologies and Communications, University of National and World Economy, ORCID 0000-0002-4867-0586.

seasonal trends. Subsequently, it delves into the specific monthly dynamics of the top keywords, before concluding with a summary of the findings on the dual influence of seasonality and event-driven developments on the Bulgarian media agenda.

Methodology

The methodological approach of this paper is grounded in the contemporary challenges of digital data processing. The modern digital landscape is defined by the rapid integration of advanced technologies such as Artificial Intelligence [1] and novel user interfaces [2]. This technological shift necessitates new approaches for managing large-scale, cloud-based Big Data solutions [3] and, critically, for processing the vast amounts of unstructured data generated by web and social platforms [4]. This paper builds directly on this context. It specifically applies established principles of modern web data processing to handle the collection of online information [5]. Furthermore, it employs a systematic approach to organizing this information, aligning with core concepts of data cataloging [6]. To extract insights from the collected data, the research utilizes a quantitative approach that draws from several specific analytical methods. These include techniques for automated content analysis [7], intelligent tagging and search [8], natural language processing [9], and semantic analysis [10]. This combination of methods allows for the transformation of raw, unstructured web data into the meaningful, quantitative insights presented in this paper. Figure 1 illustrates the methodological approach of the research.



Source: Author's diagram illustrating the research process

Figure 1: The Methodological Flowchart

This paper employs a quantitative methodology based on web scraping and data analysis to examine publication trends in Bulgarian online media. The process is divided into two main stages: data collection and data analysis.

Data Collection

The dataset for this study was compiled by scraping publicly available information from eight prominent Bulgarian media platforms: Bulgaria ON AIR, Bloomberg, BNR, BNT, bTV, DW, Eurocom, and Nova. The data collection covers a complete one-year period, from April 1, 2024, to March 31, 2025, to ensure the inclusion of all seasonal variations.

The web scraping process was configured to extract the titles and publication timestamps of articles from the selected media sources. This focus on publication titles allows for a high-level analysis of the primary topics being presented to the public, as headlines are designed to encapsulate the most newsworthy aspect of a story.

Data Analysis

The collected raw data was processed and aggregated for analysis. The analysis phase consisted of two primary components:

- Aggregated publication analysis: the total number of publications from all eight media sources was aggregated on a monthly basis. This approach was used to establish a baseline for overall media output and to identify broad seasonal trends, such as peaks of activity or quieter periods (e.g., summer months, holidays).
- Keyword frequency analysis: the second component involved a detailed analysis of the publication titles. The entire corpus of collected titles was processed to identify the most frequently occurring and relevant terms. This data-driven approach resulted in the identification of the top 10 keywords: "Bulgaria," "Sofia," "Ukraine," "Trump," "USA," "Elections," "Russia," "Israel," "NATO," and "Europe." The monthly frequency of each of these keywords was then tracked individually to observe the dynamics of specific topics and their correlation with real-world events.

The application of this two-phase methodological approach – spanning from broad data collection to specific keyword frequency analysis – produced a granular dataset. The following section presents and discusses these results, detailing the identified seasonal patterns and event-driven fluctuations in the Bulgarian media's publication volume and thematic focus.

Results and Discussion

This section transitions from the outlined methodology to the empirical findings of the study. It presents a comprehensive analysis of the data collected from the eight selected Bulgarian media platforms. The analysis is structured to first provide a high-level overview of the aggregate monthly distributions to identify seasonal trends, followed by a granular analysis of the ten most frequent keywords, illustrating how editorial focus shifts in response to specific, event-driven national and international developments. This multi-layered approach provides a quantitative basis for understanding the dynamics of the Bulgarian digital news landscape.

Monthly Publication Trends

The initial data collection phase aggregated the total number of web publications from each of the eight selected media outlets for the 12-month period. This provided a baseline understanding of the relative output of each platform. To identify seasonal patterns, the total number of publications from all eight platforms was aggregated on a monthly basis (Table 1).

Table 1: Monthly Distribution of Publications

Month	Publications
April 2024	20526
May 2024	20210
June 2024	20444
July 2024	19936
August 2024	17736
September 2024	18144
October 2024	22037
November 2024	21125
December 2024	20224
January 2025	22740
February 2025	21632
March 2025	23368

Source: Author's data obtained from web scraping

The monthly distribution reveals a relatively consistent content generation throughout the year, with clear seasonal fluctuations. Noticeable peaks occur in March 2025 (23,368 publications) and January 2025 (22,740 publications), likely due to increased coverage related to significant political, economic, or social events common at the beginning of the year. Conversely, August 2024 registers the lowest activity (17,736 publications), reflecting typical seasonal patterns associated with vacation periods and reduced overall media output. September 2024 similarly shows relatively lower numbers, suggesting a gradual return to standard publication volumes following the summer season.

Keyword Fluctuation Analysis

The following section details the monthly frequency of the top 10 most popular keywords identified in the publication titles. This analysis highlights how editorial focus shifts in response to specific national and international events. Table 2 presents the monthly frequency of occurrences of the word "Bulgaria".

Table 2: Monthly Occurrences of "Bulgaria"

Month	Titles
April 2024	743
May 2024	756
June 2024	862
July 2024	649
August 2024	614
September 2024	863
October 2024	871
November 2024	709
December 2024	678
January 2025	691
February 2025	695
March 2025	877

Source: Author's data obtained from web scraping

The monthly distribution of "Bulgaria" reveals moderate fluctuation. Peaks appear in March 2025 (877 titles), October 2024 (871 titles), and September 2024 (863 titles), corresponding with significant domestic events. Lower frequencies are registered in the summer (July and August), reflecting a typical seasonal reduction in media coverage. Table 3 presents the monthly frequency of occurrences of the word "Sofia".

Table 3: Monthly Occurrences of "Sofia"

Month	Titles
April 2024	614
May 2024	802
June 2024	664
July 2024	438
August 2024	391
September 2024	621
October 2024	764
November 2024	687
December 2024	647
January 2025	691
February 2025	824
March 2025	633

Source: Author's data obtained from web scraping

The analysis of "Sofia" highlights clear variations. Significant peaks are observable in February 2025 (824 titles) and May 2024 (802 titles), suggesting increased coverage of events in the capital. Noticeably lower coverage occurred during the summer (July and August), following a similar seasonal pattern to "Bulgaria". Table 4 presents the monthly frequency of occurrences of the word "Ukraine".

Table 4: Monthly Occurrences of "Ukraine"

Month	Titles
April 2024	541
May 2024	577
June 2024	589
July 2024	478
August 2024	351
September 2024	370
October 2024	397
November 2024	647
December 2024	499
January 2025	360
February 2025	1006
March 2025	1003

Source: Author's data obtained from web scraping

The frequency of "Ukraine" displays considerable variability. A significant peak is evident in February 2025 (1,006 titles) and March 2025 (1,003 titles), which could be linked to critical geopolitical or military developments. Conversely, a noticeable drop appears during the late summer, suggesting decreased media emphasis. Table 5 presents the monthly frequency of occurrences of the word "Trump".

Table 5: Monthly Occurrences of "Trump"

Month	Titles
April 2024	90
May 2024	113
June 2024	92
July 2024	552
August 2024	186
September 2024	268
October 2024	243
November 2024	1034
December 2024	341
January 2025	1151
February 2025	1191
March 2025	1113

Source: Author's data obtained from web scraping

The frequency of "Trump" strongly correlates with political events. A remarkable increase occurs from November 2024 (1,034 titles) through February 2025 (1,191 titles). This period coincides with critical political events, most likely connected to the US presidential elections and their aftermath. Table 6 presents the monthly frequency of occurrences of the word "USA".

Table 6: Monthly Occurrences of "USA"

Month	Titles
April 2024	452
May 2024	435
June 2024	358
July 2024	373
August 2024	346
September 2024	307
October 2024	412
November 2024	645
December 2024	409
January 2025	717
February 2025	855
March 2025	895

Source: Author's data obtained from web scraping

The distribution of "USA" exhibits noticeable variability, with distinct peaks toward the end of 2024 and the beginning of 2025. Increased coverage is observed from November 2024 (645 titles) to March 2025 (895 titles), reflecting heightened media attention driven by major political events, including the US presidential elections and post-election dynamics. Table 7 presents the monthly frequency of occurrences of the word "Elections".

Table 7: Monthly Occurrences of "Elections"

Month	Titles
April 2024	566
May 2024	390
June 2024	622
July 2024	303
August 2024	254
September 2024	522
October 2024	669
November 2024	629
December 2024	191
January 2025	165
February 2025	197
March 2025	263

Source: Author's data obtained from web scraping

The distribution of "Elections" shows patterns related to political events. Notable peaks are evident in October 2024 (669 titles) and November 2024 (629 titles), likely coinciding with the US presidential elections. Another increase is observed in June 2024 (622 titles), a period which may reflect regional or European political events. Markedly lower coverage from December 2024 onwards suggests a post-election period. Table 8 presents the monthly frequency of occurrences of the word "Russia".

Table 8: Monthly Occurrences of "Russia"

Month	Titles
April 2024	537
May 2024	293
June 2024	198
July 2024	133
August 2024	203

September 2024	342
October 2024	707
November 2024	271
December 2024	124
January 2025	213
February 2025	128
March 2025	190

Source: Author's data obtained from web scraping

The occurrences of "Russia" exhibit moderate fluctuations. Notable peaks are evident in March 2025 (580 titles) and November 2024 (448 titles), likely correlating with significant geopolitical developments. Lower frequencies appear during the summer (July and August), reflecting reduced editorial emphasis. Table 9 presents the monthly frequency of occurrences of the word "Israel".

Table 9: Monthly Occurrences of "Israel"

Month	Titles
April 2024	537
May 2024	293
June 2024	198
July 2024	133
August 2024	203
September 2024	342
October 2024	707
November 2024	271
December 2024	124
January 2025	213
February 2025	128
March 2025	190

Source: Author's data obtained from web scraping

The monthly occurrences of "Israel" show substantial fluctuations. The most prominent peaks appear in October 2024 (707 titles) and April 2024 (537 titles), likely indicating major political, military, or diplomatic events. Coverage declines sharply during other months, reaching its lowest points in December 2024 and February 2025. Table 10 presents the monthly frequency of occurrences of the word "NATO".

Table 10: Monthly Occurrences of "NATO"

Month	Titles
April 2024	208
May 2024	263
June 2024	300
July 2024	261
August 2024	255
September 2024	184
October 2024	247
November 2024	280
December 2024	266
January 2025	329
February 2025	280
March 2025	280

Source: Author's data obtained from web scraping

The frequency of "NATO" remains relatively stable, with moderate fluctuations. Slight peaks are evident in January 2025 (329 titles) and June 2024 (300 titles), possibly connected to NATO summits or significant policy announcements. The lowest frequency occurs in September 2024 (184 titles). Table 11 presents the monthly frequency of occurrences of the word "Europe".

Table 11: Monthly Occurrences of "Europe"

Month	Titles
April 2024	180
May 2024	333
June 2024	230
July 2024	185
August 2024	144
September 2024	212
October 2024	236
November 2024	279
December 2024	187
January 2025	266
February 2025	407
March 2025	424

Source: Author's data obtained from web scraping

The occurrences of "Europe" demonstrate moderate fluctuations. A significant increase occurs in February 2025 (407 titles) and March 2025 (424 titles), indicating increased media interest in European Union activities or regional developments. Another peak in May 2024 (333 titles) may correspond to EU-related events. The lowest number is seen in August 2024.

In summary, the presented data analysis confirms a clear model of dual influence on the Bulgarian online media agenda. On one hand, the aggregate publication data demonstrates a distinct and predictable seasonal cycle, characterized by a significant decline in media activity during the late summer period (August). On the other hand, the detailed keyword frequency analysis reveals a high degree of reactivity to specific, high-impact events. This is most evident in the sharp peaks in coverage for topics like "Trump" and "USA," correlating directly with the U.S. election cycle, and the surges in mentions of "Ukraine" and "Israel," coinciding with key geopolitical developments. These findings collectively illustrate the dynamic interplay between predictable, seasonal rhythms and the unpredictable, event-driven nature of the modern news cycle.

Conclusion

This paper provides empirical evidence that the Bulgarian digital news agenda operates under a dual influence: it is anchored by a predictable seasonal cycle but is simultaneously and powerfully steered by the unpredictable cadence of major national and global events. Future research could expand upon this methodology by moving beyond title-based keyword analysis to full-text sentiment analysis, or by comparing these established media trends against the corresponding discourse on social media platforms to identify potential disparities or time lags in the public conversation.

Acknowledgement

This work was financially supported by the UNWE Research Programme (Research Grant No. 22/2024/A).

References

1. Dimitrov, K. (2025). The Growing Role of VR And AI Industry to Transform Soft Skills Development. *Trakia Journal of Sciences*, 23(2), 298-303.

2. Mihova, V. (2023). Expanding The Possibilities of Information System for Working with Voice. *Innovative Information Technologies for Economy Digitalization (IITED)*, (1), 152-157.
3. Kovacheva, M. (2025). Evaluating Challenges and Opportunities Associated with Cloud-Based Big Data Solutions. *Nauchni trudove*, (2), 175-188.
4. Velkova, I. (2023). Unstructured social media data processing with artificial intelligence. *Industry 4.0*, 8(2), 65-67.
5. Tabov, Y. (2024). Web Data Processing in the Digital Age: Challenges and Solutions. *Innovative Information Technologies for Economy Digitalization (IITED)*, (1), 223-228.
6. Stefanov, G. (2023). Data Cataloging in Big Data Environment. *Innovative Information Technologies for Economy Digitalization (IITED)*, (1), 102-107.
7. Belev, I. (2024). Combining digital business automation technologies. *Innovative Information Technologies for Economy Digitalization (IITED)*, (1), 305-312.
8. Marzovanova, M. (2015). Intelligent Tagging and Search as a Fully Automated System. *Economic Alternatives*, (2), 95-100.
9. Radoev, M. (2024). Application of Natural Language Processing in Accounting. *Innovative Information Technologies for Economy Digitalization (IITED)*, (1), 107-113.
10. Lazarova, V. (2025). Artificial Intelligence as a Service. *Ikonomiceski i Sotsialni Alternativi*, (2), 128-141.

ИНТЕРНЕТ НА ОБЕКТИТЕ И ДИГИТАЛНИ ДВОЙНИЦИ В ОБЩЕСТВЕНИТЕ ПОРЪЧКИ: МОДЕЛ ЗА УМЕН КОНТРОЛ И ПРЕВЕНЦИЯ НА КОРУПЦИОННИ ПРАКТИКИ

Internet of Things and Digital Twins in Public Procurement: A Model for Smart Control and Corruption Prevention

Теодор Буров¹

email: teodor.burov@unwe.bg¹

Абстракт

Настоящият доклад представя концептуален модел за прилагане на технологии от типа Интернет на обектите (IoT) и Дигитални двойници (Digital Twins) в процеса на възлагане и изпълнение на обществени поръчки. Моделът обединява структурирани данни от електронната платформа ЦАИС ЕОП и сензорни данни от реалното изпълнение на проектите, за да създаде интерактивен дигитален близък на поръчката. Чрез алгоритми за машинно обучение се анализират ключови рискови индикатори – вид процедура, критерии за подбор, прогнозни стойности, концентрация на участници и отклонения в срокове и бюджети. България разполага с отлична основа за развитие на подобен модел благодарение на отворените данни, публикувани от ЦАИС ЕОП, и възможностите за експорт в CSV и XLS формати. Предстоящото публикуване на данни в OCDS (Open Contracting Data Standard) ще улесни изграждането на аналитични платформи, визуализации и граждански мониторинг. Липсата на интеграция на България в Public Procurement Data Space (PPDS) е силен аргумент за оригиналността и необходимостта от национален модел. Данните могат да бъдат използвани от институции като КПК, КЗК, ГДБОП за анализ на картелни практики и корупционни рискове, както и от граждански организации за независим надзор. Този подход превръща контрола върху обществените поръчки от реактивен в интелигентен, предсказуем и прозрачен процес, подпомагащ изграждането на доверие и интегритет в публичните разходи.

Abstract

The present report introduces a conceptual model for applying Internet of Things (IoT) and Digital Twin technologies in the process of awarding and implementing public procurement contracts. The model integrates structured data from the national e-procurement platform CAIS EOP with sensor data from the actual execution of projects, creating an interactive digital twin of each procurement procedure. Machine learning algorithms are used to analyze key risk indicators such as the type of procedure, selection criteria, estimated values, concentration of bidders, and deviations in deadlines and budgets. Bulgaria has a strong foundation for developing such a model thanks to the open data published by CAIS EOP and the possibility to export datasets in CSV and XLS formats. The forthcoming publication of data in the OCDS (Open Contracting Data Standard) format will further facilitate the development of analytical platforms, visualizations, and civic monitoring tools. The absence of Bulgarian integration into the Public Procurement Data Space (PPDS) provides a strong argument for the originality and necessity of a national-level model. The data can be used by institutions such as the Commission for Counteracting Corruption, the Commission on Protection of Competition, and the Cybercrime Directorate to analyze cartel practices and corruption risks, as well as by civil society organizations for independent oversight. This approach transforms public procurement oversight from

¹ Докторант, катедра „Финансов контрол“, Финансово-счетоводен факултет, УНСС, ORCID 0009-0007-8998-2985, email: Teodor.Burov@unwe.bg.

a reactive process into an intelligent, predictive, and transparent system, strengthening trust and integrity in public spending.

Ключови думи: интернет на обектите, дигитални двойници, стандарт за отворени данни при обществените поръчки (OCDS), превенция на корупцията.

JEL: H57, D83, O33, H83

Въведение

Дигитализацията на публичния сектор се превърна в ключов фактор за повишаване на прозрачността, ефективността и отчетността на публичните институции. В областта на обществените поръчки, където управлението на значителни финансови ресурси е свързано с висок риск от нередности и корупционни практики, внедряването на интелигентни технологии за контрол е не само технологична иновация, но и инструмент за укрепване на доверието в публичните процеси. През последното десетилетие Европейският съюз насърчава изграждането на общи стандарти и инфраструктури за обмен и анализ на данни в обществените поръчки. Сред тях особено значение имат Open Contracting Data Standard (OCDS), разработен от *Open Contracting Partnership*, и инициативата на Европейската комисия Public Procurement Data Space (PPDS), целяща интеграция на националните платформи с европейските бази данни като TED (Tenders Electronic Daily). Тези усилия създават предпоставки за повишаване на прозрачността и съпоставимостта на данните между държавите членки.

Въпреки това, България все още не е интегрирана в PPDS, а публикуваните данни за страната се ограничават единствено до процедурите над европейските прагове, отразени в TED. Националната система Централизирана автоматизирана информационна система „Електронни обществени поръчки“ (ЦАИС ЕОП) предоставя значително по-богат и детайлен масив от информация, включващ целия жизнен цикъл на обществените поръчки – от планирането и откриването до изпълнението и измененията на договорите. Това създава възможност за изграждане на национален модел за интелигентен контрол, който да използва отворените данни от ЦАИС ЕОП и предстоящото въвеждане на OCDS формат за периодично публикуване.

Едновременно с това, технологичните тенденции в сферата на Интернет на обектите (Internet of Things – IoT) и Дигиталните двойници (Digital Twins) откриват нови перспективи за управление и наблюдение на реални процеси чрез виртуални модели и сензорни данни в реално време. Докато IoT осигурява събиране и предаване на информация от физическата среда, дигиталните двойници позволяват създаване на динамични симулации, чрез които може да се проследява, прогнозира и оптимизира поведението на даден обект или процес. Приложени в контекста на обществените поръчки, тези технологии биха могли да осигурят интелигентен, предсказуем и прозрачен контрол върху възлагането и изпълнението на договори.

Настоящият доклад представя концептуален модел, който комбинира принципите на IoT и дигиталните двойници с потенциала на отворените данни от ЦАИС ЕОП, за да се създаде умна система за превенция на корупционни рискове и анализ на ефективността на обществените поръчки. Основната цел е да се предложи рамка за прилагане на технологии за реално-времеви мониторинг и аналитичен контрол, която надгражда съществуващите европейски модели и отговаря на специфичните особености на българската среда.

Задачите на изследването включват:

- анализ на концепциите за Интернет на обектите и дигиталните двойници и тяхната приложимост в публичния сектор;
- идентифициране на възможностите за използване на националните отворени данни (ЦАИС ЕОП, OCDS) за изграждане на дигитални модели на обществените поръчки;
- формулиране на архитектура за интелигентен контрол, базирана на рискови индикатори и машинно обучение;
- очертаване на потенциала за институционално и гражданско приложение на подобен модел.

В резултат се цели да се демонстрира, че при наличие на подходящи стандарти за публикуване и интегриране на данни, България може да изгради собствен модел на дигитален контрол и прозрачност, съвместим с европейските инициативи, но адаптиран към националния контекст. Този подход не само ще подпомогне ефективното управление на публичните ресурси, но и ще създаде устойчива основа за изграждане на доверие, интегритет и отговорност в процеса на обществените поръчки.

Теоретични и концептуални основи

Концепцията за Интернет на обектите (IoT) и приложението ѝ в публичния сектор

Интернет на обектите (Internet of Things, IoT) представлява мрежа от физически устройства, оборудвани със сензори, софтуер и комуникационни технологии, които събират и обменят данни с минимална човешка намеса. В основата си IoT позволява наблюдение и управление на физически процеси в реално време, което води до по-висока ефективност, предсказуемост и прозрачност.

В публичния сектор IoT намира приложение в различни области: интелигентна инфраструктура, управление на трафика, мониторинг на енергийна ефективност, градско планиране и дори здравеопазване. Тези технологии дават възможност за изграждане на „умни градове“ и интелигентни системи за управление на публичните ресурси, при които решенията се основават на реални данни, а не на субективни оценки.

В контекста на обществените поръчки IoT може да се използва за проследяване на изпълнението на договори – например чрез сензори, които подават данни за напредъка на строителни дейности, доставки или услуги. Тези данни могат да се комбинират със структурирана информация за самата процедура (оферти, участници, цени), което отваря път към интегрирани модели на контрол и оценка.

Дигиталните двойници (Digital Twins) – от индустриални симулации към управление на обществени процеси

Понятието „дигитален двойник“ се отнася до виртуално представяне на реален обект, система или процес, което позволява неговото наблюдение, анализ и оптимизация. Дигиталните двойници възникват в индустрията – например при проектирането на машини, транспортни системи и производствени линии – но все по-често се прилагат и в управлението на сложни обществени системи.

Основната им функция е да създадат динамична симулация, която отразява реалното състояние на обекта в даден момент, като използва данни от различни източници: сензори, бази данни, административни регистри и аналитични модели. В контекста на обществените поръчки това означава, че всеки проект или договор може да има свой „дигитален близък“, който обединява информация за:

- планиран бюджет, прогнозни стойности и критерии за възлагане;
- участници и подизпълнители;
- хронология на изпълнението, срокове и плащания;
- отклонения от първоначалните параметри.

Този подход позволява възпроизвеждане на реалните процеси в цифрова среда и прилагане на аналитични алгоритми за откриване на аномалии, оценка на риска и прогнозиране на резултати. По този начин дигиталният двойник се превръща в инструмент за интелигентен контрол и превенция на нарушения в реално време.

Синергия между IoT и дигиталните двойници за контрол в обществените поръчки

Комбинирането на IoT и дигиталните двойници представлява следващ етап в еволюцията на дигиталния контрол. Докато IoT осигурява потоци от сензорни данни, дигиталният двойник ги интерпретира в контекста на конкретната обществена поръчка, позволявайки интерактивен анализ и симулация.

Тази синергия има няколко ключови предимства:

- Непрекъснат мониторинг – данните за изпълнение се подават автоматично, без забавяне.
- Идентификация на рискови отклонения – чрез анализ на ключови индикатори (срокове, цени, концентрация на участници, промени в договора).
- Подкрепа за вземане на решения – чрез прогнозни модели и визуални панели.
- Интеграция на институционален и граждански контрол – достъп до визуализирани данни в реално време.

В резултат контролът върху обществените поръчки може да се трансформира от последващ и документален към динамичен и предсказуем, което съществено намалява възможностите за корупционни практики и неефективно разходване на публични средства.

IoT и дигиталните двойници в контекста на прозрачността и интегритета

Връзката между дигиталните технологии и интегритета на публичния сектор се разглежда все по-често в рамките на концепцията за „data-driven integrity“ – интегритет, основан на данни. Изследвания на OECD, World Bank и European Commission посочват, че използването на структурирани и достъпни данни позволява не само откриване на нередности, но и предварително прогнозиране на рискове, което е същността на съвременния контрол.

В този контекст дигиталните двойници и IoT се явяват инструмент на управлението по интегритет, като осигуряват прозрачност, автоматизация и аналитичен капацитет. Прилагането им в обществените поръчки ще позволи:

- обективна оценка на изпълнението на договори;
- наблюдение на съответствието със законовите изисквания;
- откриване на потенциални конфликти на интереси и картелни съглашения;
- предоставяне на достоверна информация за граждански мониторинг.

Така концепцията за дигитален контрол не се свежда единствено до техническа модернизация, а представлява нова управленска философия, при която прозрачността и отчетността се постигат чрез технологии, данни и отворени алгоритми.

Данни, инфраструктура и национален контекст

Националната система ЦАИС ЕОП като основа за дигитален контрол

България разполага с една от най-структурираните платформи за електронни обществени поръчки в Европейския съюз – Централизираната автоматизирана информационна система „Електронни обществени поръчки“ (ЦАИС ЕОП). Системата е разработена и поддържана от Агенцията по обществени поръчки (АОП) и осигурява пълна дигитализация на възлагателния процес, включително публикуване на обявления, подаване на оферти, сключване и изпълнение на договори. ЦАИС ЕОП съдържа богат масив от структурирани данни, който покрива всички етапи от жизнения цикъл на обществената поръчка:

- планиране и прогнозна стойност;
- откриване на процедурата и публикувани документи;
- подаване и оценка на оферти;
- възлагане на изпълнител;
- сключване и изменения на договора;
- проследяване на изпълнението.

Наличието на централизиран достъп и унифициран формат на публикуваната информация позволява автоматизирано извличане и анализ на данни, което създава предпоставки за изграждане на модели за ранно предупреждение и превенция на рискове. Допълнителен потенциал представлява възможността за експорт на данни в CSV и XLS формати, което улеснява машинната обработка и интеграцията с външни аналитични системи.

Отворени данни и инициативи за прозрачност: примерът на OpenTender.eu

В допълнение към националната система, България присъства в международни проекти за прозрачност като OpenTender.eu, създаден в рамките на инициативата DIGIWHIST и поддържан

от организацията *Open Knowledge Foundation*. Платформата предоставя възможност за сравнителен анализ между държавите членки на ЕС, като използва данни от националните портали за отворени данни и от TED (Tenders Electronic Daily). За България източникът на данни, използван от OpenTender.eu, е националният портал за отворени данни (data.egov.bg). Към момента порталът публикува два основни набора – *данни за сключени договори* и *данни за изменения на договори*, предоставяни от ЦАИС ЕОП. Въпреки че това представлява важна стъпка към прозрачност, информационният обхват е ограничен, тъй като липсват данни за планиране, процедури, подадени оферти и реално изпълнение. Освен това данните не се обновяват в реално време, а последната налична информация в OpenTender.eu за България е към края на 2023 г. Тази структура на данните е недостатъчна за създаване на дигитални двойници, които изискват пълна и динамична картина на възлагателния процес. Без информация за предварителните етапи (обявяване, оферирание) и за последващото изпълнение е невъзможно да се изгради симулационен модел, който да проследява реалните зависимости и отклонения. Следователно, за да бъде приложен концептуалният модел на интелигентен контрол, е необходимо публикуване на данни във формат, който осигурява цялостен и актуален жизнен цикъл на поръчката.

Ролята на стандарта OCDS за изграждане на пълна картина на процеса

Решението на този проблем се съдържа в стандарта OCDS (Open Contracting Data Standard), разработен от *Open Contracting Partnership* и прилаган от редица държави и международни организации.

OCDS дефинира унифициран модел за структуриране и публикуване на данни за обществени поръчки, който обхваща всички етапи от жизнения цикъл – *планиране, откриване, оферирание, възлагане, договаряне, изпълнение и изменения на договорите*.

Основното предимство на OCDS е, че позволява:

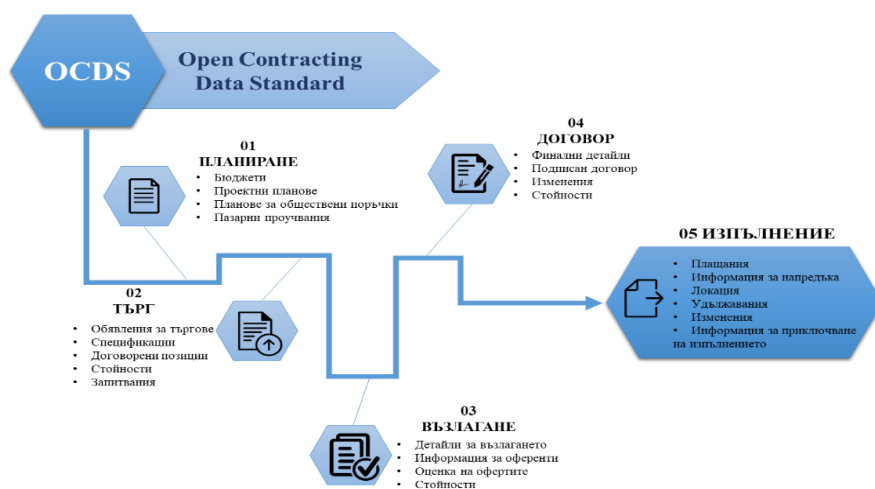
- интеграция между различни системи (напр. национални платформи и европейски бази данни);
- машинно четими данни, подходящи за анализ и автоматизирано извличане на рискове;
- повторна употреба на данни за изследвания, визуализации и граждански мониторинг.

Въвеждането на OCDS в България, което предстои като функционалност в ЦАИС ЕОП, ще осигури необходимата инфраструктура за прилагане на модели, базирани на дигитални двойници и IoT.

Този формат ще позволи създаването на национална база данни с реално-времеви актуализации, която може да бъде интегрирана с аналитични и симулационни системи, поддържащи интелигентен контрол и превенция на корупционни практики.

Българският контекст спрямо PPDS: предизвикателства и възможности

Липсата на интеграция на България в Public Procurement Data Space (PPDS) не е технологичен, а по-скоро политически и управленски въпрос. Към настоящия момент PPDS извлича информация за страната единствено от TED, което означава, че процедурите под европейските прагове остават извън обхвата на европейския анализ. В същото време националната система ЦАИС ЕОП предоставя пълни и детайлни данни за тези поръчки. Това обстоятелство може да се разглежда не като слабост, а като възможност за създаване на специфичен национален модел, който отразява особеностите на българската практика и може да служи като пример за интегриране „отдолу нагоре“ (*bottom-up integration*) в PPDS. По този начин България би могла да предложи иновативен подход, при който отворените данни и стандартът OCDS се използват за създаване на реално действащ дигитален близък на обществените поръчки – модел, който впоследствие може да бъде синхронизиран с европейската инфраструктура.



Източник: Систематизация на автора

Фигура 1: OCDS – пълен жизнен цикъл на поръчката

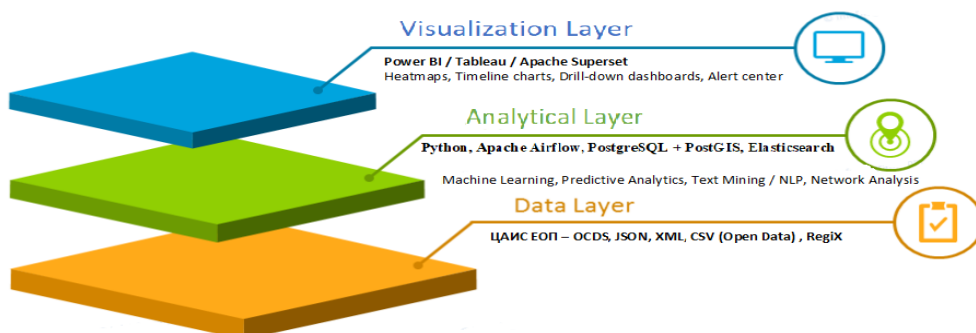
Предложен модел за интелигентен контрол

Концепция и архитектура на модела

Предложеният модел за интелигентен контрол в обществените поръчки се основава на интеграция между структурирани административни данни и сензорна информация от реалното изпълнение на проектите. Неговата основна цел е да се създаде дигитален двойник на обществената поръчка, който отразява в реално време състоянието и напредъка на процесите по възлагане и изпълнение.

Моделът включва три основни слоя:

- Данни и интеграция (Data Layer) – обединява данни от ЦАИС ЕОП, OCDS, националния портал за отворени данни и външни сензорни източници (напр. IoT устройства за мониторинг на строителни дейности, доставки или услуги). Този слой гарантира съвместимост и синхронизация между административните и оперативните потоци от данни.
- Аналитичен слой (Analytical Layer) – реализира обработка и анализ чрез алгоритми за машинно обучение и предиктивен анализ, насочени към идентифициране на рискови модели и отклонения. Тук се изчисляват рискови индикатори, „червени флагове“ и се извършват симулации върху дигиталния близък на поръчката.
- Визуализационен и контролен слой (Visualization & Control Layer) – осигурява визуално представяне на състоянието на поръчките чрез табла (dashboards), карти и времеви линии. Данните могат да бъдат достъпни както за контролни институции (КПК, КЗК, ГДБОП), така и за граждански организации и медии в режим на прозрачност.



Източник: Систематизация на автора

Фигура 2: Архитектура на модела за интелигентен контрол в обществените поръчки

Ключови рискови индикатори и „червени флагове“

В основата на аналитичния слой стоят рискови индикатори, които се изчисляват въз основа на исторически данни и текуща информация от изпълнението. Някои от индикаторите са общовалидни за всички държави членки на ЕС, докато други са специфични за българския контекст. Примери за ключови индикатори включват:

- ✓ Договор, сключен без конкуренция
- ✓ Пряко договаряне без обявление
- ✓ Критерии за избор само най-ниска цена
- ✓ Повтарящи се изпълнители
- ✓ Необичайни изменения на договор
- ✓ Минимална разлика или равенство между прогнозни и договорни стойности
- ✓ Голяма разлика между цената на победителя и останалите участници
- ✓ Цена за наем vs. Цена за покупка на същата стока
- ✓ Един участник подава оферти с голяма разлика в цените при различни възложители
- ✓ Съкратени срокове за оферти, избор на изпълнител и/или за изпълнение
- ✓ Непубликуване на договор и/или приложенията

Всеки индикатор може да бъде нормализиран и оценен по скала на риска (например 0–1), което позволява агрегиране на данните в композитен индекс на интегритета на обществените поръчки. При превишаване на зададени прагове системата може автоматично да генерира аларми („red flags“) и уведомления за отговорните институции.

Дигиталният двойник като инструмент за симулация и прогнозен контрол

Дигиталният двойник на обществената поръчка представлява виртуална реплика на реалния процес, която се актуализира непрекъснато чрез данни от административни и IoT източници. Чрез него могат да се извършват симулации на различни сценарии – например:

- какви ще бъдат ефектите при забавяне на изпълнението с 20%;
- как би се отразила промяна в прогнозната стойност или критериите за оценка;
- какви вероятности от нарушение или конфликт на интереси произтичат от повтаряемостта на участниците.

Тази симулационна среда позволява прилагане на предиктивен контрол – откриване на потенциални нарушения преди те да се случат, а не след приключване на процедурата. Освен това, дигиталният близък може да се използва като инструмент за обучение и тестване на антикорупционни алгоритми, което е особено ценно за бъдещи системи на изкуствен интелект в публичния сектор.

Институционални и граждански приложения

Моделът предлага двупосочна употреба:

- Институционален контрол – предоставяне на данни и визуални панели на органи като Комисията за защита на конкуренцията (КЗК), Комисията за противодействие на корупцията (КПК), Главна дирекция „Борба с организираната престъпност“ (ГДБОП) и Сметната палата. Тези институции могат да използват системата за анализ на картелни практики, отклонения в ценообразуването или нередности при възлагането.

- Граждански мониторинг – достъп до обобщени визуализации и отворени данни чрез публичен портал. Това ще насърчи участието на гражданското общество и ще засили доверие в системата на обществените поръчки.

По този начин моделът реализира концепцията за „умно управление на интегритета“, при което технологиите, данните и прозрачността действат като взаимно подсилващи се фактори.

Очаквани ползи и принос

Прилагането на модела би имало няколко стратегически ползи:

- Преминаване от реактивен към предиктивен контрол;
- Намаляване на човешкия фактор при оценката на риска;

- Повишаване на доверието в обществените разходи чрез достъпна и разбираема визуализация на данните;
- Създаване на база за бъдещи интеграции с PPDS и европейски инициативи за прозрачност.

В научен аспект приносът се състои в разработването на интегративен модел, съчетаващ принципите на IoT, дигиталните двойници и отворените данни в обществените поръчки – област, която до момента не е разглеждана системно в българската литература.

Заклучение и насоки за бъдещи изследвания

Дигиталната трансформация на обществените поръчки представлява един от най-значимите процеси в развитието на съвременното публично управление. Настоящият доклад очерта концептуален модел, базиран на интеграцията между Интернет на обектите (IoT), дигиталните двойници (Digital Twins) и отворените данни, чрез който контролът върху обществените поръчки може да бъде трансформиран от реактивен в интелигентен, превантивен и прозрачен процес. Предложената рамка демонстрира как съчетаването на реално-времеви данни от изпълнението на договорите с структурирана информация от електронните платформи (ЦАИС ЕОП) позволява изграждането на динамичен дигитален близък на всяка обществена поръчка. Този близък не само отразява текущото състояние на договора, но и дава възможност за симулации, прогнозен анализ и автоматично откриване на рискове чрез алгоритми за машинно обучение. България притежава значителен потенциал за развитие на подобен модел. Наличието на пълна електронна система за възлагане (ЦАИС ЕОП), възможностите за експорт в отворени формати (CSV, XLS) и предстоящото въвеждане на OCDS (Open Contracting Data Standard) създават стабилна основа за национален подход към дигиталния контрол. Липсата на пълна интеграция с Public Procurement Data Space (PPDS) следва да се разглежда не като слабост, а като възможност за разработване на независим национален модел, който да служи като пилотен пример за адаптиране към европейските рамки. Предимствата на предложената система са многопластови:

- Повишаване на прозрачността и отчетността на възлагащите органи;
- Намаляване на възможностите за корупционни практики и картелни съглашения;
- Улесняване на контролните институции (КПК, КЗК, Сметна палата, ГДБОП) чрез автоматизирани анализи и визуализации;
- Активизиране на гражданското общество чрез достъп до разбираеми данни и визуални панели.

Научният принос на изследването се състои в предложението на иновативен интегриран модел за дигитален контрол, който съчетава технически, организационни и етични измерения на управлението на обществени ресурси. Той поставя основите за изграждането на нова парадигма на контрол – базирана на данни, алгоритми и предсказуемост, вместо на ръчен одит и последваща реакция.

В заключение, интелигентният контрол в обществените поръчки не е просто технологична новост, а необходима управленска трансформация, която поставя данните, прозрачността и доверието в основата на публичните политики. Внедряването на подобен модел в България би могло да превърне страната в иноватор в областта на дигиталния публичен контрол в рамките на Европейския съюз.

Литература

1. Агенция по обществени поръчки (АОП). (2024). *Централизирана автоматизирана информационна система „Електронни обществени поръчки“ (ЦАИС ЕОП)*. <https://eop.bg>
2. Министерство на електронното управление. (2024). *Национален портал за отворени данни – Data.egov.bg*. <https://data.egov.bg>

3. Open Knowledge Foundation. (2023). *OpenTender.eu – Public Procurement Analytics Platform*. <https://opentender.eu>
4. Open Contracting Partnership. (2024). *Open Contracting Data Standard (OCDS) Documentation*. <https://standard.open-contracting.org>
5. European Commission. (2023). *Public Procurement Data Space (PPDS): Towards a European Data Infrastructure for Public Procurement*. Brussels: European Commission.
6. OECD. (2022). *Integrity and Transparency in Public Procurement: Using Data for Accountability*. Paris: OECD Publishing.
7. World Bank. (2023). *Digital Government and Data-Driven Integrity*. Washington, D.C.: The World Bank Group.
8. Kritzinger, W., Karner, M., Traar, G., Henjes, J., & Sihn, W. (2018). *Digital Twin in Manufacturing: A Categorized Literature Review and Classification Framework*. IFAC-PapersOnLine, 51(11), 1016–1022.
9. Atzori, L., Iera, A., & Morabito, G. (2017). *The Internet of Things: A Survey*. Computer Networks, 54(15), 2787–2805.

НЕВИДИМИТЕ БАРИЕРИ ВЪВ ВИЗУАЛИЗАЦИЯТА НА ДАННИ: КАК ДОСТЪПНОСТТА ОФОРМЯ ПОТРЕБИТЕЛСКОТО ИЗЖИВЯВАНЕ

The Invisible Barriers in Data Visualization: How Accessibility Shapes the User Experience

Лиляна Петрова Сапунджиева¹,
e-mail: lilyanasapundzhieva@gmail.com

Абстракт

Визуализацията на данни е един от най-разпространените и мощни инструменти за предаване на информация, но често изключва предназначената за нея публика. „Невидимите бариери във визуализацията на данни: как достъпността оформя потребителското изживяване“ изследва как дизайнерските решения като използване на цветове, интерактивност, етикетирание и сложност – могат неволно да създадат препятствия за потребители със зрителни, когнитивни или моторни увреждания. Изхождайки от принципите на приобщаващия дизайн и стандартите за достъпност, докладът откроява основните UX предизвикателства, с които хората с увреждания се сблъскват при работа с обичайни графики, табла и инфографики. Чрез примери и потребителски прозрения се показва, че достъпната визуализация не е просто изискване за съответствие със стандартите, а двигател за по-добър, по-ясен и по-съпричастен дизайн. В крайна сметка изследването аргументира, че чрез идентифициране и премахване на тези „невидими бариери“ дизайнерите могат да създават преживявания с данни, които са по-ефективни за всички.

Abstract

Data visualization is one of the most widespread and powerful tools for conveying information, yet it often excludes the very audience it is meant for. “Invisible Barriers in Data Visualization: How Accessibility Shapes User Experience” explores how design choices—such as color use, interactivity, labeling, and complexity—can unintentionally create obstacles for users with visual, cognitive, or motor impairments. Building on the principles of inclusive design and accessibility standards, the report highlights key UX challenges that people with disabilities face when interacting with common charts, dashboards, and infographics. Through examples and user insights, it demonstrates that accessible visualization is not merely a compliance requirement but a driver of better, clearer, and more empathetic design. Ultimately, the study argues that by identifying and removing these “invisible barriers,” designers can create data experiences that are more effective for everyone.

Ключови думи: Information and Internet Services, Computer Programs, Communication, Disability, Technological Change

Невидимите бариери във визуализацията на данни: как достъпността оформя потребителското изживяване

Въведение – Визуализирането на данни и допирните точки с достъпността

Визуализацията на данни е сред най-въздействащите начини за предаване на информация в съвременния дигитален свят. Диаграмите, таблата и инфографиките превръщат сухите числа в истории, които помагат за откриването на закономерности, за взимането на решения, за осмислянето на сложни процеси. В ерата на едно такова информационното пренасищане именно

¹ Докторант в катедра ИТК, УНСС, e-mail: lilyanasapundzhieva@gmail.com

визуалната комуникация се превърна в универсален език между данните и човека. *Ефективната визуализация не просто представя данни, а изгражда мост между знанието и разбирането (1).*

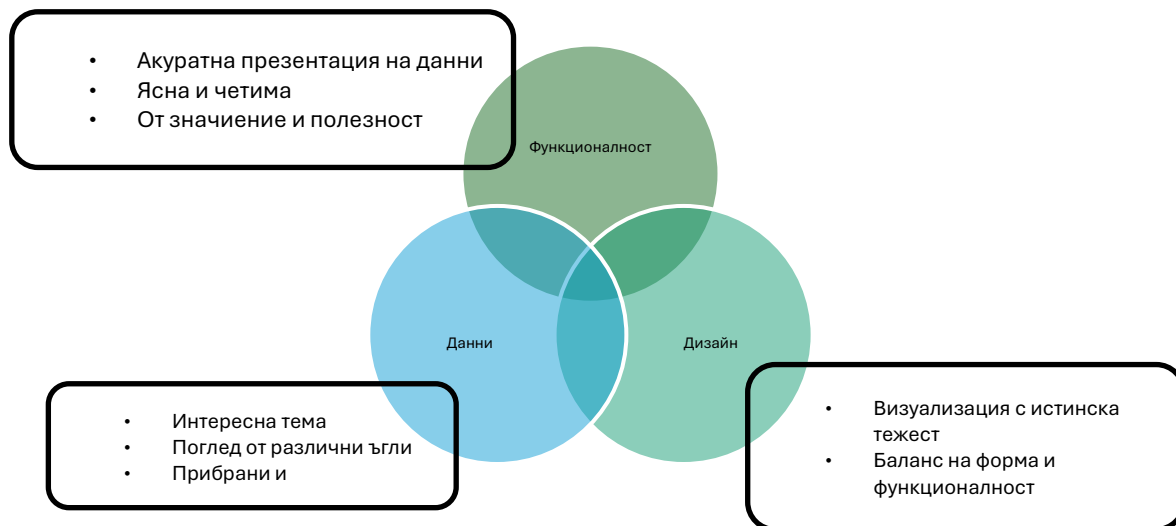
Ала всъщност този мост невинаги е достъпен за всички потребители. Дизайнът, който се основава единствено на визуални елементи — цвят, форма, анимация — може да се превърне в бариера за хора със зрителни, когнитивни и/или двигателни затруднения. Примерите са многобройни и навсякъде: графики, които разчитат само на червено и зелено за разграничаване на категории — те използват символическото значение на тези два цвята, но също така откъсват потребителите с цветна слепота; диаграми с нисък контраст; интерактивни табла, които не могат да се използват без мишка като средство за навигация по визуализацията. Според изследването на Elavsky, Bennett и Moritz (2022), повечето визуализации в Интернет днес нарушават основни принципи на достъпност, което прави данните „видими“ само за определена част от обществото.

Ролята на цветовете, етикетите, интерактивността и структурата в дизайна на визуализации е ключова, което обаче не бива да означава, че те не трябва да бъдат използвани с внимание към различните начини, по които хората възприемат и обработват информацията. Както отбелязва Kent Eisenhuth (Smashing Magazine), прилагането на принципа *accessibility-first* не само помага на хората с увреждания, но подобрява потребителското изживяване за всички — включително в ситуации на ниско осветление, умора или ограничен достъп до устройства.

Достъпността във визуализацията на данни не е само техническо изискване, а главна част от дизайн принципите, основани на емпатия и всеобщо разбиране. Тя насърчава мисленето за разнообразието от човешки възприятия още в началните етапи на съставянето на презентацията на данни. Международните стандарти като WCAG и ARIA осигуряват рамка за това мислене, но реалната промяна настъпва, когато дизайнерите на потребителските изживявания започнат да виждат достъпността не като задължение, а като неразделна част от добрия дизайн.

„Визуализацията на данни трябва да говори на всички, не само на тези, които могат да я видят.“

Тази мисъл стои в основата на темата на настоящия доклад — как могат да се премахнат или поне намалят невидимите бариери във визуализацията на данни, така че информацията да бъде разбираема, използвана и достъпна за всеки човек, независимо от техните способности.



Фигура 1: Елементите на една добра визуализация

Визуализацията на данни под формата на диаграми, табла (дашбордове) и инфографики

предава послания чрез използване на цветове, интерактивност, етикетирание и сложност на потребители с зрителни, когнитивни и/или двигателни затруднения прилагайки принципите на инклузивния дизайн и спазването на стандартите за достъпност.

Проблемът – Невидимите бариери във визуализацията на данни

Какво точно са бариерите? Бариера може да се определи просто като всяко препятствие, пречка или спънка, която възпрепятства хората да правят това, което искат или им е необходимо да правят. Пример за това може да бъде опит да бъде пресечена една река – понякога няма мост, понякога мостът е счупен, а понякога има знак за забрана на преминаването на моста. За хората и специфично децата с интелектуални увреждания тези бариери могат значително да повлияят на тяхната способност да учат, да се развиват и да участват пълноценно в обществото.

Бариерата обикновено се разделя на три основни категории: социални бариери (коренящи се в нагласи и предразсъдъци), психологически бариери (възникващи от вътрешни конфликти) и физически бариери (създадени от заобикалящата среда. Всяка категория представя уникални предизвикателства, но разбирането им е първата стъпка към тяхното преодоляване.

• Социалните бариери

В своята същност социалните бариери възникват, когато обществото възприема увреждането като „разлика“, а не като част от човешкото разнообразие. Тази перспектива създава верижна реакция на изключване. Например, когато дете със синдром на Даун влезе в класната стая, вниманието често се насочва към това, което то не може да прави, вместо към потенциала и уменията, които носи със себе си. Този тип мислене оформя начина на взаимодействие, ограничава възможностите и намалява смисленото участие в общността. В дългосрочен план, социалните бариери поддържат култура на разделение, в която различието се възприема като предизвикателство, а не като ценност.

• Психологическите бариери

Психологическите бариери се коренят в вътрешните нагласи и емоционални реакции – както на хората с увреждания, така и на тяхната среда. Те често се проявяват под формата на:

- **Оттегляне от социални ситуации:** Избягване на взаимодействия, за да се предотврати потенциално неудобство или отхвърляне.
- **Научена безпомощност:** Вярата, че усилията няма да доведат до успех, което води до отказ от опити.
- **Страх от провал:** Избягване на предизвикателства, за да се защити крехкото самочувствие.
- **Негативен вътрешен диалог:** Мисли, които подсилват ограниченията, вместо възможностите.

Трагедията на психологическите бариери е, че те често ограничават човека повече от самото увреждане. Например, дете може да има когнитивните способности да усвои нови умения, но да бъде задържано от собствените си съмнения и страхове.

• Техническите бариери

Техническите бариери представляват практическите и инфраструктурните ограничения, които възпрепятстват пълноценното участие в обществото. Те са най-видимата форма на „невидими“ препятствия, защото често се възприемат като незначителни детайли в дизайна или планирането – докато не засегнат някого пряко.

Физическите бариери включват липсата на достъп до пространства и услуги:

- **Архитектурни препятствия:** Сгради с единствено стълби, тесни врати, високи плотове или лошо осветление.
- **Проблеми с транспорта:** Автобуси без подемници за инвалидни колички, метро станции без асансьори, паркинги без обозначени достъпни места.
- **Недоглеждания в градското планиране:** Тротоари без наклони, пешеходни сигнали без аудио указания, временни строежи, блокиращи достъпни маршрути.

В дигиталната среда бариерите често са по-незабележими, но не по-малко сериозни:

- Уеб достъпност: Уебсайтове без алтернативен текст за изображения, без съвместимост с екранни четци или с нисък контраст.
- Софтуер и интерфейси: Приложения, които не поддържат навигация чрез клавиатура или гласови команди.
- Информационна недостъпност: Видеа без субтитри или аудиоописания, документи без структурирани заглавия и метаданни.

Разбулените невидими бариери:

1. Дизайнерите на визуализацията на данни зависят много на символизма на цветовете, които използват
 2. Ниския контраст за създаване на естетика на презентацията за сметка на добрия и достъпен прочит на данните
 3. Табла с прекалено много елементи и сложни оформления, които биха създали дискомфорт и дават чувство за пренатруфеност на екрана
 4. Липсващ alt текст, тъй като той не се проектира на екрана
 5. Анимациите, които отклоняват вниманието на потребителя от същинското съдържание на визуализацията
 6. Липсата на поддръжка на клавиатура за пълното изживяване на потребителя
 7. Недостатъчен контекст при презентирането на визуализацията на данните
- “Проблемите с достъпността не засягат само хора с увреждания – те те товарят когато си под стрес или си уморен, когато използваш мобилното си устройство или дори когато си в лошо осветена стая“*

Значението на достъпността

Достъпността във визуализацията на данни е ключов аспект на съвременния дизайн, който не трябва да се разглежда като нишов проблем, а като глобално предизвикателство. Според данни на Световната здравна организация около 15% от населението на света живее с някаква форма на увреждане, което означава над 1 милиард души, които могат да срещнат бариери при разбирането на информация, която мнозина от нас приемат за даденост. Това подчертава, че визуализациите на данни, които не вземат предвид нуждите на всички потребители, могат да изключат значителна част от аудиторията още в самото начало(2).

Достъпността обаче не се ограничава само до хора с постоянни увреждания. Тя засяга и всеки потребител в различни ежедневни ситуации. Например, около 1 на 12 човека от мъжки пол има някаква форма на дефицит на цветово зрение, което прави диаграми с червено-зелени комбинации трудни за разбиране (3). Също така фактори като стрес, умора или гледане на мобилно устройство при силна слънчева светлина могат временно да намалят способността на всеки да възприема и интерпретира данни.

Истината е, че достъпността подобрява преживяването за всички. Ясните цветови контрасти, опростените оформления, четимият текст и описателните надписи не само подпомагат хората с увреждания, но и улесняват разбирането на данните за всички потребители, независимо от контекста (4). Достъпният дизайн създава доверие, показва емпатия и отразява висок професионализъм. Колкото повече хора могат да видят и разберат визуализацията на данните, толкова по-силно и въздействащо става посланието, което тя предава.

В този смисъл, достъпността не е просто въпрос на спазване на стандарти или регулации, а интегрална част от добрия дизайн. Тя гарантира, че данните достигат до максимално широка аудитория, превръщайки информацията в истинска сила за информиране и вземане на решения.

„Когато дизайнът работи за хора с често срещани проблеми, той работи по-добре за всички нас“

Проучвания и примери за важността на достъпността във визуализацията на данни

Представените в този доклад са две проучвания да различни теми, с общ корен – как хората с увреждания възприемат визуализациите на данни и как те успяват да работят с тях най-ефективно.

Проучване 1 - Предизвикателства с данните за хора с интелектуални и развойни затруднения (IDD) (5)

Ограничения, свързани с увреждания и данъчна грамотност

Всеки човек с интелектуални и развойни увреждания (ID/IDD) има уникални силни и слаби страни, което води до различни преживявания при работа с данни. Много хора с IDD нямат възможност или мотивация да развият данъчната си грамотност поради липса на образование или вредни стереотипи. Участниците често съобщаваха, че данните са трудни за разбиране и обработка. Например, интервюирания във интервю 2 казва: *"То е като когато лекар говори с теб на медицински термини – почти трябва да кажеш: обясни ми по-просто."*

Интервюирания в интервю 11, застъпник със значителен опит, посочва, че липсата на данъчна грамотност е сериозно предизвикателство и обучението по данни трябва да се интегрира в ежедневното решаване на проблеми. Различните типове IDD и нивата на тежест водят до различни трудности при обработката на данни: хора с аутизъм са податливи на когнитивно претоварване, а други, като P8 с Fetal Alcohol Syndrome, имат проблеми със задържането на информация.

Интеграцията на данни от различни източници също е предизвикателство, а някои участници изразяват тревога за точността на данните и дезинформацията. Основни предизвикателства са:

1. Данните са абстрактни и хората с IDD имат малко или никакво образование, за да ги разбират.
2. Работата с данни изисква когнитивни умения като задържане, припомняне и интегриране на информация, в които хората с IDD срещат трудности.
3. Обучението по данни често се случва в изолация, а не в реалния живот, което води до ниска осведоменост и интерес.

Ролята на грижещите се

Хората с по-сериозни увреждания често разчитат на грижещи се за достъп до данни. Източниците на информация обикновено са хора, а не технологии, което намалява личната им автономия. P9 получава информация предимно чрез помощта на другите, а P7 и P14 показват уязвимост при вземане на решения и потенциална експлоатация.

Освен това, грижещите се и застъпниците използват данните като инструмент за авторитет и разказване на истории в здравеопазване, образование и правна среда. Визуализацията на данни може да подсили посланията, но липсата на обучение често пречи на създаването на графики и истории, базирани на данни. Съвместни и „безтехнологични“ визуализации – например визуални календари, сортиране на дрехи или графично контролиране на порции – могат да подобрят данъчната грамотност, ангажираността и осведомеността при вземане на ежедневни решения.

При комуникационни трудности съвместната визуализация може да насърчи размисъл и критично мислене между хората с IDD и околните, създавайки пространство за общи разговори и взаимно разбиране чрез данни.

Проучване 2 - Защо комбинирането на текст и визуализация може да подобри байесовото разсъждение: перспектива на когнитивното натоварване (6)

Целта на изследването е била да се проучи ефектът от комбинирането на текст и визуализация върху байесовото разсъждение и когнитивното натоварване. Резултатите показват, че самият формат на представяне оказва малко до средно влияние върху точността на байесовото разсъждение, като визуализациите сами по себе си подобряват резултатите в сравнение с текстовото съдържание. Комбинирането на текст и визуализация води до леко, но значимо подобрене спрямо използването само на текст. Що се отнася до когнитивното натоварване, включването на визуализации не увеличава трудността на задачите, а в някои случаи дори намалява субективното усещане за усилие. Това подкрепя заключението, че визуализациите ефективно допълват текста, подпомагайки по-доброто разсъждение, без да претоварват потребителите. Важно е да се отбележи, че самите тенденции в точността не обясняват напълно наблюдаваните подобрения, като когнитивните фактори играят ключова роля.

Методи за достигане на добро ниво на достъпност при визуализацията на данни

Съществуват няколко ключови принципа за създаване на достъпни визуализации на данни, базирани на утвърдени насоки за достъпност. Следвайки ги, дизайнерите на визуализации на данни би трябвало частично, ако не и изцяло да решат проблемите за четимост на данните. (7)

1. Предоставете текстово резюме на визуализацията.

Текстовото описание помага на потребители, които намират визуализацията за объркваща, и осигурява информация за тенденции и модели на данните за хора с нарушено зрение. Освен това улеснява индексването от търсачките.

2. Достъпна таблица с данни.

Данните трябва да са налични в лесно четим формат на таблица, който е удобен както за визуални, така и за невизуални потребители, и може да подобри резултатите от търсенето.

3. Достатъчен контраст и отделяне на елементите.

Контрастът между елементите и фона е от съществено значение. W3C WCAG 2.1 препоръчва поне 3:1 за големи елементи и 4,5:1 за малък текст. Разделяне чрез граници, пространство, модели или линии може да увеличи четливостта, но трябва да се използва умерено.

4. Не разчитайте само на цвета за предаване на информация.

Използвайте етикети, символи и подсказки, за да подпомогнете хора с различни форми на далтонизъм. Около 8% от мъжете по света имат проблеми с възприемането на цветовете, затова цветовете трябва да са лесно различими.

5. Осигурете четлив текст.

Избягвайте малък шрифт, осигурете контраст и разстояние между редовете. Предпочитат се sans serif шрифтове като Verdana и Helvetica, с избягване на курсив.

6. Използвайте прост и разбираем език.

Ясният, ясен и кратък език помага на широк кръг потребители, включително такива с когнитивни затруднения или несъвършен владеещ езика.

7. Предпочитайте познати и прости визуализации.

Сложните или новаторски графики могат да объркат потребителите. Ако използвате такъв тип визуализация, осигурете текстово описание и обмислете разделяне на големи графики на по-малки.

8. Интерактивните визуализации трябва да са достъпни чрез клавиатура.

Това улеснява потребители с моторни затруднения и осигурява достъпност чрез различни помощни технологии.

9. Позволете изключване на анимации и движението.

Някои потребители, например с вестибуларни нарушения, могат да се разсейват или да изпитват дискомфорт. Уважавайте настройките на „prefers-reduced-motion“ в CSS.

10. Осигурете отзивчивост на различни размери на екрана и нива на увеличение.

Съдържанието трябва да се адаптира към различни екрани и увеличения, за да избегне изрязване или припокриване на елементи, а интерактивните графики могат да се преоформят, например вертикална графика да стане хоризонтална на малък екран.

Стандарти и средства за достъпно визуализиране на данни

Стандартите, посочени по-долу, са основните насоки за осигуряване на достъпност при дигитално съдържание и визуализации на данни, обхващат както световно, така и европейско ниво и гарантират адекватното презентиране на данни за всички потребители.

WCAG 2.1 / 2.2 (Web Content Accessibility Guidelines)(8)

Издадени от W3C, тези насоки представляват глобален стандарт за дигитална достъпност. В контекста на визуализации те се фокусират върху четири основни принципа: възприемане, оперативност, разбираемост и надеждност. За възприемане се изисква достатъчен цветови контраст, предоставяне на текстови алтернативи и субтитри за мултимедия. Оперативността включва възможност за навигация с клавиатура и избягване на съдържание, което може да причини пристъпи или дискомфорт. Разбираемостта се гарантира чрез ясна структура, логични етикети и навигация, а надеждността чрез съвместимост с асистивни технологии. Ключов извод е, че визуализациите трябва да осигуряват правилен контраст, текстови еквиваленти и пълна навигация с клавиатура или екранен четец.

ARIA (Accessible Rich Internet Applications)(9)

Също издадена от W3C, ARIA насочва към достъпността на динамично съдържание като табла и интерактивни графики. Основните практики включват използване на ARIA роли за описание на графики, добавяне на ARIA етикети и описания за различни типове графики и данни, както и ARIA Live Regions за обявяване на промени в интерактивните визуализации.

Section 508 (U.S. Rehabilitation Act)

Този стандарт, издаден в САЩ, има за цел дигиталното съдържание и софтуер във федералния сектор да бъдат достъпни. Той често се приема като базов стандарт по света. В контекста на визуализации изискването е данните да са достъпни чрез текст, клавиатура и асистивни технологии.

EN 301 549 (EU Accessibility Standard)(10)

Издаден от Европейския съюз, този стандарт определя изисквания за ICT достъпност, съвместими с WCAG и Section 508. Той обхваща публични сайтове и аналитични инструменти в ЕС, като гарантира, че цифровите решения са достъпни за всички потребители.

ISO 9241-171 (Ергономика на човеко-системното взаимодействие)

Този международен стандарт от ISO се фокусира върху достъпността на софтуерни интерфейси. Практическата насока е дизайнът да осигурява визуална, слухова и моторна достъпност, което е особено важно при създаването на инструменти за визуализация на данни.

Принципите на инклузивния дизайн насочват към създаване на дигитални продукти, достъпни за всички потребители, независимо от техните способности или ограничения. Те препоръчват проектиране както за ситуационни, така и за постоянни увреждания, тъй като „Проектирането за достъпност не е само за хора с постоянни увреждания; то е полезно за всички, включително за тези с ситуационни ограничения“(11). Осигуряването на гъвкавост в представянето на съдържанието, например чрез текстови обобщения за визуално съдържание, и избягването на предположението, че потребителят възприема информацията само чрез един сензорен канал, са ключови аспекти на подхода. Универсалният дизайн за учене (UDL)

разширява тези принципи, като предлага различни начини за представяне на информацията – чрез текст, звук или тактилна обратна връзка, защото „Универсалният дизайн за учене осигурява множество начини за представяне, ангажиране и изразяване, за да подкрепи различни учащи се“ (12). Допълнително, контролни списъци за достъпност на визуализациите на данни, създадени от различни инструменти и организации, обобщават препоръките на WCAG и добрите UX практики, улеснявайки прилагането на инклузивен дизайн в практиката.

Заклучение – от просто данни към достъпна визуализация на тях от всички нас

Достъпността е съществена част както от дигиталния, така и във физическия свят, която гарантира пълноценно участие на всички и създава приобщаващи среди, в които всеки има равен достъп до информация, услуги и възможности. Приемането на принципите на достъпност позволява на дизайнерите и разработчиците да създават по-иновативни и удобни решения, които водят към по-инклузивен и справедлив свят.

Въпреки това, дизайнът на визуализации често е ориентиран към хора без увреждания, което ограничава възможността на лица с интелектуални и развиващи се разстройства да се ангажират ефективно с информацията. С нарастването на значението на данните за вземане на решения, осигуряването на достъп до информация става критично за самостоятелните решения на всички потребители. Общността на хора с различни заболявания и затруднения търси повече достъпни инструменти, а организациите и компаниите проявяват нарастващ интерес към подобряване на комуникацията. Когнитивно достъпните визуализации не само подпомагат хората със затруднения, но също така откриват нови пътища за по-добро разбиране и вземане на решения за всички потребители.

References

1. Zaheer, S., (2018). *Inclusive Data Visualization: Designing for Cognitive and Visual Accessibility*. International Journal of Engineering Technology Research & Management, 02(05). (Zaheer, S., 2018. Inclusive Data Visualization: Designing for Cognitive and Visual Accessibility. International Journal of Engineering Technology Research & Management, 02(05)).
2. World Health Organization, (2023). *Disability and Health*, Geneva, WHO. (World Health Organization, 2023. Disability and Health, Geneva: WHO).
3. Brennan, M., (2020). *Color Vision Deficiency: Implications for Design*, New York, Design Press. (Brennan, M., 2020. Color Vision Deficiency: Implications for Design, New York: Design Press).
4. Few, S., (2013). *Data Visualization: Past, Present, and Future*, Oakland, Analytics Press. (Few, S., 2013. Data Visualization: Past, Present, and Future, Oakland: Analytics Press).
5. Wu, K., Petersen, E., Ahmad, T., Burlinson, D., Tanis, E. S., & Albers Szafir, D., (2021). *Understanding Data Accessibility for People with Intellectual and Developmental Disabilities*, in Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems, Yokohama, Japan. (Wu, K. et al., 2021. Understanding Data Accessibility for People with Intellectual and Developmental Disabilities, in Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems, Yokohama, Japan).
6. Bancelhon, M., Wright, A., Ha, S., Crouser, J., & Ottley, A., (2023). *Why Combining Text and Visualization Could Improve Bayesian Reasoning: A Cognitive Load Perspective*. arXiv preprint arXiv:2302.00707. (Bancelhon, M., Wright, A., Ha, S., Crouser, J., & Ottley, A., 2023. Why Combining Text and Visualization Could Improve Bayesian Reasoning: A Cognitive Load Perspective. arXiv preprint arXiv:2302.00707).
7. Moseng, Ø., (2021). *10 Guidelines for DataViz Accessibility*. Highcharts Blog. (Moseng, Ø., 2021. 10 Guidelines for DataViz Accessibility. Highcharts Blog).

8. W3C, (2018). *Web Content Accessibility Guidelines 2.1*, W3C. (W3C, 2018. Web Content Accessibility Guidelines 2.1, W3C).
9. W3C, (2014). *Accessible Rich Internet Applications (ARIA) 1.1*, W3C. (W3C, 2014. Accessible Rich Internet Applications (ARIA) 1.1, W3C).
10. European Telecommunications Standards Institute, (2018). *EN 301 549: Accessibility requirements for ICT products and services*, ETSI. (European Telecommunications Standards Institute, 2018. EN 301 549: Accessibility requirements for ICT products and services, ETSI).
11. W3C, (2018). *Web Content Accessibility Guidelines 2.1*, W3C. (W3C, 2018. Web Content Accessibility Guidelines 2.1, W3C).
12. CAST, (2018). *Universal Design for Learning Guidelines*, CAST. (CAST, 2018. Universal Design for Learning Guidelines, CAST).

ПРИЛОЖНИ АСПЕКТИ НА ИНТЕГРАЦИЯ МЕЖДУ СИСТЕМИ ЗА УПРАВЛЕНИЕ НА СЪДЪРЖАНИЕ В ОРГАНИЗАЦИЯТА (ЕСМ) И СИСТЕМИ ЗА УПРАВЛЕНИЕ НА БИЗНЕС ПРОЦЕСИ (BPM)

Use cases for integrating Enterprise Content Management systems with Business Process Management systems

Иван Белев¹

e-mail: i.belev@unwe.bg¹

Абстракт

Представяне на практични аспекти и сценарии за комбиниране на две от технологиите за дигитализация и автоматизация в организацията - системи за управление на съдържание (Enterprise Content Management) и системи за управление на бизнес процеси (Business Process Management).

Abstract

Use cases for integration between two of the main technologies for digital business automation - Enterprise Content Management (ECM) and Business Process Management (BPM).

Ключови думи: дигитализация, автоматизация, технологии, сценарии, BPM, ECM.

JEL: O33, L86, M15

Дигитализация и автоматизация

В предишно изследване авторът разглежда термина „дигитализация и автоматизация в организацията“, който произлиза от на английския термин “digital business automation”. Изведена е дефиницията на термина като: „Използването на компютърни технологии за дигитализиране и автоматизиране на различни аспекти от дейността на организациите с цел постигане на по-висока производителност, по-добри бизнес резултати и/или по-висока клиентска удовлетвореност“. Изведени са следните дейности, които спадат към това понятие:

- дигитализация и автоматизация на процеси;
- дигитализация и управление на съдържание и документи;
- дигитализиране и управление на бизнес правила и решения;
- автоматизиране на повтаряеми стъпки;
- извличане на информация и знание за процесите в организацията (process mining);
- други.

Авторът на претендира, че списъкът с дейности по-горе е изчерпателен. В настоящото изследване се посочват някои приложни аспекти на интеграция между две от технологиите, които се използват за реализиране на посочените сценарии, а именно:

¹ Главен асистент, доктор. Катедра Информационни технологии и комуникации, Университет за национално и световно стопанство.

- системи за управление на съдържание в организацията (ECM)
- системи за управление на бизнес процеси (BPM)

Дигитализация и автоматизация на процеси

В редица предходни изследвания авторът разглежда дигитализацията и автоматизацията на процеси и технологиите за управление на бизнес процеси (Business Process Management). Често използвана дш автора дефиниция за Управлението на бизнес процеси е следната: дисциплиниран подход за откриване, проектиране, изпълнение, документиране, проследяване, контрол и измерване на автоматизирани и неавтоматизирани бизнес процеси с цел постигане на трайни и целенасочени резултати, отговарящи на стратегическите цели на организацията. Управлението на бизнес процеси може да включва също така целенасочено, колаборативно и технологично-подпомогнати действия по описание, подобрене и управление на цялостни бизнес процеси, които водят до постигането на бизнес резултати, създават стойност и позволяват на организацията да постига бизнес целите си по-лесно.

Дигитализация и управление на съдържание и документи

Авторът свързва дигитализацията и управление на съдържание и документи в организациите, особено в по-големи организации, с термина Enterprise content management (ECM) или Управление на съдържанието на предприятието. Предходни изследвания по темата водят до възможността да се дефинира следното определение: съвкупност от стратегии, методи, инструменти, процеси, умения и технологии за управление на пълния жизнен цикъл на съдържанието в организацията – създаване/извличане, откриване, съхраняване, управление, опазване, търсене/намиране, публикуване и предоставяне.

Приложни аспекти на интеграция между системи за управление на съдържание в организацията (ECM) и системи за управление на бизнес процеси (BPM)

Тази част от изследването разглежда примерни сценарии за комбиниране на описаните по-горе две технологии за дигитализация и автоматизация в организациите. Предложените приложни аспекти са изведени от опита на автора в подобни инициативи, както и на база на технологичните характеристики на избраните технологии.

Осигуряване на съдържание от организацията по време на изпълнение на бизнес процесите

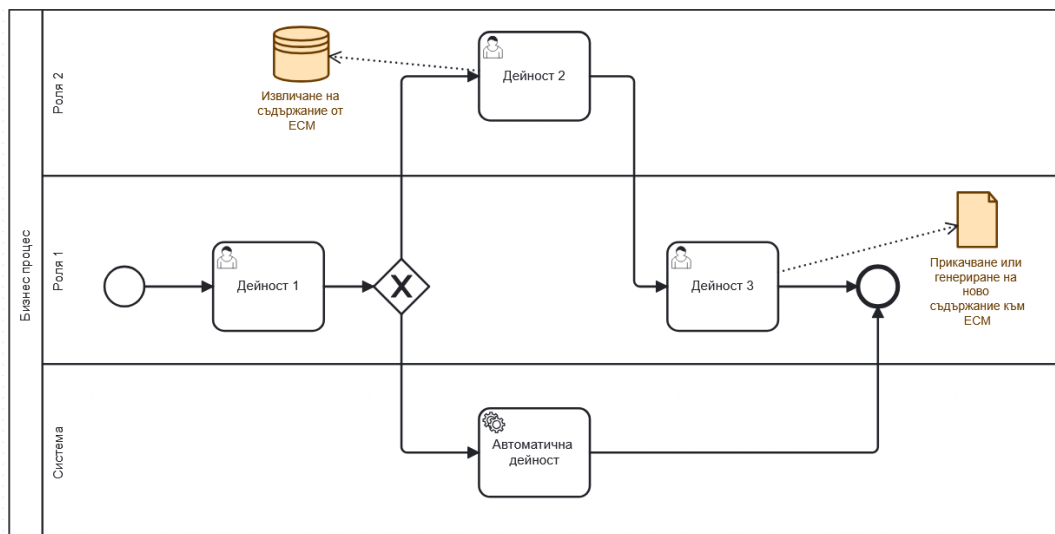
Често срещан сценарий за комбиниране на системи за управление на съдържание в организацията (ECM) и системи за управление на бизнес процеси (BPM) е осъществяването на интеграция между двата типа системи с цел осигуряване на следното:

- показване на документи и съдържание на организацията по време на стъпките на изпълнение на процесите в организацията;
- създаване и прикачване на нови документи и съдържание по време на стъпките на изпълнение на процесите в организацията.

Посочените два практически аспекта на комбиниране на двете технологии в организацията целят да се постигне максимално улеснение на потребителите при работа със съдържание на организацията по време на изпълнение на стъпките от процесите, които те изпълняват. Интеграцията за постигане на такова улеснение може да бъде в две посоки:

- от ECM към BPM чрез извличане (Read) на документи и съдържание и показването им в стъпките на процесите;
- от BPM към ECM чрез прикачване (Write) документи и съдържание от стъпките на процесите в ECM системата.

На следващата фигура (Фигура 1) е представен примерен процес, без да е описан конкретен сценарий, в който са нанесени две точки на интеграция по описаните по-горе два сценария:



Фигура 1: Осигуряване на съдържание от организацията по време на изпълнение на бизнес процесите

Управление на жизнения цикъл на документите

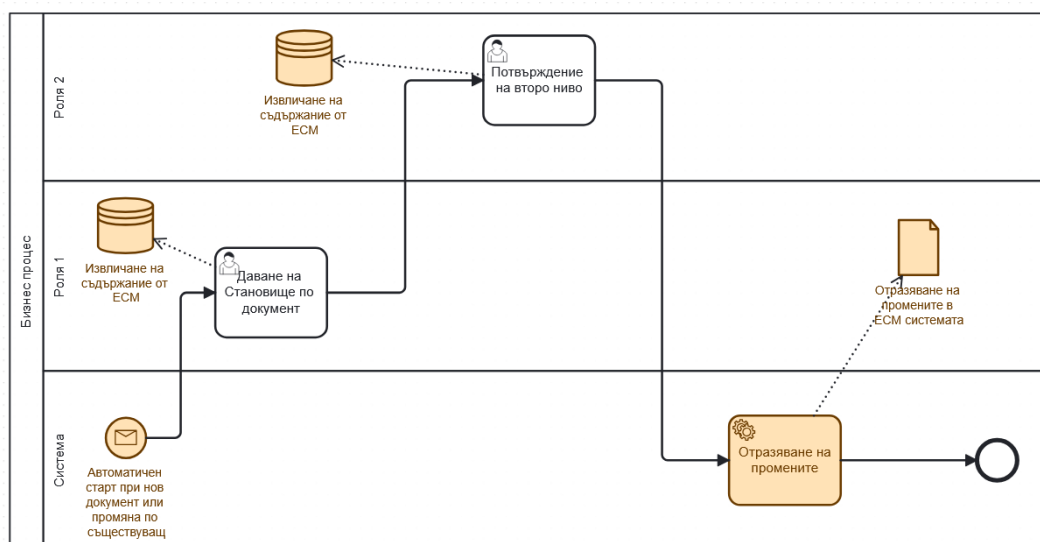
Друг често срещан сценарий за комбиниране на системи за управление на съдържание в организацията (ЕСМ) и системи за управление на бизнес процеси (BPM) е осъществяването на интеграция между двата типа системи с цел осигуряване на жизнения цикъл на съдържанието в организацията чрез:

- изпълнение на процеси по одобрение на документи и съдържание в организацията;
- изпълнение на автоматизирани процеси, свързани със съхранение, архивиране и изтриване на съдържание в организацията.

Посочените два практически аспекта на комбиниране на двете технологии в организацията целят да се постигне максимална прозрачност и гъвкавост на процесите по управление на жизнения цикъл на документите и други видове съдържание в организацията. Интеграцията за постигане на това обикновено включват:

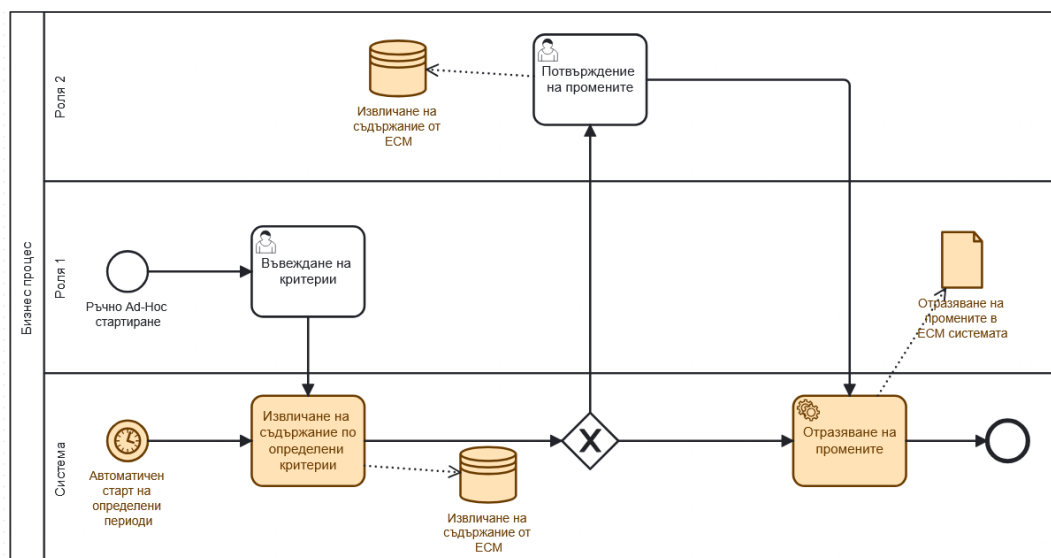
- автоматизирано създаване на процеси по одобрение или друга обработка в BPM при настъпване на дадени обстоятелства в ЕСМ системата като например: при създаване на нов документ от даден тип, който изисква одобрение преди да бъде публикуван или утвърден за използване в организацията;
- ръчно или автоматизирано създаване на процеси в BPM системата при настъпване на обстоятелства, свързани с управление на жизнения цикъл на дадени документи или съдържание – например при необходимост от регулярно архивиране или премахване на документи след изтичане на срока за тяхното съхранение.

На следващата фигура (Фигура 2) е представен примерен процес по одобрение на ново съдържание/документ или промени по съществуващ такъв. Процесът се инициира автоматично при възникването на необходимост от одобрение. Данните от ЕСМ се извличат и предоставят в стъпките на процеса. В зависимост от решението на участниците в процеса, към ЕСМ автоматизирано се отразяват крайните резултати.



Фигура 2: Изпълнение на процеси по одобрение на документи и съдържание в организацията

На следващата фигура (**Фигура 3**) е представен примерен процес по регулярни действия, свързани с премахване/архивиране на съдържание според правилата за съхранение на организацията. Процесът се иницира автоматично при възникването на определени за това периоди, както и ръчно при възникване на извънредна необходимост. Данните от ЕСМ се извличат и предоставят в стъпките на процеса. В зависимост от решението на участниците в процеса, към ЕСМ автоматизирано се отразяват крайните резултати.



Фигура 3: Ръчно или автоматизирано създаване на процеси по управление на жизнения цикъл на документи или съдържание

Получаване на цялостен поглед върху обект в организацията

Авторът разглежда още един често срещан сценарий за комбиниране на системи за управление на съдържание в организацията (ЕСМ) и системи за управление на бизнес процеси (BPM), а именно – за осигуряване на цялостен поглед върху даден обект в организацията (т.нар. 360° View). Това на практика означава да се използва системата за бизнес процеси, в която да се комбинира информация от всички свързани системи в организацията, които дават данни за даден

обект, включително и ЕСМ системата, от която да се извлекат всички документи и съдържание, свързано с този обект. След събиране на всички данни те се предоставят на едно място (в стъпките на процесите на организацията) на потребителите, които трябва да разполагат с тази информация.

Като заключение на изложеното може да се каже, че комбинирането на технологиите за управление на бизнес процеси (BPM) и технологиите за управление на съдържание (ЕСМ) в организацията е възможно в различни варианти, които може да доведат до синергичен ефект от употребата им. В следващи разработки авторът има интерес да продължи да анализира възможностите за комбиниране на технологии за дигитализация и автоматизация в организацията в други възможни варианти и комбинации за постигане на още по-голям ефект.

References

1. Belev, Ivan, (2024), *Combining digital business automation technologies*, Innovative Information Technologies for Economy Digitalization (IITED), issue 1, p. 305-312, <https://EconPapers.repec.org/RePEc:nwe:iitfed:y:2024:i:1:p:305-312>.
2. Milev, P., & Tabov, Y. (2023). From Unstructured Data to Insights: Understanding the Role of ChatGPT in the Rising Trend of AI Chatbots in Web Publications. *Godishnik na UNSS*, (1), 17-29.
3. IBM Business Process Manager Knowledge Center, IBM, from https://www.ibm.com/support/knowledgecenter/SSFPJS_8.5.6/com.ibm.wbpm.wle.editor.doc/topics/cbpmmandcasemgmt.html
4. *Definition of Enterprise Content Management (ECM)* - Gartner Information Technology Glossary. (n.d.). Gartner. <https://www.gartner.com/en/information-technology/glossary/enterprise-content-management-ec>

ТЕОРЕТИЧНИ АСПЕКТИ НА СТАТИСТИЧЕСКИЯ АНАЛИЗ НА НЕСТРУКТУРИРАНИ ДАННИ ОТ ЧАТБОТОВЕ ЗА УСЪВЪРШЕНСТВАНЕ НА КЛИЕНТСКОТО ПРЕЖИВЯВАНЕ

Theoretical aspects of statistical analysis of unstructured data from chatbots for
improving customer experience

Биляна Голешова¹

e-mail: bilyana.goleshova@unwe.bg¹

Абстракт

В днешния дигитализиран свят една от компютърните програми, която спестява много време и усилия както на потребителите, така и на компаниите е чатбота. Той улеснява обслужването на клиенти чрез автоматизация и, едновременно с това, генерира голям обем високоразмерни, неструктурирани текстови данни. Настоящият доклад има за цел да създаде теоретична и методологическа рамка за статистически анализ на разговорите в чатбот чатове, като преобразува неструктурирания диалог в статистически измерими величини за нуждите на икономическия анализ. Тъй като традиционните статистически методи не са достатъчни за извличане на икономическа оценка от такъв тип текст за подобряване на начина, по който се обслужват клиентите, се приложи подход, който се фокусира върху усъвършенствани методи за обработка на естествен език (NLP) като: токенизация, лематизация, стоп думи, векторизация, анализ на настроеността и тематично моделиране. Приложи се и статистика за извличане на изводи и заключения и прогнозни модели, за да се разбере дали промяната в чатбота наистина води до удовлетворяване на потребителите и да се прогнозира как ще бъде в бъдеще. Даде се светлина и върху това как анализа на такива данни подобрява клиентското изживяване.

Накрая се достигна до заключението, че целта на доклада е напълно изпълнена. Той представи нужният теоретичен фундамент като показва връзката между статистическия анализ на езикови данни и възможността за вземане на информирани решения, които усъвършенстват потребителското преживяване, увеличават клиентската лоялност и подобряват процесите.

Abstract

In today's digitalized world, one of the computer systems that saves a lot of time and effort for both users and the company is the chatbot. It facilitates customer service through automation and at the same time creates a large volume of high-dimensional, unstructured text data. This report aims to create a theoretical and methodological framework for statistical analysis of chatbot conversations by transforming unstructured dialogue into statistically measurable quantities for the needs of economic analysis. Since traditional statistical methods are not sufficient to extract economic value from this type of text to improve the way customers are served, an approach was applied that focuses on advanced natural language processing (NLP) methods such as: tokenization, lemmatization, stop words, vectorization, sentiment analysis, and thematic modeling. Inferential statistics and predictive models were also applied to understand whether the change in the chatbot really leads to user satisfaction and to predict how it will be in the future. It also sheds light on how analyzing such data improves the customer experience.

¹ Докторант, катедра „Статистика и иконометрия“, факултет „Приложна информатика и статистика“, Университет за национално и световно стопанство

PhD student, Department of Statistics and Econometrics, Faculty of Applied Informatics and Statistics, University of National and World Economy

Finally, it was concluded that the purpose of the report was fully fulfilled. It presented the necessary theoretical foundation by showing the connection between the statistical analysis of linguistic data and the ability to make informed decisions that enhance the user experience, increase customer loyalty and improve processes.

Ключови думи: неструктурирани данни, статистически анализ, чатбот, клиентско преживяване
JEL: C45, C55, D91, M31

Въведение

С напредването на технологиите се улесни и начинът, по който организациите в сферата на обслужване на клиенти взаимодействат със своите потребители. Обработвайки милиони разговори ежедневно, чатботовете генерират огромни количества неструктурирани текстови данни под формата на диалог. Този вид данни са безценни, но скрити. Безценни са, защото издават емоциите, желанията и болките на потребителите, благодарение на което компаниите имат възможност да разберат нуждите им и да усъвършенстват клиентското преживяване. Същевременно са известни като скрити, защото теоретичните аспекти на ефективното извличане на знание от тях остава предизвикателство. Традиционните статистически методи са създадени за работа със структурирани количествените данни и изпитват затруднения при срещата със сложния и неструктуриран естествен език. Поради това настоящият доклад има за цел да се представят теоретичните основи и методологии за статистически анализ на неструктурирани чатбот данни с фокус върху усъвършенстването на клиентското преживяване.

Теоретични основи на изследването

1. Преглед на литературата по темата

Значително нарастващата употреба на чатботове и ключовата роля на анализа на неструктурирани данни, генерирани от тези системи, за оптимизиране на клиентското преживяване, са водеща тема, която намира силен отзвук в съвременната научна литература. Доказателство за мащабността и важността на чатбота в сферата на обслужване на клиенти са докладите на Juniper Research, които посочват, че чатботовете ще спестят милиарди долари на индустрията и ще обработват над 50% от всички комуникации в електронната търговия до средата на десетилетието.

В едно от своите изследвания свързано с големите данни българският автор Dimitrova V. (2025) изразява мнение, че това, което е измеримо и класифицируемо може да се преобразува в данни и че в повечето случаи данните се събират, обработват и анализират с идеята да се подобри разбирането на определена система. Тези нейни твърдения важат с пълна сила относно извличането, обработването и анализа на неструктурирани данни от чатботове, с цел да се подобри системата, а това да доведе и до усъвършенстване на клиентското изживяване.

Настоящият преглед на литературата разглежда темата през погледа на изследователите в три основни направления:

➤ Клиентското преживяване и чатбот взаимодействия

Основополагащи са твърденията за потребителското преживяване на Verhoef, P. C. (2009) и Lemon, K. N. & Verhoef (2016). Според Verhoef, P. C. (2009) то е цялостен, взаимосвързан подход, който се формира не само от директните взаимодействия по време на покупката, но и от индиректните взаимодействия, като включва когнитивни, емоционални, поведенчески, сетивни и социални реакции на клиента. По-късно Lemon, K. N. & Verhoef, P. C. (2016) стигат до заключението, че чатботът е само една точка на контакт, която трябва да бъде безпроблемно управлявана и интегрирана в цялото пътуване на клиента, за да се постигне положително преживяване.

Основен извод на Gnewuch, U., Morana, S., & Maedche, A. (2017) е, че успехът на чатботовете в обслужването на клиенти зависи от точния баланс между ефективност и социални умения, както и от точното им позициониране спрямо типа заявка.

Chung, M., Ko, E., Joung, H., & Kim, S. J. (2020) също взимат отношение по въпроса смятайки, че дизайнът и характеристиките на чатбота оказват влияние върху удовлетвореността и намеренията за бъдещо поведение на потребителите.

➤ Обработка на естествен език в анализа на чатбот данни

Позицията на Jurafsky, D., & Martin, J. H. (2023) относно обработката на естествен език в анализа на чатбот данни е следната: Естественият език предоставя основните методи и инструменти за превръщане на неструктурираните текстови данни, генерирани от чатбот взаимодействията, в структурирана, приложима информация, като позволява на системите да разбират, интерпретират и класифицират намеренията на клиента, чувствата и темите на дискусията, с цел оптимизиране на клиентското преживяване. Това, което позволява на чатбот данните да бъдат анализирани и превърнати в стратегическа информация е, че благодарение на естествения език се разбира целта на клиента, намеренията и чувствата му.

Два аргумента на Pang, B., & Lee, L. (2008) и Liu, B. (2020) към анализа на чатбот данни се фокусират върху анализа на настроеността и фактологичното извличане на знания от текст. Основното мнение на Pang, B., & Lee, L. (2008) е, че NLP може систематично да класифицира и извлича субективна информация от текст, като определя дали изразеното мнение в чатбот разговор е положително, отрицателно или неутрално, както и каква е силата на това определеното чувство. Liu, B. (2020) констатира, че обработката на естествен език е незаменима за извличането на детайлно знание от чатбот разговорите, чрез техники като анализ на настроеността и извличане на мнения, което позволява на бизнеса да идентифицира и количествено да измери не само общото настроение на клиента, но и кои конкретни аспекти на услугите или продуктите предизвикват тези реакции.

Последната теза, която се разгледа по въпроса е на Adamopoulou, E., & Moussiades, L. (2020). Твърдят, че чрез компонентите си за разбиране и генериране естественият език е технологичният стълб, който определя типа, функционалността и сложността на всеки чатбот, като по този начин влияе върху качеството и количеството на данните, които могат да бъдат анализирани за подобряване на клиентското преживяване.

➤ Статистически и машиннообучаващи подходи за анализ на неструктурирани данни

Превръщането на свободен, словесен текст в структурирана форма позволява използването на статистически методи и алгоритми за машинно обучение

Твърдението на Creswell, J. W., & Plano Clark, V. L. (2017), отнасящо се до начинът, по който машиннообучаващите подходи трябва да бъдат комбинирани в изследване за постигане на по-ефективни резултати е следното: Най-силното и изчерпателно разбиране на смесени явления, като клиентското преживяване с чатботове, се постига чрез прилагане на смесени методи, при които количествените и качествените подходи се обединяват, анализират и интегрират в рамките на едно проучване.

Наблюдава се сходство в твърденията на разгледаната литература, което дава стабилно основание за извършване на теоретично изследване по темата в настоящия доклад. Стана ясно, че успешното функциониране на чатботовете зависи от баланса между ефективност и социална комуникация, както и от техния дизайн и въздействие върху удовлетвореността на клиента.

Единодушно е мнението, че обработката на естествен език е ключова технология, която позволява неструктурираните текстови данни от чатбот взаимодействия да бъдат трансформирани в структурирани (със стратегически приложима информация), което води до прилагането на статистически и машиннообучаващи методи. Технологията дава възможност системите да интерпретират намеренията, емоциите и темите на клиентските съобщения, което е от съществено значение за оптимизиране на клиентското преживяване. Най-пълното и ефективно разбиране на сложни явления като клиентското преживяване се постига чрез интегрирането на количествени и качествени подходи в рамките на едно цялостно изследване, като за целта могат да се използват смесените методи.

2. Теоретична рамка

Теоретичната рамка в настоящото изследване обхваща важни и основополагащи за разработката понятия, които следва да бъдат разяснени:

- Чатботът е софтуерно приложение, базирано на изкуствен интелект, което симулира човешки разговор. Той комуникира с потребителите чрез текстови или гласови методи, използвайки обработка на естествен език, с цел разбиране намеренията на потребителя и предоставяне на автоматизирани отговори или изпълняване на задачи.
- Клиентското преживяване представлява съвкупността от всички възприятия, емоции и оценки, които един клиент формира с дадена марка или компания от първоначалния им контакт до следпродажбеното обслужване.
- Неструктурираните данни са данни, които нямат предварително дефинирана структура, което ги прави трудни за съхранение, обработка и анализ.

След като се изясниха основните понятия - чатбот, клиентско преживяване и неструктурирани данни, вече е възможно да се представи начинът, по който настоящата теоретична рамка съчетава тези елементи, структурирайки анализа в три основни, взаимосвързани области, а именно: статистически анализ на неструктурирани данни, обработка на естествен език и управление на клиентското преживяване.

Теоретичната основа за статистическия анализ на неструктурирани данни се базира на текстов анализ. Текстовият анализ е процес на извличане на скрита информация от текст, което позволява последваща статистическа обработка (регресионен анализ).

Основната цел на обработката на естествен език е да даде възможност на компютрите да разбират, интерпретират и генерират човешки език по смислен начин. NLP е много важна за анализа на неструктурирани данни. В този доклад той включва предварителна обработка, количествено представяне, анализ на настроеността и тематично моделиране.

Управлението на клиентското преживяване обхваща емоционалните, когнитивните и поведенческите реакции на клиента към взаимодействието с дадена марка. Оценява се чрез метрики като удовлетвореност (CSAT), лоялност (NPS) и усилия на клиента (CES) и как те се влияят от качеството на обслужване. От поведенческите аспекти се включва това как емоциите, несигурността и очакванията на клиента, изразени в чата, влияят на бъдещото му икономическо поведение. Много е важна и концептуализацията на успешен разговор (решен проблем, удовлетворение).

Методология за статистически анализ на чатбот данни

Методологията за статистически анализ на чатбот данните е систематизиран процес, който превръща необработените потребителски диалози в полезни заключения, необходими за подобряване на потребителското изживяване в чатбота. Достигането до целта изисква преминаването през следните етапи:

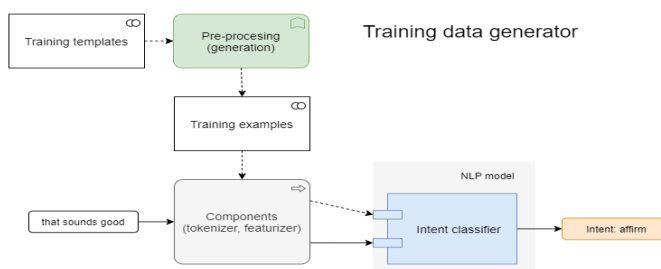
1. Етапи на подготовка на данните

- Предварителна обработка на чатбот данни

Това е първият етап при анализ на данни на разговори събирани от чатбот. Идеята е суровите чатбот разговори да се трансформират в стандартизиран масив от данни, готов за статистически анализ (напр. за изчисляване на честоти, анализ на тематиката или оценка на удовлетвореността). За целта:

- 1) Прави се токенизация. Токенизацията е разделяне на текста от диалога на по-малки смислови единици по известни като токени, т.е. целия разговор се разбива на отделни думи или фрази. Това е необходимо, защото статистическият анализ се нуждае от брой срещания на отделни думи или фрази и не може да работи с цели изречения.
- 2) Премахват се стоп думи. Стоп думите са често срещани, но неинформативни думи, които често потребителите използват в чата като: „с“, „на“, „в“, „и“ и др. Те с нищо не допринасят за смисъла на съобщението или за статистическия анализ. С тяхното премахване се предотвратява доминирането на тези думи в честотния анализ и се намалява обема на данните.
- 3) Извършва се лематизация. Това е процес на превръщане на думите в тяхната основна форма, по известна като лема. Идеята е думи с еднакъв смисъл да се броят заедно (Например говоря, говорих, говорил).

Тези етапи по предварителна обработка на данните е представен визуално с помощта на следната диаграма:



Източници: „ Intersog (2023-2024)“

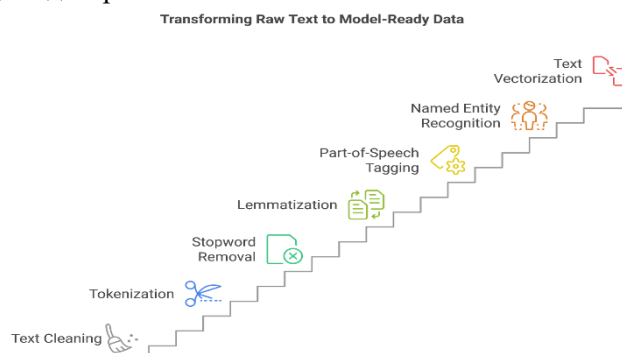
Фигура 1: „ Предварителна обработка на чатбот данни“

➤ Количествено представяне (векторизация) на чатбот данни

След предварителната обработка е необходимо думите да се превърнат в числови стойности, с цел улесняване на по-нататъшния анализ. Известни техники за това са:

- 1) Bag-of-Words изветна на български като торба с думи създава речник от всички уникални думи в данните, като брои колко пъти се среща определена дума в дадено съобщение
- 2) TF-IDF (Term Frequency-Inverse Document Frequency) на български честота на термина - обратна честота на документа освен честотата на думата в дадено съобщение, отчита и колко рядка е тази дума в целия набор от данни. С което дава по-голяма тежест на специфичните, важните думи в текст.
- 3) Word2Vec (Word to Vector) или вектор на думи улавя семантичните връзки и позволява количествени операции. Например кралица - жена, принц - мъж.
- 4) BERT Embeddings (Bidirectional Encoder Representations from Transformers), в превод двупосочни представяния на енкодера от трансформъри най-точно прихваща нюансите и контекста в сложните чатбот разговори. Една и съща дума има различен вектор в зависимост от нейното значение в конкретното изречение.

За по-ясна представа на читателя подготвителният етап е представен в обобщен вид, визуално на следващата диаграма:



Източници: „ Chaitanya Sawant, (2025),“

Фигура 2: „Етапи на подготовка на данните“

2. Описателна статистика на чатбот данни

Описателната статистика на чатбот данни е процес по събиране, обобщаване и представяне на ключови показатели от тези данни, с цел да се разбере какво се случва с потребителите. Без да представят сложни изводи и прогнози дава обобщена картина за ефективност, за поведението на потребителите и къде има нужда от подобрене.

➤ Честотен анализ на чатбот данни

Процес на изброяване, на това колко често се срещат конкретни елементи в даден набор от данни. При чатботовете честотният анализ се извършва върху почистени и токенизирани текстове, които обменят информация между потребителя и системата, както и върху метаданните (времето нужно за отговор). Идеята е да се открият често срещани думи или фрази, които да разкрият предпочитанията на потребителите, с цел да се оптимизира самия чатбот, за да се усъвършенства клиентското преживяване.

1) Честота на потребителските намерения

Това е анализ на най-често срещаните нужди на потребителя, известна метрика в статистиката като мода (най-често срещата тема). Обикновено измерва защо потребителя започва разговор, което позволява на екипа да приоритизира подобренията. Честотата на потребителски намерение се представя таблично. Таблицата е с категории и проценти на съответни разговори попадащи в отговарящата им категория.

2) Честота на ключови думи

Определя кои смислови думи се срещат най-често в потребителските чат разговори след като са премахнати стоп думите, за които се спомена по-горе в разработката. Освен това с метода TF-IDF се дава по-голяма тежест на наистина важните думи от разговора. Целта на този анализ е да идентифицират неочаквани теми или някакви пропуски в обучението на чатбот. Честотата на тези думи обикновено се представя списъчно (първите 15 думи например) или с облаци от думи.

3) Честота на пътищата на разговора

Изследва кои са най-често срещаните пътища и последователност от стъпки преди разговорът да бъде завършен или прехвърлен към човек (въпрос от клиента - отговор от потребителя). По този начин разкрива неефективни моменти или бъгове в логиката на чатбота. Резултатите се представят с помощта на Sankey диаграма или матрица. Sankey диаграмата проследява как потребителите преминават през различни категории въпроси и действия в разговорите с чатбота.

➤ Разпределение на дължината на разговора и времето за реакция

Извършва се, за да оцени ефективността на системата и качеството на потребителското изживяване. За целта в статистиката се използват централната тенденция и вариацията.

1) Разпределение на дължината на разговора

Дължината на разговора се измерва обикновено като общия брой съобщения - реплики, разменени между потребителя и робота в рамките на един чат разговор. Централната тенденция показва колко съобщения средно са необходими за приключване на един разговор, като се предпочита ниската средна стойност, защото показва бързо и директно разрешаване на проблема. Високата може да означава, че чатботът не разбира веднага въпроса. Обаче в случай, че средната стойност е много по-висока от медианната, става ясно, че има малък брой много дълги и сложни разговори, които изкривяват средното. Стандартното отклонение, от своя страна, показва колко силно се различават отделните разговори от средния брой реплики. Високото стандартно отклонение показва непоследователност, чатботът може да е много ефективен за рутинни въпроси, но много неефективен за сложни или нови теми и тази неефективност трябва да се поправи.

2) Разпределение на времето за реакция

Времето за реакция измерва закъснението между изпращането на потребителското съобщение и генерирането на отговор от чатбота. То визуализира нивото на техническа производителност и потребителско изживяване. Централната тенденция измерва времето в милисекунди или секунди, което отнема на чатбота да отговори. То трябва да бъде ниско, за да е ефективно. Чрез стандартното отклонение се разкрива нестабилност в система на чатбот. При положение, в което средното време е ниско, голямото стандартно отклонение показва, че част от потребителите получават много бавен отговор. Това може да се случи при силно натоварване.

3. Методи за анализ на съдържанието на данни от чатбота

Този тип модели имат задачата да извлекат модели, тенденции и съществено важна информация от текстовите диалози между потребителите и робота, с идеята да се подобри

клиентското преживяване. Основните методи за анализ на съдържанието на чатбот данни се основават на обработката на естествен език и следва да бъдат разгледани.

➤ Тематично моделиране на чатбот данни

Тематичното моделиране се фокусира върху намирането на латентни теми в голям набор от текстови данни от чатбот, без предварително дефинирани категории. Това помага за приоритизиране актуализациите на съдържанието и подобряване на конкретни отговори на робота.

Известни подходи за осъществяване на този метод са:

1) LDA (Latent Dirichlet Allocation) или Латентно разпределение на Дирихле

Това е алгоритъм, който приема, че всеки чат-разговор е комбинация от няколко теми, а всяка тема е комбинация от токени. И настоящият въпрос е какво всъщност прави това разпределение в чатбота? То преминава през всички думи в разговора и вкарва всяка дума към случайно избрана от него тема. След това започва да преразглежда вкараните думи, като търси по-добро разпределение, с цел да напасне така думите по темите, че да стане ясно защо се виждат точно определени думи, групирани в точно определени разговори. Например диалог в чатбот съдържаща топка, волейбол и кош е спортна, затова местенето на този разговор в раздел „спорт“ увеличава логиката на цялата система за разпределение. Идеята е така получените теми да са ясни и смислени за анализатора.

2) NMF (Non-Negative Matrix Factorization), в превод Неотрицателна матрична факторизация

Този подход разлага голямата матрица на думите на две по-малки матрици, за по-лесно интерпретиране. Получават се две матрици: едната показва важността на всяка тема в съответен разговор, а другата - важността на всяка дума в дадена тема. Важно е да се отбележи, че този алгоритъм изисква всички числа да са неотрицателни, за да може получените теми да са адитивни и лесни за разбиране.

➤ Анализ на настроенията (Sentiment Analysis)

Прави се, за да се определи емоционалната нагласа на потребителя в определен чатбот - разговор (позитивна, неутрална, негативна). Целта е бързо да се идентифицират нивата на удовлетвореност и да даде приоритет на проблемите, които роботът не е успял да разреши.

Два са подходите за sentiment анализ, а именно:

1) Лексикални подходи (Lexicon-Based)

Този метод разчита на предварително създадени речници от думи, които са разделени в две категории - позитивни и негативни. Всеки речник съдържа списък от думи и асоциации с оценки за настроение. Например „гримьор“ - 0, „опитен“ - (+2), „неопитен“ - (-2). Алгоритъмът сканира съобщението на клиента, след това събира оценките от настроенията на всяка дума и ги сумира или осреднява, с цел да изчисли общото настроение на целия разговор.

2) Подходи базирани на машинно обучение (Machine Learning)

При чатботовете тези подходи се използват за класифициране, регресия и прогнозиране. Този модел използва обучени техники (като Наивен Бейс, SVM, Deep Learning), които се прилагат за разбиране на по-сложни езикови структури. Работи като първо се обучава върху голям набор от данни от разговори, които вече са отбелязани от хора като „позитивни“, „неутрални“ и „негативни“. Той се научава не само от отделните думи, а от последователностите от думи и цялостния контекст на изречението. Така, когато моделът види ново съобщение от чатбота, прилага вече научените модели, за да предскаже с максимална точност какво е общото настроение. Важно е да се отбележи, че за разлика от лексикалните подходи, тези разбират сарказми.

4. Статистика за изводи и заключения и прогностични модели за чатбот данни

➤ Статистика за изводи и заключения

Статистиката за изводи и заключения е много важна при анализа на чатбот данни, защото позволява да се правят достоверни изводи и заключения за цялата потребителска база, основавайки се на данни от извадка (обикновено това са обикновен период, група разговори). Тя

нахвърля обобщената картина, която представя описателната статистика и дава възможност да се разбере, дали промяната в чатбота наистина води до удовлетворяване на потребителите. Необходимо е да се направи следното:

1) Тестване на ефективността

Целта е да се провери дали иновациите в дизайна на диалога, във функционалността или алгоритъма за разбиране на естествен език водят до статистическо значимо подобрене. За да се достигне до целта се прави A/B тестване, което включва хи-квадрат тест и t - тест. Чрез сравнение между старата и новата версия на чатбот група с хи-квадрат теста се взема решение, дали са успешно разрешени запитванията в чата, а с помощта на t - теста средният рейтинг на удовлетвореност и времето за разрешаване на проблема на клиента. Удовлетвореността на клиента се представя с помощта на доверителни интервали, за да се направи надеждно заключение за цялата потребителска база.

2) Оценка на обръкването

Оценява се какво въздействие има върху клиента обръкването на чатбота. Когато не разбира и прехвърля запитванията към агент. Това се прави с помощта на корелационен и регресионен анализ.

➤ Прогностични модели

Задачата на прогностичните модели за чатбот данни е да предскаже бъдещото клиентско удовлетворение или риска от неуспешен разговор, въз основа на текущи или исторически данни от диалози. Това помага на системата да вземе предварителни мерки в реално време, преди клиента да остане недоволен.

Прилагат се различни техники от машинното обучение и статистиката за предвиждане на резултата от чатбот - разговора.

1) Регресионен анализ (Regression Analysis)

Регресионния анализ помага за установяване на количествена връзка между характеристиките на разговора и степента на удовлетвореност. За анализа се използват променливи като брой съобщения преди разрешаване на проблема, време за реакция на робота, брой обръквания.

2) Анализ на тоналността (Sentiment Analysis)

Използва се за непрекъснато оценяване на емоционалното състояние на клиента като класифицира всяко съобщение като положително, неутрално, отрицателно. В случай, че се влоши тоналността, рискът от неуспешен разговор се повишава и може да се достигне до неудовлетвореност на клиента. Подходите за този тип анализ се разгледаха в подточка „Методи за анализ на съдържанието на данни от чатбота“ и няма да се обръща отново внимание върху тях.

3) Дървета на решенията и случайни гори (Decision Trees & Random Forests)

Използват се при нелинейни данни, като създават поредица от правила, с дървовидна структура, за достигането на достоверни прогнози. Например при положителен тон, време за реакция на робота от 1 секунда, брой съобщения 3 се прогнозира за висока клиентска удовлетвореност.

4) Анализ на времеви редове (Time Series Analysis)

Прилага се за прогнозиране на количествен обем от данни във времето. Анализирайки исторически времеви ред прогнозира за тенденцията на развитие, сезонност, цикличност. Така предвижда броя на разговорите, които ще пристигнат в бъдещето, което помага за оптимизиране на разпределението на ресурсите.

Приложение на анализа на чатбот данни за усъвършенстване на клиентското преживяване

1. Клиентски трудности

Тематичното моделиране и анализа на настроението имат задачата да разкрият съдържанието и емоционалния заряд на потребителското запитване, с което се идентифицират повтарящи се проблеми и области за усъвършенстване. Например ако по-голям брой от чат

разговорите съдържат тема „грешка“, „не мога да вляза“, „проблем с паролата“, с тематичното моделиране системата автоматично открива, че има проблем с достъпа и го определя като клиентска трудност. Ако пък има връщане на продукт и тематичното моделиране в комбинация със сентимент анализа отчетат висок процент отрицателен емоционален тон на потребителя, се констатира, че процесът по връщане е сложен и бавен - отново клиентска трудност.

2. Подобряване на диалоговите взаимодействия

С помощта на статистическите методи се измерва ефективността на чатбота и се идентифицират местата където той се проваля. По-високият процент на неуспех на работата да разбере намерението на потребителя, сигнализира, че модела не е достатъчно обучен да разбира естественият език. Тези провали са причина за обогатяването на тренировъчните данни на робота и добавяне на нови синоними, с цел да се увеличи точността на отговорите. Ако пък определена тема води до постоянно прехвърляне към агент, означава че логиката на робота или отговорът по тази тема не са достатъчни и е необходимо да се интегрира API (Application Programming Interface, в превод програмен интерфейс за приложения) за директно извличане на информация и прецизиране на отговора на робота.

3. Персонализирано и предсказуемо обслужване

С помощта на клъстерният анализ може да се сегментират клиентите според тяхното поведение и нужди разкрити в чат разговора. Например в групи като „Лоялни клиенти“, „Проблемни клиенти“, или пък „Информационни клиенти“. След като се идентифицира групата, роботът може да предостави персонализирани отговори още в началото на разговора. Например, ако клиентите от групата на „информационните“, често питат за промоции, робота може да изпревари нуждите им като предложи промоция преди потребителя да е попитал, с което създава положително клиентско преживяване.

Предизвикателства и бъдещи насоки при анализа на чатбот данни

1. Предизвикателства

➤ Методологични предизвикателства

Този тип предизвикателства са насочени към разбирането на естествения език и работата с големи и неструктурирани данни.

1) Обработката на сарказъм, жаргон и неясноти в контекста

На системите за анализ на настроението им е трудно да разпознават, когато положителни думи се използват в негативен смисъл. Затрудняват се и когато клиентите използват съкращения, емотикони или жаргон. За да могат да анализират такива е необходимо да бъдат обучени с реалистични данни (с включен шум). Още едно предизвикателство е разбирането на местоимения и препратки. Липсата на това разбиране води до тематични грешки.

2) Мащабируемост при обработката на големи обеми от данни

Сложните методи като дълбочинно обучение, обработка на естествен език, клъстерен анализ изискват силна разпределителна изчислителна среда, за да се приложат. Процесите трябва да са оптимизирани за скорост и автоматизираны, за да се получи анализ в реално време, който е важен за своевременна реакции при проблеми.

3) Проблем с качеството на данните

Необходими са строги процеси за предварителна обработка на данните, тъй като те са нестандартизирани и шумни (премахване на символи, дублиражи, попълване на липсващи стойности).

➤ Законодателни аспекти

Анализът на чатбот данни включва лична, чувствителна и емоционално натоварена информация за клиентите, за това трябва да се спазва общият регламент за защита на данните на ЕС. Така няма да се нарушава неприкосновеността и личния живот на потребителите.

2. Бъдещи насоки

➤ Необходимо е да се интегрира генеративен изкуствен интелект (Големи езикови модели) в анализа на чатбот разговори и усъвършенстването на клиентското преживяване.

Големите езикови модели могат да извършват много по-точен и ясен анализ на настроението, както и тематично моделиране от обичайните алгоритми. Те могат да обобщават дълги разговори, които след това да бъдат статистически анализирани. Освен това имат възможности

да констатираят набор от провали на чатбота и да предложат по добри и точни контекстуални отговори за обучението на робота. Големите езикови модели умеят да идентифицират пропуски, анализирайки хиляди негативни разговора и да предложат промени в езика и правилата.

- В бъдеще чатбот данните да се комбинират със структурирани данни от други източници (история на покупките, посетени страници на уебсайта, честота на използване на приложението).

Това позволява да се направи клъстерен анализ като се идентифицират проблемите и корелациите между поведението на потребителите онлайн и нуждата от поддръжка на системата.

Заклучение

Настоящият доклад обстойно изследва теоретичните аспекти на статистическият анализ на неструктурирани данни от чатбот разговори, с идеята да се направи подобрене на клиентското преживяване и да се избегне достигането до фрустрация на потребителя.

При прегледа на литературата на изследването се разгледа темата през погледа на различни изследователи и се констатира тяхното мнение по въпроса. В теоретичната рамка се обясни значението на базовите понятия по темата (неструктурирани данни, клиентско преживяване и чатбот) и се хвърли светлина върху нейния обхват, след което се изложиха основните теоретични подходи използвани в анализа. Обърна се внимание на предварителната обработка на данните (токенизация, стоп думи, лематизация), която е важна за превръщането на неструктурирания текст във вид подходящ за анализ. Към подготовката, за превръщането на думите в числови стойности се включи и векторизацията, която обхваща методи като: Bag-of-Words, TF-IDF, BERT Embeddings и Word2Vec. Впоследствие се представиха техники за честотен анализ на чатбот данни и се хвърли светлина върху методите за анализ на съдържанието на такива данните. Както стана ясно, методите за анализ на съдържанието на чатбот данни са изключително важни, за извличането на модели, тенденции и съществено важна информация от диалозите между потребителите и робота, с идеята да се подобри клиентското преживяване. Този анализ се фокусира върху техники като тематично моделиране, благодарение на което се намират латентни теми в голям набор от текстови данни от чатбот, без предварително дефинирани категории и анализ на настроението, с който се определя емоционалната нагласа на потребителя и се идентифицират нивата на удовлетвореност. На последно място се разгледаха статистически методи за анализ на изводи и заключения и методи за извършването на прогнози. Статистическите методи за изводи и заключения са от особено значение за анализа, защото дават възможност да се дефинира краен и достоверен извод за това, дали промяната в чатбота наистина води до удовлетворяване на потребителите и въз основа на това се правят бъдещи прогнози. Накрая в една точна се обясни как всички тези техники за анализ на данни от чатбот допринасят за подобряването на клиентското преживяване.

След като се представи нужният теоретичен фундамент, докладът показва връзката между статистическият анализ на езикови данни и възможността за вземане на информирани решения, които подобряват потребителското преживяване, увеличават клиентската лоялност и подобряват процесите. От това следва, че поставената цел, да се представят теоретичните основи и методологии за статистически анализ на неструктурирани чатбот данни с фокус върху усъвършенстването на клиентското преживяване, е напълно постигната.

Основен принос на доклада към теоретичната база е, че представя цялостен теоретико-методологичен модел, който систематизирано свързва статистическите методи и обработката на естествен език с управлението на клиентското изживяване.

Източници

1. Adamopoulou, E., & Moussiades, L. (2020). An Overview of Chatbot Technology, AI & Society, 35(3), pp. 535 - 547.

2. Allahyari, M., Pouriyeh, S., Assefi, M., Safaei, S., Trippe, E. D., Gutierrez, J. B., & Kochut, K. (2017). A brief survey of text mining: Classification, clustering and extraction techniques
3. Blei, D. M., Ng, A. Y., & Jordan, M. I. (2003). Latent Dirichlet Allocation, *Journal of Machine Learning Research*, 3, pp. 993 - 1022.
4. Chung, M., Ko, E., Joung, H., & Kim, S. J. (2020). Chatbot e-service and customer satisfaction, *Journal of Business Research*, 117, pp. 587 - 595.
5. Creswell, J. W., & Plano Clark, V. L. (2017). *Designing and Conducting Mixed Methods Research* (3rd ed.), SAGE publications, Chapter 1 - 4.
6. Dimitrova, V., (2025). Modeling Big Data Using Python, *Challenges to the Statistical Science and the Business Analysis Profession in the Era of Big Data and Artificial Intelligence: Conference Proceedings*, Sofia, pp. 30 - 40. ISBN 978-619-232-923-5.
7. Gatzioufa, P., & Saprikis, V. (2022). A literature review on users' behavioral intention toward chatbots' adoption, *Applied Computing and Informatics*, 18(1), 1 - 17.
8. Gnewuch, U., Morana, S., & Maedche, A. (2017). Towards Designing Cooperative and Social Conversational Agents for Customer Service. In *Proceedings of the 38th International Conference on Information Systems (ICIS 2017)*, Seoul, South Korea.
9. <https://intersog.com/blog/strategy/three-methods-of-pre-processing-data-in-chatbot-development/>
10. Jurafsky, D., & Martin, J. H. (2023). Chatbots and dialogue systems. In *Speech and Language Processing* (3rd ed.), Draft, Chapter 15, pp. 309 - 330
11. Kaczorowska-Spychalska, D. (2019). How chatbots influence marketing. *Management*, 23(1), pp. 251 - 270.
12. Lee, D. D., & Seung, H. S. (1999). Learning the parts of objects by non-negative matrix factorization, *Nature*, 401(6755), pp. 788 - 791.
13. Lemon, K. N., & Verhoef, P. C. (2016). Understanding customer experience throughout the customer journey, *Journal of Marketing*, 80(6), pp. 69 - 96.
14. Liu, B. (2012). *Sentiment analysis and opinion mining*, Morgan & Claypool Publishers, ISBN 9781608458851
15. Liu, B. (2020). *Sentiment Analysis: Mining Opinions, Sentiments, and Emotions* (2nd ed.), Cambridge University Press, Chapter 1 - 4, pp. 1 - 113.
16. Pang, B., & Lee, L. (2008). Opinion mining and sentiment analysis. *Foundations and Trends in Information Retrieval*, 2(1 - 2), pp. 1 - 135.
17. Sawant, C. (2025, January 3). From messy text to model-ready data: A guide to NLP preprocessing, Medium.
18. Schöbel, S., Schmitt, A., Benner, D., & Janson, A. (2021). Future directions for chatbot research: An interdisciplinary research agenda, *Computing*, 103(10), 2915 - 2942
19. Verhoef, P. C., et al. (2009). Customer experience creation: Determinants, dynamics and management strategies, *Journal of Retailing*, 85(1), 31 - 41.

SEMANTIC SEARCH BASED ON EMBEDDING MODELS

Vanya Lazarova¹

e-mail: vlazarova@unwe.bg

Abstract

In this paper, very briefly the embedding models and the capabilities they provide for semantic search are introduced. In the paper is also presented the workflow of semantic search and the Semantic search in SQL database that are extended with data type VECTOR and the operations with vectors. The benefit for the end user is that he can find relevant documents in his organization's database by comparing query embeddings and documents.

Key words: LLMs, vector search, vector databases, embedding vectors

JEL: C30

1. Embedding models

The embedding models are kinds of large language models (LLMs) that generate for a given text (it could be a single sentence, but also could be a business document, journal paper, chapter of book, and so on) a vector that represents the coordinates of this text in the multidimensional space of the model.

Models are named „language” or „linguistic“, but some LLMs can represent images, audio and video objects, not only text.

We will consider further exclusively models that represent text.

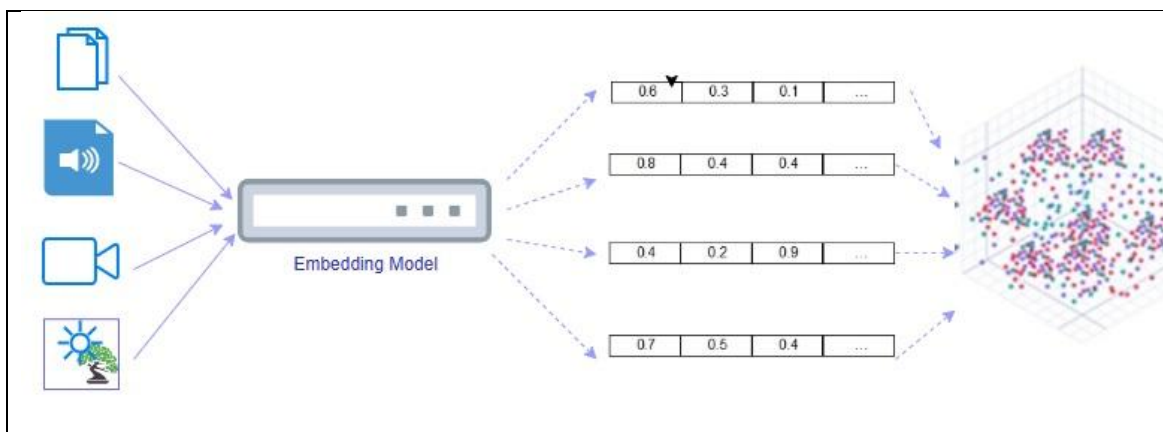


Fig. 1. Embedding models can represent different objects – text, audio, video, images

The embedding vectors capture semantic meaning and context through numerical representations of data. Data with similar semantic meaning have embeddings that are closer together. Embeddings enable a wide range of applications, including:

- 1) **Search with queries on natural language:** Find relevant documents within large databases, like legal document retrieval or enterprise search, by comparing the embeddings of queries and documents.
- 2) **Retrieval-Augmented Generation (RAG):** Enhance the quality and relevance of generated text by retrieving and incorporating contextually relevant information into the context of a model.
- 3) **Clustering and Categorization:** Group similar texts together, identifying trends and topics

¹ Доц. д-р към катедра ИТК, УНСС, email: vlazarova@unwe.bg

within your data.

4) **Classification:** Automatically categorize text based on its content, such as sentiment analysis or spam detection.

5) **Text Similarity:** Identify duplicate content, enabling tasks like web page deduplication or plagiarism detection.

In the paper we will consider the usage of embedding vectors, which can be applied in almost any information systems of small and middle companies - search in databases with queries on natural language.

The leaderboard of the site **huggingface.co** <https://huggingface.co/spaces/mteb/leaderboard> (8) compares 100+ text embedding models across different languages and task types.

In this paper we recommend the multilanguage model of Google, **gemini-embedding-exp-03-07**. It generates embedding vectors with the following characteristics:

- **Dimensionality:** Each embedding vector consists of **3,072 dimensions**. It allows truncating the original 3,072-dimensional vectors to a smaller size, enabling a balance between accuracy and storage efficiency.
- **Input Token Limit:** The model supports input texts up to **8,192 tokens** in length.
- **Multilingual Support:** The model supports over 100 languages, making it versatile for various linguistic applications.

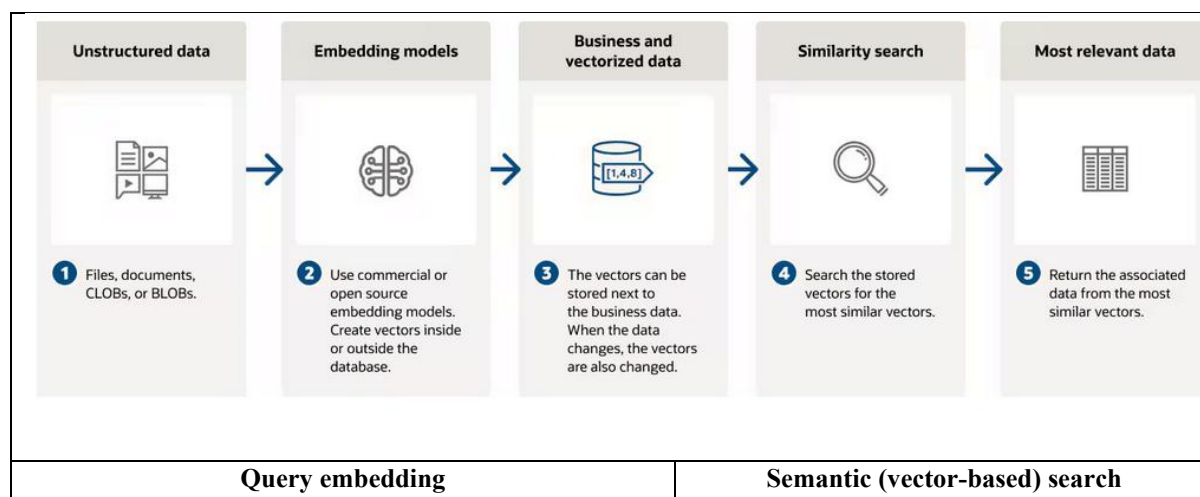
This model is **available free** through the Gemini API. Details of this embedding model are described on the Google webpage:

<https://developers.googleblog.com/en/gemini-embedding-text-model-now-available-gemini-api/>.

2. Semantic search based on embedding models

2.1. Workflow of semantic search

Semantic search is more meaningful than simple keyword search, as it seeks to understand the deeper meaning and intent behind the query. LLMs make it possible to go beyond keyword matching. The information about the relationships between words, synonyms, and entities, stored in LLM, allows to interpret an user's request articulated in natural language.



Source: <https://www.oracle.com/my/database/ai-vector-search/> (7)

Fig.2. The workflow of vector search

The process of vector-based semantic search in database has two phases:

Query embedding: The LLM converts the user’s query into an *embedding vector*, that is a numerical representation capturing its semantic meaning. If the document is large (e.g. a book), it’s typically broken into smaller chunks (e.g. chapters), each mapped as a unique point in a high-dimensional vector space.

Semantic search: The embedding vector of the query is compared to the *embedding* vectors of the documents. The documents with the highest similarity score are provided to the user as matching results.

2.2. Semantic search in SQL database

SQL databases supporting VECTOR type and VECTOR INDEX

Since 2023 most SQL servers are extended with data type VECTOR and the operations with vectors. The vector functions, vector data type and vector indexes, allow developers to perform efficient searches and analyses on large volumes of vector data.

Embedding models generate typical vectors with 1 536 or 3 072 dimensions.

The embedding model **gemini-embedding-exp-03-07** generates vectors with **3 072** dimensions, but could be truncated to **1 536**

Two SQL databases supporting VECTOR type and VECTOR INDEX

	data type VECTOR(n) n - dimensions of embedding vectors	VECTOR INDEX algorithm
Microsoft SQL Server 2025	VECTOR(n) n: max 32,768 dimensions Optimized for vectors to 2 048 vector elements are of FLOAT32 Each row of Embedding column contains an area of float32	HNSW
Oracle database 23ai	VECTOR(n, <i>element_type</i>) n: max 32 768 dimensions Optimized for vectors to 4 096 <i>element_type</i> : FLOAT32, FLOAT64 Each row of Embedding column contains a VARBINARY	HNSW

CREATE table with a column of VECTOR type and a VECTOR INDEX for this column

Microsoft SQL 2025	CREATE TABLE Documents (Id INT PRIMARY KEY, Content NVARCHAR(MAX), Embedding VECTOR(1536)); CREATE VECTOR INDEX idx_Documents_Embedding ON Documents (Embedding) WITH (DISTANCE_FUNCTION = 'COSINE' -- or 'EUCLIDEAN');
Oracle Database 23ai	CREATE TABLE Documents (Id NUMBER PRIMARY KEY, Documtent CLOB, Embedding VECTOR(1536)); CREATE VECTOR INDEX idx_Documents_Embedding ON Documents (Embedding)

	WITH PARAMETERS (dimensions = 1536, metric = 'COSINE' -- or 'L2');
--	---

The vector search in SQL is performed by operator (v1) <-> (v2) or function VECTOR_DISTANCE(v1,v2, search_mode).

These function and operator are not standard in SQL. The considered SQL databases differ in the way the choice between exact and approximate search is pointed.

	operator (v1) <-> (v2)	function VECTOR_DISTANCE(v1,v2, search_mode)
• Microsoft SQL 2025	(v1) <-> (v2) search mode depends on the VECTOR INDEX	search_mode: COSINE for cosine distance EUCLIDEAN for Euclidean (L2) distance
• Oracle Database 23ai	(v1) <-> (v2) search mode depends on the VECTOR INDEX	search_mode: COSINE for cosine distance L2 for Euclidean (L2) distance

In both SQL servers, the **approximate search** requires two conditions:

- 1) VECTOR INDEX must be created
- 2) In the SELECT statement, a restriction on the number of the selected rows must be defined:
 - TOP n -- in Microsoft SQL 2025
 - FIRST n -- in Oracle Database 23ai

The choice between exact and approximate search is different in these two SQL servers using books database.

• **SELECT in Microsoft SQL 2025**

-- Exact search: function VECTOR_DISTANCE() must be used SELECT TOP 5 BookId, Title, VECTOR_DISTANCE (Book_embedding_vector, @query_vector, 'COSINE') as Distance FROM books ORDER BY Distance;
-- Approximate search: operator <-> must be used. SELECT TOP 5 BookId, Title, FROM books ORDER BY Book_embedding_vector <-> @query_vector;

• **SELECT in Oracle SQL Database 23ai**

-- Exact search: In FETCH clause the key word EXACT must be used. SELECT BookId, Title, VECTOR_DISTANCE (Book_embedding_vector, :query_vector, COSINE) as Distance FROM books ORDER BY Distance FETCH EXACT FIRST 5 ;
-- Approximate search: In FETCH clause the key word APPROX must be used. SELECT BookId, Title, VECTOR_DISTANCE (Book_embedding_vector, :query_vector, COSINE) as Distance FROM books ORDER BY Distance FETCH APPROX FIRST 5 ;

3. Conclusion

In the text above, we have indicated a wide range of applications that can use Embeddings. We only looked at a small part – search with queries on natural language and find relevant documents within large databases.

There are some concepts used in the code - Euclidean vector distance and Cosine vector distance - that are not explained in details, but this is beyond the scope of this report. More detailed information on the matter can be obtained from the article „Artificial Intelligence as a Service“ (5).

We hope this will be useful for practical developments in the field of artificial intelligence as a service.

References

1. Hristov, G. (9 2024 r.). Improving the Quality of Financial Information Through Machine Learning. Economic Alternatives, pp. 529-540.
2. Belev, I. (2023). Object-Oriented Ecm Approach, Innovative Information Technologies for Economy Digitalization (IITED), University of National and World Economy, Sofia, Bulgaria, issue 1, pages 7-14.
3. Mihova, V. (2023). Expanding The Possibilities Of Information System For Working With Voice, Innovative Information Technologies for Economy Digitalization (IITED), University of National and World Economy, Sofia, Bulgaria, pages 152-157.
4. Velkova, I. (2023). Reducing Advertising Fraud In Digital Marketing With Artificial Intelligence. Innovative Information Technologies for Economy Digitalization (IITED) (pages 79-88). Sofia: University of National and World Economy, Bulgaria.
5. Lazarova, V. (2025). Artificial Intelligence as a Service. Ikonomiceski i Sotsialni Alternativi, (2), 128-141.
6. Stefanov, Geno (2024). Integrating Amazon Web Services (AWS) and Hadoop for Big Data processing, Innovative Information Technologies for Economy Digitalization (IITED), University of National and World Economy, Sofia, Bulgaria, pages 154-161.
7. oracle.com. Oracle AI Vector Search. Retrieved from <https://www.oracle.com/asean/database/ai-vector-search/>
8. huggingface.co. Embedding Leaderboard. Retrieved from <https://huggingface.co/spaces/mteb/leaderboard>

СЪВРЕМЕННИ АНАЛИТИЧНИ ПОДХОДИ ЗА ОЦЕНКА НА РЕЗУЛТАТИТЕ ОТ ДИГИТАЛНИ МАРКЕТИНГОВИ КАМПАНИИ

Modern Analytical Approaches for Evaluating the Results of Digital Marketing Campaigns

Гергана Тодорова¹,
e-mail: gtodorova811@gmail.com¹

Абстракт

Разходите за маркетинг и реклама на съвременните компании представляват все по-голям процент от общите оперативни средства. Динамично развиващата се дигитална среда налага взимането на бързи и верни решения по отношение на ефективност, бюджет и възвращаемост, което превръща изкуствения интелект в обещаващ инструмент за управление на разходите. Настоящият доклад представя съвременни подходи за оценка на резултатите от маркетинговите дейности, включително измерване на възвращаемостта. Интелигентните системи дават възможност за задълбочен анализ с цел правилно разпределение на инвестираните средства за оперативни и маркетингови дейности в бюджета на компанията и позволяват ясно да се определи тежестта на разходите, от които може да се очаква възвращаемост, в контекста на бюджетиране и финансов резултат. Използването на точните инструменти подпомага взимането на информирани управленски решения и може да подобри финансовите резултати на компанията.

Abstract

Marketing and advertising expenses represent an increasingly significant share of companies' overall operational budgets. The dynamic digital environment requires quick and accurate decisions regarding efficiency, budgeting and return on investment, making artificial intelligence a promising tool for cost management. This paper presents modern approaches to evaluating the results of marketing activities, including measuring return of investment. Intelligent systems enable in-depth analysis aimed at proper allocation of marketing and operational funds within corporate budgets and help to clearly determine which expenses are expected to generate returns in the context of budgeting and financial performance. The use of appropriate analytical tools supports informed managerial decision-making and can significantly improve a company's financial results.

Ключови думи: дигитален маркетинг, изкуствен интелект, анализ на ефективност, възвращаемост, бюджетиране.

JEL: M31, C55, O33

Увод

Разходите за маркетинг и реклама на съвременните компании представляват все по-значителен дял от общите оперативни средства, а ефективното им управление изисква адаптивност и бърза реакция (Radev & Naydenov, 2021). Дигиталната трансформация променя начина, по който организациите взаимодействат с потребителите, както и начина, по който измерват ефективността на своите маркетингови дейности. В условията на бърз технологичен напредък и силна конкуренция, необходимостта от прецизна оценка на резултатите от маркетинговите кампании е по-голяма от всякога (Al Khaldy, Al-Obaydi, & al Shari, 2023).

¹ Докторант към катедра Индустриален бизнес, Бизнес факултет, УНСС, ORCID: 0009-0007-5873-4217

В съвременния дигитален свят маркетинговите кампании преминават от креативна дейност към строго аналитичен процес, където именно данните и тяхната интерпретация играят решаваща роля. Организациите, които желаят да достигнат до релевантна аудитория, да я задържат и да извлекат максимална стойност от нея, трябва да прилагат съвременни аналитични подходи, които да измерват, прогнозират и оптимизират всяка стъпка на маркетинговия план. В този контекст интеграцията на изкуствен интелект (AI), машинно обучение и BI платформи превръщат маркетинга в дисциплина, базирана на данни, а не само на интуиция (Kaufmann & Reinecke, 2025).

Използването на AI системи дава възможност за анализ на огромни обеми от разнородни данни — от уеб анализи, CRM системи, социални мрежи и мобилни приложения — с цел автоматично откриване на шаблони на потребителско поведение и препоръчване на оптимални действия в маркетинговата среда. Така компаниите по-бързо идентифицират неефективни кампании, пренасочват разходите си и повишават възвръщаемостта на инвестициите (ROI) чрез целенасочена оптимизация (Al Khaldy, Al-Obaydi, & al Shari, 2023).

Теоретични основи на оценката на маркетинговите резултати

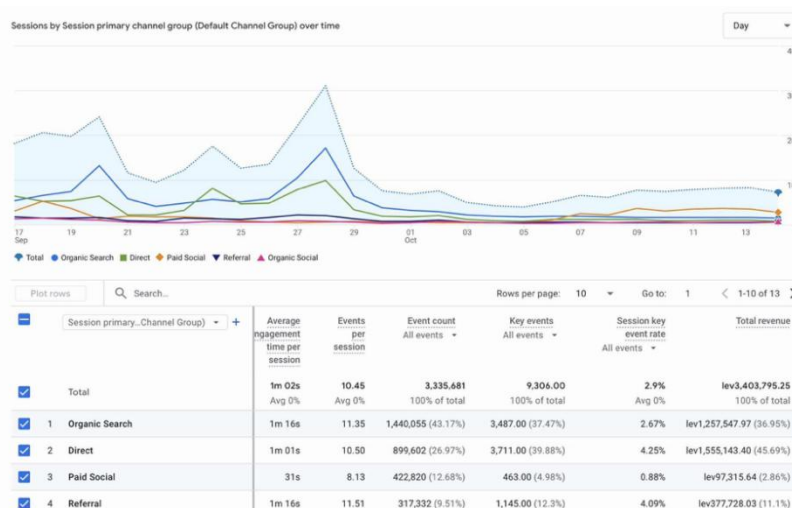
Оценката на ефективността на маркетинговите дейности е сред най-сложните управленски задачи. Класическите показатели като ROI (Return on Investment), CPA (Cost per Acquisition), CPL (Cost per Lead) и CLV (Customer Lifetime Value) служат като основни инструменти за измерване на ефективността, но често не отразяват пълната картина в динамична дигитална среда (Saura, Palos-Sánchez, & Cerdá Suárez, 2017). В дигиталния маркетинг се появяват нови предизвикателства – многоканалност, фрагментирано потребителско поведение, зависимост от алгоритми на социалните мрежи и невъзможност за проследяване на потребителите в реално време поради ограничения в бисквитките и регулациите за защита на личните данни (Pantelić, 2022; Sholihah et al., 2024). Това налага комбинирането на класически финансови показатели с данни от уеб, социални платформи, CRM системи и вътрешни финансови отчети. Чрез подобен подход може да бъде направена по-точна оценка на влиянието на маркетинговите кампании върху реалните бизнес резултати (Saura, Palos-Sánchez, & Cerdá Suárez, 2017).

Съвременни подходи за анализ на ефективността

С развитието на технологиите, компаниите разполагат с големи обеми данни, които предоставят нови възможности за анализ. Сред основните съвременни подходи се открояват:

- Анализ на придобиване на потребители

Придобиването на потребители представлява основополагащ елемент от стратегията на дигиталния маркетинг, тъй като определя основата на бъдещата ангажираност и дългосрочната стойност на клиентите (Chougala, 2024). Този анализ се фокусира върху идентифициране на ефективните канали за привличане на нови потребители, оценка на тяхната стойност и оптимизация на маркетинговите усилия. Основната цел на анализа на придобиването е да се определи кои канали и кампании водят до най-голям брой качествени потребители (López García et al., 2019). Това включва измерване на ключови показатели като Cost Per Acquisition (CPA), Customer Acquisition Cost (CAC) и Conversion Rate (CR), които дават информация за ефективността на всяка маркетингова инициатива. Инструментите за анализ на придобиването на потребители, вградени в платформи като Google Analytics, Meta Analytics и др., предоставят възможност за автоматично събиране, обработка и визуализация на данни от различни канали. Чрез тях компаниите могат да следят в реално време ефективността на кампаниите, да идентифицират най-успешните стратегии и да предприемат корективни действия, когато резултатите не отговарят на очакванията.



Фигура 1- Анализ за придобиване на потребители в Google Analytics

- **Прогнозен анализ**, който използва исторически данни и алгоритми за машинно обучение, за да предвиди бъдещи резултати от кампаниите.

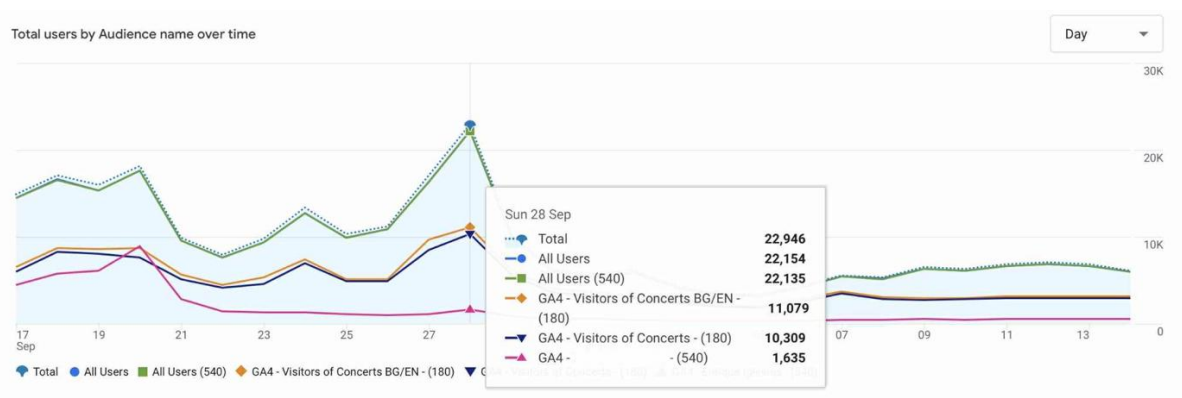
Прогнозният анализ позволява на организациите да предвиждат бъдещи резултати и поведение на потребителите въз основа на исторически данни. В условията на динамична дигитална среда, където потребителите преминават през множество точки на контакт и канали, прогнозният анализ помага да се идентифицират тенденции, които не са видими при стандартните ретроспективни отчети. Същността на прогнозния анализ се крие в използването на математически модели и алгоритми за машинно обучение, които обработват големи обеми разнородни данни. Например, чрез прогнозен анализ маркетинговият екип може да предвиди кои потребители имат висок риск от отлив през следващите месеци и да предприеме превантивни действия – напр. насочени кампании за задържане. Прогнозният анализ се използва и за оптимизация на бюджетите за маркетинг. Компаниите следва да адаптират своите модели на приходи, за да отговорят на условията на средата, което изисква гъвкавост и динамична оценка на маркетинговите усилия (Radev & Deyanov, 2023). Чрез моделиране на бъдещи резултати компаниите могат да определят кои канали и кампании ще донесат най-висока възвръщаемост на инвестициите (ROI) и да пренасочат ресурсите от по-малко ефективни активности. Освен това алгоритмите за прогнозиране позволяват идентифициране на нови сегменти потребители, които имат потенциал за висока стойност, но не са били открити чрез традиционни аналитични методи. Това повишава ефективността на маркетинговите стратегии и подкрепя вземането на стратегически решения (Al Khaldy, Al-Obaydi, & Al Shari, 2023).



Фигура 2 - Прогноза за брой потребители, съпоставена в реален брой потребители на уеб сайт // Google Analytics

- **Кохортен анализ**, който групира потребителите според поведението им и проследява възвръщаемостта от конкретни кампании спрямо отделните сегменти.

Чрез кохортен анализ се оценява поведението на потребителите във времето, базирано на модели на задържане и ангажираност. Този подход се основава на групиране на потребители според определени характеристики или период на придобиване, след което се проследява тяхното поведение в различни времеви интервали. Основното предимство на кохортния анализ е възможността да се идентифицират специфични тенденции, които биха останали незабелязани в общия низ с данни. Той позволява да се разграничат различни модели на поведение и да се изследва как различни кампании и маркетингови стратегии влияят на ангажираността и задържането на отделни групи потребители. Например, кохортен анализ може да покаже, че потребителите, придобити чрез платена реклама в социални мрежи, демонстрират висока активност през първите три месеца, но значително намаляват ангажираността си след шестия месец, докато потребителите, привлечени чрез органичен трафик, запазват стабилна активност през по-дълъг период. Кохортният анализ също така е ценен инструмент за изчисляване на дългосрочната стойност на клиента (Customer Lifetime Value – CLV). Чрез проследяване на отделните кохорти, компаниите могат да определят кои канали за придобиване водят до най-ценните клиенти и кои стратегии изискват допълнителна оптимизация. Тази информация позволява фокусиране на ресурсите върху групите с най-голям потенциал за генериране на доход и задържане, което води до повишаване на възвръщаемостта на маркетинговите инвестиции (López García, Lizcano, Ramos & Matos, 2019).



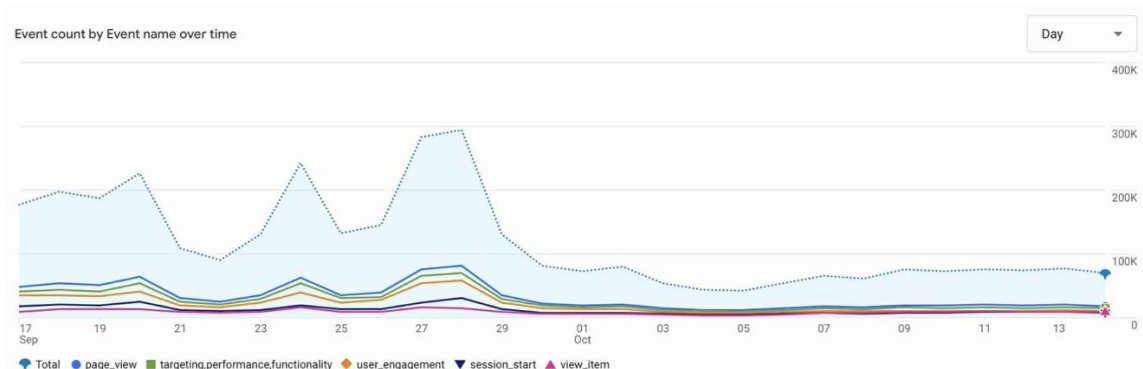
Фигура 3 - Кохортен анализ на групи потребители, които бизнесът ръчно е създавал в Google Analytics

- **Анализ на ангажираността**, който измерва взаимодействието с брандовото съдържание в социални мрежи, блогове и уебсайтове.

Анализът на ангажираността измерва до каква степен потребителите взаимодействат с продукта, услугата или марката. Той включва множество показатели, които дават представа за качеството на връзката между потребителя и организацията, включително време, прекарано на уебсайта или в приложението, честота на взаимодействията, брой повторни посещения, кликания върху съдържание и взаимодействия в социалните мрежи. Анализът на тези данни позволява да се оцени не само количеството трафик, но и неговото качество, което е решаващо за определяне на стойността на клиента (Walker, 2021). Например, потребители, които редовно отварят имейл кампании и взаимодействат със съдържанието в социалните мрежи, могат да бъдат класифицирани като високоефективни клиенти, докато потребителите, които посещават сайта еднократно, но не предприемат действия, изискват различна стратегия за ангажиране.

Анализът на ангажираността не се ограничава до еднократни измервания, а често се комбинира с кохортен и прогнозен анализ, за да се проследи динамиката на поведението на потребителите

във времето и да се предскаже тяхната бъдеща активност. Освен това, различни модули с изкуствен интелект или машинно обучение в платформите допринасят за автоматизацията на този процес, като идентифицират модели на поведение, които предполагат висок риск от отлив или, обратно, висок потенциал за повторни действия. Въз основа на тези данни могат да се предложат целеви оферти, персонализирано съдържание и оптимизирани кампании, които максимизират стойността на потребителя (Chougala, 2024).



Фигура 4 - Анализ на ангажираността при предварително зададени "събития" в уеб сайт // Google Analytics

Изкуствен интелект и машинно обучение в оценката на маркетинговите резултати

Изкуственият интелект се превръща във важен инструмент в процеса на анализ на маркетинговите резултати. Чрез **машинно обучение, обработка на естествен език и ИИ**, системите могат автоматично да разпознават закономерности в потребителското поведение, да откриват неефективни разходи и да прогнозират потенциалната възвръщаемост от кампании (Al Khaldy, Al-Obaydi & Al Shari, 2023; Haleem, Javaid, Singh & Suman, 2022).

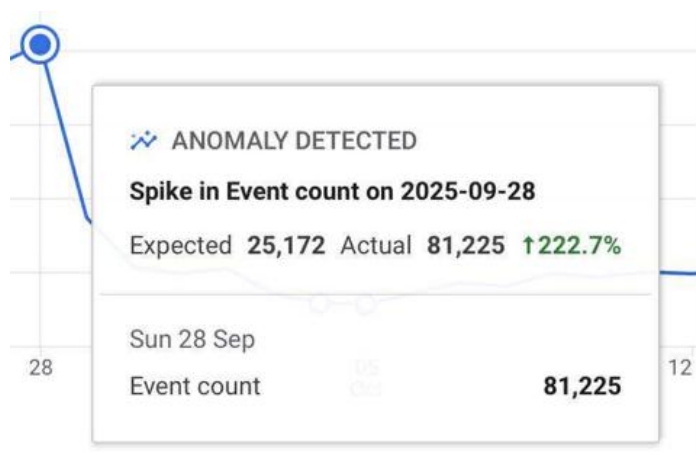
Приложенията на AI включват:

- автоматично анализиране на кампанийни данни и идентифициране на успешни и неуспешни канали;
- модели за прогнозиране на потребителски интерес и вероятност за конверсия;
- откриване на аномалии (например фалшив трафик или неефективни бюджети);
- автоматично препоръчване на оптимално разпределение на разходите между различни канали.

Тези решения спомагат за **намаляване на човешкия фактор** и **ускоряване на вземането на решения**, като осигуряват надеждна база за бюджетиране и стратегическо планиране (Al Khaldy, Al-Obaydi & Al Shari, 2023).

BI системи и интегрирани аналитични решения

Business Intelligence (BI) системите продължават да бъдат фундаментална част от процеса на оценка на маркетинговите резултати. Те позволяват визуализация и обобщаване на данни от различни източници чрез табла (dashboards), графики и автоматизирани отчети. Съвременните бизнес интелигентни системи вече започват да „говорят“ на езика на изкуствения интелект – AI анализира и прогнозира, докато BI визуализира и комуникира резултатите. Така BI инструментите предоставят на мениджърите **интуитивен и обективен поглед върху ефективността**, което подпомага навременното пренасочване на ресурси към по-успешни инициативи.



Фигура 5 - Засичане на аномалия в Google Analytics

Изводи и насоки за бъдещи изследвания

Измерването на ефективността на маркетинговите дейности отдавна не може да се основава единствено на традиционни финансови показатели. Изкуственият интелект и интегрирането му в нови или подобрени технологични решения разкриват огромен брой възможности за интегрирано и динамично оценяване на маркетинговите резултати. Въпреки че цялата сфера е обхваната от завиден оптимизъм в използването на изкуствен интелект, предлаган на пазара от големите компании, пред маркетинг анализаторите стоят сериозни предизвикателства. Следва да се вземат предвид всички възможни грешки, които биха могли да останат незабелязани и да бъдат погрешно интерпретирани, следователно – да доведат до съществени вреди за бизнес организациите. Макар интегрираните бизнес решения в сферата на маркетинга да стават все по-надеждни, човешкият фактор е тук, за да остане още известно време. В днешно време, достигането на определени изводи, което би отнело X време и Y брой служители, се случва мигновено с помощта на „дигитален помощник“. Усилията на всички доставчици на маркетинг услуги и реклама в онлайн пространството са насочени именно към оптимизация и намаляване на времето за анализ, като по този начин компанията може да бъде по-склонна да инвестира повече средства в същинска реклама, вместо в обучаване и наемане на още повече хора, които да следят за нейното адекватно изпълнение. Взимайки предвид изложеното в доклада и в предходните редове от заключението, можем с увереност да твърдим, че комбинацията от AI, BI и методологичен подход, ориентиран към данни, може да подобри стратегическото управление на маркетинга и води до оптимизация (или реструктуриране) на разходите. Бъдещите изследвания следва да се насочат както към разработване на автоматизирани системи за оценка на възвръщаемостта, които обединяват маркетингови и финансови данни, така и към по-прозрачно и ефективно управление на потоците с информация, които водят до взимането на подобни решения.

Литература

1. Chougala, M. K. (2024). The role of digital marketing in driving customer acquisition, personalization, and brand trust in fintech services. *Journal of Computational Analysis and Applications*, 33(6), 1018–1023. <https://doi.org/10.28945/6467>
2. Haleem, A., Javaid, M., Singh, R. P., & Suman, R. (2022). Artificial intelligence (AI) applications for marketing. *Journal of Marketing Analytics*, 10(1), 1–15. <https://doi.org/10.1057/s41270-022-00115-8>
3. Kaufmann, P., & Reinecke, S. (2025). AI in marketing analytics – 5 levers for data driven business growth. <http://hdl.handle.net/20.500.14171/123423>
4. López García, J. J., Lizcano, D., Ramos, C. M. Q., & Matos, N. (2019). Digital marketing actions that achieve a better attraction and loyalty of users: An analytical

- study. *Journal of Theoretical and Applied Electronic Commerce Research*, 11(6), 130–d238353. <https://doi.org/10.4067/S0718-18762019000600130>
5. Pantelić, O. (2022). Cookies implementation analysis and the impact on user privacy. *Sustainability*, 14(9), 5015. <https://doi.org/10.3390/su14095015>
 6. Radev, R., & Naydenov, A. (2021). Ad hoc approach of cost management during the time of COVIDization of the economy. *SHS Web of Conferences*, 120, 02007. https://www.researchgate.net/publication/354093639_AD_HOC_approach_of_cost_management_during_the_time_of_COVIDization_of_the_economy
 7. Radev, R., & Deyanov, D. (2023). Ad hoc approach of cost management during the time of COVIDization of the economy. *SHS Web of Conferences*, 120, 02007. https://www.researchgate.net/publication/354093639_AD_HOC_approach_of_cost_management_during_the_time_of_COVIDization_of_the_economy
 8. Rogers, D. L. (2016). *Digital transformation playbook: Rethink your business in digital age*. Columbia University Press.
 9. Saura, J. R., Palos Sanchez, P., & Rios Martin, M. A. (2017). Business intelligence and analytics for improving the performance of marketing decisions. *Journal of Business Research*, 70, 453–457. <https://doi.org/10.1016/j.jbusres.2016.08.004>
 10. Vollrath, M. D., & Villegas, S. G. (2021). Avoiding digital marketing analytics myopia: Revisiting the customer decision journey as a strategic marketing framework. *Journal of Marketing Analytics*, 10(2), 106–113. <https://doi.org/10.1057/s41270-021-00112-5>
 11. Walker, R. (2021). Customer engagement and analytics in digital marketing. *Journal of Interactive Marketing*, 53, 22–35. <https://doi.org/10.1016/j.intmar.2020.10.001>
 12. Al Khaldy, M., Al Obaydi, B., & Al Shari, A. (2023). The impact of predictive analytics and AI on digital marketing strategy and ROI. *International Journal of Business Analytics*, 10(2), 45–60. <https://doi.org/10.4018/IJBAN.20230401>

ИНТЕГРИРАНЕ НА ИЗКУСТВЕН ИНТЕЛЕКТ В УПРАВЛЕНИЕТО НА ЧОВЕШКИТЕ РЕСУРСИ: РЕЗУЛТАТИ ОТ ПИЛОТНО ИЗСЛЕДВАНЕ В БЪЛГАРСКИ ОРГАНИЗАЦИИ

Милена Миленкова¹,
e-mail: mmilenkova@unwe.bg,

Абстракт

Докладът представя резултатите от пилотно анкетно изследване проведено с 52 организации от различни сектори, които функционират на територията на Р България. Основната цел на изследването е да анализира нагласите, опита и предизвикателствата, които срещат организациите в България при прилагане на ИИ в управлението на човешките ресурси. Данните са събрани през периода юни – септември 2025 г.

Abstract

The report presents the findings of a pilot survey conducted among 52 organizations from various sectors operating within the territory of the Republic of Bulgaria. The primary objective of the study is to examine the attitudes, experiences, and challenges encountered by organizations in Bulgaria in the process of implementing artificial intelligence in human resource management. The data were collected between June and September 2025.

Ключови думи: Управление на човешките ресурси (УЧР), изкуствен интелект (ИИ), иновативн о управление на човешките ресурси (ИУЧР).

Keywords: Human Resource Management (HRM), Artificial Intelligence (AI), innovation in human resource management (IHRM)

JEL: M12, M54, O15

1. Въведение

Днес, когато говорим за изкуствен интелект, често мислим за бъдещето – за това какво предстои. В същото време процесите, които се развиват тук и сега, са изключително динамични и дълбоки, до степен че често липсва време за задълбочен анализ и осмисляне на настъпващите промени. Това не е просто технологична трансформация – това е промяна в културата на вземане на решения и в начина, по който се разбира и управлява човешкият потенциал.

„Управление на човешките ресурси без човешки ресурси? – Иновативното управление в ерата на изкуствения интелект – това ли е?“

Докладът представя резултатите от пилотно анкетно изследване, проведено сред 52 организации, функциониращи на територията на Република България. Основната цел на изследването е да анализира нагласите, опита и предизвикателствата, пред които са изправени организациите в България при прилагането на изкуствен интелект в управлението на човешките ресурси.

Изкуственият интелект постепенно, но устойчиво навлиза в процесите и подходите, чрез които се управляват хората в организациите. Макар и все още в начален етап, интересът към дигитализацията и автоматизацията в тази област, както и необходимостта от научна обосновка

¹ Докторант към катедра „Управление“ на УНСС

и емпирични данни, изискват развитието на изследвания и практики в синхрон с реалните потребности на организациите.

През последните години наблюдаваме нарастваща интеграция на изкуствения интелект в управлението на човешките ресурси на глобално ниво. Водещи международни проучвания показват, че ИИ се използва не само за автоматизация на административни процеси, но и за стратегически функции като анализ на ангажираността, прогнозиране на текучеството и персонализиране на обучението. В настоящата изследователска рамка българските организации все по-често осъзнават необходимостта от адаптация към тези глобални тенденции, като същевременно търсят баланс между технологичните възможности и човешкия фактор.

2. Методология

Използван е комбиниран изследователски подход, включващ както количествени, така и качествени методи. Изследването е проведено в периода юли – септември 2025 г. чрез онлайн анкетно проучване, изпратено до 120 организации от частния и публичния сектор. Критерият за подбор на организациите е наличието на публично достъпна информация, че те проявяват интерес към иновативни технологии и практики.

От поканените 120 организации са получени отговори само от 52 организации, които формират изследваната извадка.

Анкетният инструмент включва както категориални въпроси, така и твърдения, оценявани по скали от типа на Ликърт. Анализът обхваща показателите от Ликърт-скалите, сравнение на средни стойности и описателна статистика. В допълнение към затворените въпроси са включени и отворени, които позволяват на респондентите да споделят наблюдения и обратна връзка, базирана на реалния им организационен опит. Тази информация предоставя ценен качествен контекст и допринася за по-задълбочено разбиране на практиките и нагласите, свързани с нивото на въведените иновативни технологии в управлението на човешките ресурси.

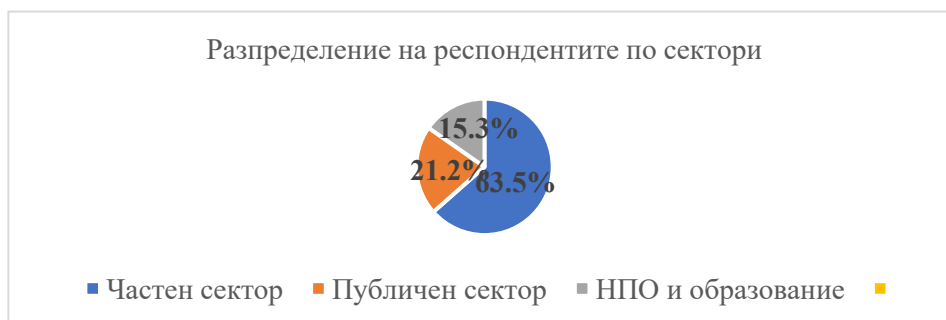
Участниците имат възможност да изразят желание за включване в последващо, по-задълбочено изследване, като оставят координати за контакт. С част от тях са проведени допълнителни насочени полуструктурирани интервюта, целящи да уточнят нагласите и да подпомогнат интерпретацията на получените резултати.

Изследването има пилотен характер и цели да постави основите за бъдещи, по-мощабни емпирични анализи. Макар броят на участващите организации да не позволява пълна генерализация на резултатите, събраните данни предоставят надеждна основа за идентифициране на основни тенденции и модели на възприемане на ИИ в HR контекста.

3. Резултати

Профил на участниците

Резултатите от изследването показват, че интересът и активността по отношение на внедряването на ИИ идват основно от бизнеса.



Фигура 1: Профил на участниците

Източник: Авторово анкетно изследване, 2025 г.

Сред участващите организации преобладават представители на частния сектор, основно от сфери като услуги, образование и информационни технологии. Налице е и участие на публични институции и неправителствени организации, което допринася за по-пълна картина на нагласите в различни организационни контексти. Повечето участници заемат ръководни или експертни позиции в областта на човешките ресурси, което гарантира информирани и професионално обосновани отговори.

Ниво на запознатост с ИИ

Данните показват, че **27%** от респондентите се самоопределят като *добре запознати* с изкуствения интелект, **41%** посочват, че имат *базови познания*, а останалите познават *само отделни приложения* на ИИ.

Тези резултати сочат към умерено ниво на осведоменост по темата, съчетано с изразен интерес и заявена потребност от допълнително обучение и практическо запознаване с възможностите на ИИ. Това показва, че в организационен контекст изкуственият интелект все още се възприема предимно като нововъзникваща, но перспективна област, чието приложение изисква повишаване на компетентността и информираността сред специалистите по човешки ресурси.

В отворените отговори се открояват няколко ключови тенденции. Част от респондентите посочват, че **запознаването с ИИ е резултат основно от самообучение, участие в семинари и неформални източници на информация**, а не от системни вътрешноорганизационни обучения. Това показва, че процесът на изграждане на компетентност в областта на изкуствения интелект е все още на индивидуално, а не на институционално ниво.

Освен това се наблюдава и **различие между сектори** – представители на технологични и ИТ компании демонстрират по-висока степен на информираност, докато в организации от публичния сектор и по-традиционни частни организации темата остава по-слабо позната.

В обобщение може да се заключи, че **запознатостта с изкуствения интелект сред организациите в България е все още начална, но потенциалът за развитие е значителен**. Налице е отчетлив стремеж към обучение и осъзнаване на необходимостта от развитие на нов тип компетентности, които ще бъдат ключови за успешната дигитална трансформация в управлението на човешките ресурси.



Фигура 2: Ниво на запознатост с ИИ на участниците

Източник: Авторово анкетно изследване, 2025 г.

Области на приложение на ИИ

Резултатите от изследването показват, че изкуственият интелект в организациите в България най-често се използва в процесите по подбор на персонал и обучение на служителите. Данните са обобщени в таблица 1.

Таблица 1: Област на използване на ИИ в УЧР от организациите в България

Област на използване	Използва се (%)	Планира се (%)	Не се използва (%)
Подбор на персонал	38.5	28.8	32.7
Обучение и развитие	30.8	42.3	26.9
Оценка на представяне	25.0	34.6	40.4
Аналитика на хората	19.2	38.5	42.3

Предимства и предизвикателства

Днес, когато алгоритъмът може да реши кого да наеме, обучи или възнагради – къде остава човекът в управлението на хора?

Средните стойности на отговорите ясно показват, че респондентите отчитат конкретни ползи от прилагането на изкуствен интелект в управлението на човешките ресурси. Най-високи оценки получават твърденията, че **ИИ улеснява рутинните дейности (4.32/5)** и **подобрява качеството на подбор (4.08/5)**. В същото време, респондентите посочват и умерено изразено опасение, че **изкуственият интелект може да представлява заплахата за човешкия фактор (3.41/5)**.

Тези резултати очертават двоен фокус в нагласите – от една страна, осъзнаване на ефективността и оптимизационния потенциал на ИИ, а от друга – предпазливост по отношение на неговото въздействие върху човешките роли и взаимоотношения в организацията.

В контекста на тези нагласи основните предизвикателства, посочени от участниците, включват **липсата на достатъчно обучени специалисти и етичните дилеми**, свързани с прозрачността и справедливостта на алгоритмите. Това потвърждава, че успехът на прилагането на ИИ в управлението на човешките ресурси зависи не толкова от самата технология, колкото от **хората, които я прилагат**, и от способността на организациите да изградят подходяща култура на доверие и отговорност.



Фигура 2: Възприемани ползи и рискове от прилагането на ИИ в управлението на човешките ресурси

Източник: Авторско анкетно изследване, 2025 г.

Наред с технологичните и организационните предизвикателства, респондентите поставят въпроса за етичните измерения на използването на ИИ – особено по отношение на прозрачността при вземане на решения, защитата на личните данни и недопускането на алгоритмична пристрастност. Тези теми изискват последователно внимание както от изследователската общност, така и от практиците по управление на човешките ресурси.

Общи нагласи

Когато се обобщят всички индикатори, резултатите показват, че отношението към изкуствения интелект е **преобладаващо положително**, със средна оценка около **4.1 от 5**. Това свидетелства за нарастващо доверие в потенциала на технологиите да подпомагат процесите по управление на човешките ресурси.

Организациите в България **не възприемат ИИ като заплаха**, а по-скоро като **партньор в процеса на вземане на решения** и инструмент за повишаване на ефективността. Тази нагласа представлява важна основа за бъдещо развитие и по-широко внедряване на изкуствения интелект в организационната практика.

Получените резултати кореспондират с глобалните тенденции, според които ИИ се възприема все повече като инструмент за подкрепа, а не за заместване на човека в организационната среда. В българските организации обаче внедряването на ИИ е все още в ранен етап, което обяснява известна предпазливост и необходимост от натрупване на практически опит.

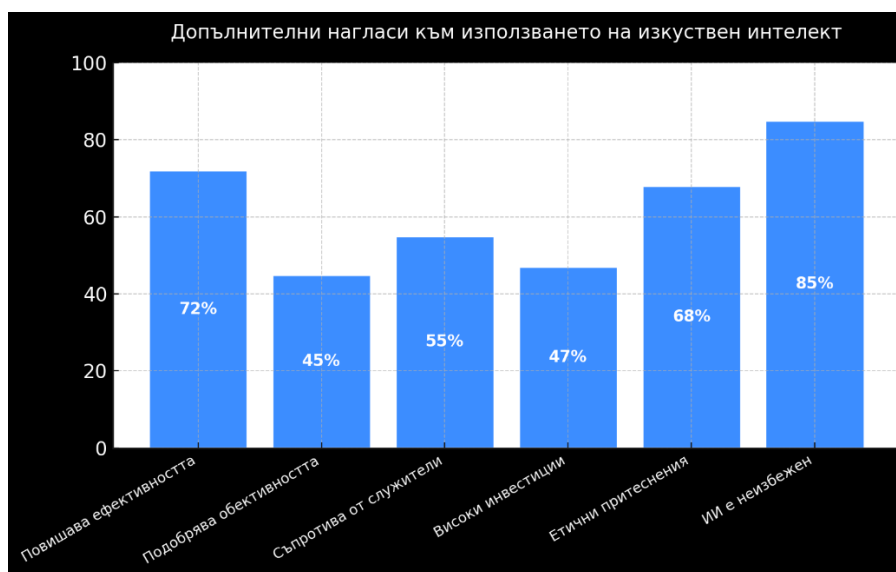
Допълнителни наблюдения и тенденции

Получените данни предоставят по-задълбочен поглед върху възприятията на организациите относно въздействието на изкуствения интелект. Около **72%** от респондентите вярват, че ИИ **повишава ефективността** на организационните процеси, докато **45%** изразяват съмнение относно **обективността при вземане на решения**, когато те се подпомагат от алгоритми.

Повече от половината (**55%**) отчитат **съпротива от страна на служителите**, което може да се разглежда като индикатор за необходимост от по-добра комуникация и изграждане на доверие в процеса на дигитална трансформация. Почти половината (**47%**) посочват **високите инвестиционни разходи** като основна бариера пред внедряването на ИИ.

Значителна част от анкетираните (**68%**) изразяват **загриженост по отношение на етичните аспекти** на използването на изкуствен интелект, включително прозрачността, защитата на личните данни и риска от алгоритмична пристрастност. Въпреки тези притеснения, **85%** от участниците са убедени, че **навлизането на изкуствения интелект е неизбежно** и че организациите следва активно да се подготвят за това.

Тези резултати показват едновременно **висока степен на осъзнатост и реализъм** – организациите разпознават потенциала на изкуствения интелект, но и ясно съзнават предизвикателствата, които съпътстват неговото прилагане.



Фигура 3: Допълнителни наблюдения и тенденции от ИИ в управлението на човешките ресурси

Източник: Авторово анкетно изследване, 2025 г.

4. Перспективи за бъдещи изследвания

Резултатите от проведеното пилотно изследване очертават необходимостта от по-задълбочени проучвания относно ролята и ефектите на изкуствения интелект в управлението на човешките ресурси в България. Необходим е по-широк емпиричен обхват, който да включва организации с различни по размер и видове дейности, както и проследяване на динамиката във времето при внедряването на интелигентни технологии.

Особен интерес представлява изследването на **взаимодействието между човешкия фактор и алгоритмичните решения**, включително темите за доверие, прозрачност и възприемана справедливост при използването на ИИ в процеси като подбор, оценка и развитие на персонала (но не само!). В тази изследователска рамка следва да се проучат и **етичните и нормативни аспекти**, които определят границите и възможностите за устойчиво прилагане на изкуствения интелект в управлението на хора.

Препоръчително е бъдещи изследвания да включват **качествени методи** – интервюта, казуси и наблюдения – с цел по-дълбоко разбиране на организационните практики и на начина, по който специалистите по човешки ресурси възприемат и адаптират технологичните решения.

Развитието на подобна изследователска линия би допринесло както за **натрупване на научни доказателства**, така и за **формиране на практически насоки** за ефективното и етично използване на изкуствения интелект в управлението на човешкия капитал в българския контекст.

Особен интерес представлява и изследването на организационните модели за въвеждане на ИИ – как се вземат решения, какви стратегии за промяна се прилагат и каква е ролята на лидерството. Анализът на тези аспекти би допринесъл за по-добро разбиране на факторите, които определят успеха или неуспеха на технологичните иновации в управлението на човешкия капитал.

5. Заключение и препоръки

Резултатите от пилотното изследване потвърждават, че изкуственият интелект постепенно се превръща в значим фактор за трансформацията на управлението на човешките ресурси в България. Нагласите на организациите са предимно положителни, като ИИ се възприема основно като инструмент за повишаване на ефективността и подобряване на качеството на вземане на решения. В същото време, съществуват ясно изразени предизвикателства, свързани с недостига на компетенции, етичните въпроси и нуждата от организационна готовност за промяна.

На база на резултатите могат да бъдат формулирани следните основни препоръки:

1. **Постепенно и отговорно внедряване на изкуствения интелект** в процесите по управление на човешките ресурси, като се отчита организационният контекст и зрелостта на технологичната инфраструктура.
2. **Инвестиции в обучение и развитие на дигитални компетентности** сред специалистите по човешки ресурси, с цел ефективно използване и критично разбиране на инструментите, базирани на ИИ.
3. **Изграждане на етична рамка и регулации**, гарантиращи прозрачност, защита на личните данни и справедливост при алгоритмичните решения.
4. **Насърчаване на партньорства между академичната общност и бизнес организациите**, с цел обмен на знания, изследвания и добри практики за устойчиво внедряване на иновации.

Тези мерки могат да подпомогнат формирането на среда, в която технологичните решения се интегрират в подкрепа на човека, а не като негов заместител — създавайки предпоставки за по-гъвкаво, ефективно и етично управление на човешкия капитал в дигиталната ера.

В заключение може да се обобщи, че **изкуственият интелект не заменя човешкия фактор, а го допълва и развива**, като подпомага по-интелигентното и ефективно управление на хората. Прилагането на ИИ в управлението на човешките ресурси е **неизбежна стъпка** от развитието на

съвременните организации, но **успехът му зависи от постигането на баланс между технология и човешки ценности.** Изкуственият интелект следва да се разглежда като **инструмент за подкрепа, а не за заместване на човека** – средство за разгръщане на потенциала, а не за неговото ограничаване.

Литература / References

1. Научната публикация „Предизвикателства пред индустриалния растеж в България“, сборник от конференцията, 2018, стр. 54-59 Електронно копие налично на: <http://www.industrialgrowth.eu/wp-content/uploads/2018/11/5.pdf> ИНТЕЛИГЕНТЕН РАСТЕЖ НА БИЗНЕСА ЧРЕЗ ИНОВАЦИИ Автор: Цветана Стоянова.
2. Fenwick A, Molnar G and Frangos P (2024) Revisiting the role of HR in the age of AI: bringing humans and machines closer together in the workplace. *Front. Artif. Intell.* 6:1272823. doi: 10.3389/frai.2023.1272823.
3. Łukasik-Stachowiak, K. (2023) ‘Uncertainties and Challenges in Human Resource Management in the Era of Artificial Intelligence’, *Scientific Papers of Silesian University of Technology. Organization & Management / Zeszyty Naukowe Politechniki Slaskiej. Seria Organizacji i Zarzadzanie*, (181), pp. 341–356. doi:10.29119/1641-3466.2023.181.23.
4. Статия „Как генеративният AI предефинира управлението на човешкия капитал: Рационализиране на HR процесите и подобряване на изживяването на служителите“ от Тина Фигероа главен директор "Човешки ресурси" в Crix, публикувана CHRO Excellence: HR Strategy & Implementation, Jul2024, Business Source Ultimate.
5. Anguelov, K., Stoyanova, T. & Tamošiūnienė, R. 2020, "Research of motivation of employees in the IT sector in Bulgaria", *Entrepreneurship and Sustainability Issues*, vol. 7, no. 3, pp. 2556-2567

ПРИЛОЖЕНИЕ НА ИЗКУСТВЕНИЯ ИНТЕЛЕКТ В ОБРАЗОВАНИЕТО - ЗА И ПРОТИВ

Application of Artificial Intelligence in Education - Pros and Cons

Наталия Маринова¹,
e-mail: n.marinova@uni-svishtov.bg

Абстракт

Преподаването, ученето и управлението на учебния процес в днешните образователни институции се трансформира от редица информационни технологии, дигитални платформи, софтуерни системи и средства за управление на класната стая, за оценяване и за анализиране на ангажираността на обучаемите и за подобряване на интерактивността на образователните ресурси, в т. ч. и от такива с изкуствен разум. Парадигмите за използване на системи с изкуствен интелект в образованието определят предназначението на такива платформи и приложения като водещи, поддържащи или упълномощаващи когнитивното обучение на учащите. В настоящият доклад се прави опит да се систематизират предимствата и недостатъците на приложението на системи с изкуствен интелект в образователния процес и да се изложат някои идеи и предложения на Европейския съюз и България за преодоляване на последните.

Abstract

Teaching, learning, and management of the educational process in today's educational institutions are being transformed by a variety of information technologies, digital platforms, software systems and classroom management tools, assessment tools, and tools for analysing learners' engagement and enhancing the interactivity of educational resources, including those with artificial intelligence. Paradigms for using artificial intelligence systems in education define the purpose of such platforms and applications as leading, supporting, or empowering learners' cognitive learning. This paper attempts to systematize the advantages and disadvantages of applying artificial intelligence systems in the educational process and to present some ideas and proposals from the European Union and Bulgaria to overcome the latter.

Ключови думи: дигитални платформи с изкуствен интелект, образование.

JEL: C88, L86, Q55

Въведение

Информационните технологии са водещ фактор в процеса на образователно придобиване на нови знания и умения, опосредствайки прехода от модел на предаване на информацията към модел на взаимно, активно, по-самостоятелно и отговорно учене. Обучителите винаги са се опитвали да намерят нови начини за адаптиране към нуждите на обучаемите чрез мотивиращи учебни програми, нови методи на преподаване, ефективни дискусии, методи за модернизиране на класните стаи, инструменти на социалните медии, софтуерни приложения за генериране и обработка на текстово и мултимедийно съдържание, за онлайн тестване на знания, за симулиране на ролеви игри и др.

Съвременните образователни институции функционират в условия на глобално и непрестанно създаване на нови факти и данни с експоненциални темпове, което изисква от тях да адаптират

¹ Associate Professor, PhD (Econ). Department of Business Informatics, D. A. Tsenov Academy of Economics, Svishstov, ORCID ID: 0000-0001-8732-7564.

начините на преподаване, изследване и обучение на учащите, с цел подобряване на тяхната конкурентоспособност. А това изисква промяна в управленските процеси и образователните нагласи и практики в посока повишаване знанията и компетенциите на обучаемите чрез предоставяне на персонализирани образователни ресурси, интегрирани IoT технологии и системи с изкуствен интелект, насочващи, поддържащи или упълномощаващи процеса на обучение (Ouyang & Jiao, 2021).

Предимства на приложението на системи с изкуствен интелект в образователния процес

Формите на приложение на системите с изкуствен интелект в обучителния процес са най-различни. На база функционалните възможности на информационните технологии, използвани за реализация на изкуствен разум (изкуствени невронни мрежи, машинно обучение, роботи, разширена реалност, виртуална реалност), Попов разделя тези форми на четири основни групи (Попов, 2023):

- дигитални платформи с елементи на изкуствен интелект;
- софтуерни решения, поддържащи иновативни концепции за обучение;
- комплексни софтуерни системи;
- средства за изпълнение на отделни функции.

В цитираната публикация последователно са систематизирани положителните страни на използването на различни образователни продукти с изкуствен интелект (интелигентен кампус, интелигентна класна стая, системи за адаптивно обучение, системи за персонализирано обучение, интелигентен преподавател, интелигентен настойник и др.) от обучаеми и обучители. За изследователските цели на настоящия доклад от особен интерес е първата форма на приложение - дигиталните платформи с елементи на изкуствен интелект – които се описват като „комплекс от интегрирани хардуерни устройства, софтуерни системи и комуникации“.

Списък на полярни по настоящем образователни платформи с изкуствен интелект е показан в Таблица 1:

Таблица 1. Популярни образователни платформи с изкуствен интелект

Приложение	Функционалност	Производител или уеб-ресурс
Khanmigo	Умен асистент за всякакви въпроси, свързани с учителската дейност	Khan Academy
Coursera Coach	Асистент с възможности за виртуална реалност, използваща машинно обучение за превод на курсове	Coursera
Даскал GPT	Специализиран чатбот с изкуствен интелект за български учители	Синдео
СмаpТест	Българска платформа за автоматизация на процеса по изпитване и оценяване	https://www.smartest.bg/
ClassBuddy	Българска платформа за създаване на интерактивни игри за усвояване и проверка на знанията на учениците	Нимеро ООД
Ask AI	Безплатна разширена търсачка	https://iask.ai/
Flawlessly AI	Инструмент за граматическа проверка	https://flawlessly.ai/
Magic School AI	Платформа за свързване на преподавателите в общност	Magic School, Inc.
Canva Magic Studio	Инструмент за автоматично генериране на дизайн	Canva
NotebookLM	Помощник за учене по дадена тема	https://notebooklm.google.com
Tutor AI	Инструмент за самообучение в дадена област	AI Tutor

Learn About	Инструмент за бързо изучаване на дадена тема	https://learning.google.com/experiments/learn-about
SciSpace	Инструмент за анализ на научна литература и проверка на твърдения	PubGenius Inc.
TalkPal AI	Приложение за учене на езици с изкуствен интелект учител	Talkpal, Inc.

Източник: Авторова систематизация.

На основа проучването на различни литературни източници и на съществуващите софтуерни решения с изкуствен интелект с възможности за трансформация на образователния процес можем да обособим **преимуществовата** от използването на такива в следните направления:

1) Намалване на тревожността и стреса у обучаемите, причинени от учене в лабораторна среда или от взаимодействие с възпитатели (Sears, 2018).

2) Получаване на подкрепа при разпознаването на емоционални нужди, на речеви и езикови проблеми, на признаци на дислексия и на двигателни затруднения у учениците от асистивни технологии за справяне с увреждания, за превод в реално време (например Google Translate) и за потапяне във виртуални светове (Google Expeditions).

3) Опосредстване усвояването на общообразователно и специфично учебно съдържание чрез насочени образователни ресурси и интелигентни системи за преподаване като Algebra Tutor, SQL-Tutor, EER-Tutor, Mathematics Tutor, eTeacher, Carnegie Learning и др.

4) Персонализирано подпомагане на обучаемите чрез софтуерни системи за адаптиране на учебни уроци (DreamBox), на съдържание (Knewton Alta), на изпитни тестове (Magoosh) и на квалификационни курсове за учене през целия живот.

5) Повишаване на интереса на учащите към учебното съдържание и научните експерименти чрез стимулационни, интерактивни (Engage VR, Mainstay, Packback) и игрови образователни платформи (Minecraft: Education Edition, Duolingo, Kahoot!, Labster и др.) и проследяване на тяхното поведение, мотивация и ангажираност към процеса на учене чрез технологии за управление на класната стая.

6) Получаване на незабавна обратна връзка за силните и слабите страни и за пропуските в знанията на учащите чрез платформи за следене участието на учениците в часовете, за предиктивен анализ и за проследяване на развитието на учащите (ClassDojo).

7) Автоматизиране на процеса на преподаване и изпитване посредством софтуер за оптимизация на графици на учителите и учениците, за оптимизация на графици на училищните дейности (AI Planner), за оценяване на тестове и задачи (Gradescope) и за откриване на плагиатство в разработките на учащите.

8) Съкращаване на времето на обучаващите за разработване на учебни планове, подготовка на уроци и финансово осигуряване на образователния процес чрез софтуерни инструменти за създаване на съдържание (Magic School AI) и за ефективно управление на бюджета (Fetchy).

9) Намалване на разходите за заплати в образователните институции чрез използване на социални работи за обучение на деца в предучилищна възраст или усвояване на определени умения (говорене на чужд език например).

10) Оптимизиране на изследователския процес по проучване, систематизиране и резюмиране на уеб-източници чрез използване на генеративни чатбот приложения (ChatGPT, Gemini, Copilot, Meta AI, Ernie Bot, Grok, Claude, DeepSeek, Mistral и др.).

11) Споделяне на добри образователни практики относно приложението на системи с изкуствен интелект в обучението с общности от обучители и/или обучаеми чрез средствата на социалните медии и специализирани уеб-сайтове (<https://prepodavame.bg/>, <https://www.educatorstechnology.com/2025/09/ai-tools-and-resources-for-teachers.html>) и др.

Проблеми в използването на системи с изкуствен интелект в образованието

Бурното развитие на генеративните системи с изкуствения интелект през последните три години задълбочи зависимостта на хората от дигитални устройства и промени работните места, професиите и компетенциите на работниците и служителите в бизнес организациите. Днешните GenAI приложения могат да създават текст, изображения, видеоклипове и музикални произведения и да автоматизират организационните задачи по нов и ефективен начин. През настоящата 2025 г. еуфорията около този възход стихна и започна процес на оценка на ефектите от инвестициите в производителите на такива и от използването им в различни сфери на икономиката, в т.ч. образователната.

Много от **общите рискове** на генеративните приложения (халюциниране на факти, изкривяване на резултати, отравяне на данни, алгоритмично пристрастие, липса на познания у потребителите за работа с GenAI инструментите и др.) важат и за образователните решения и платформи с изкуствен разум. Съвременните образователни системи обаче са подложени на двоен натиск:

- от една страна, от тях се очаква да се трансформират под влиянието на изброените по-горе форми на приложение на системи с изкуствен интелект;
- от друга страна, те са изправени пред предизвикателството да обучат бъдещите участници на пазара на труда на съществуващи или на ново появили се в резултат на внедряването на системи с изкуствен интелект професии.

Изкуственият интелект вече оказва реално въздействие върху трудовия пазар и променя трансформира, по който хората работят, учат и се преквалифицират, поставяйки ги пред безпрецедентното предизвикателство да се адаптират с по-бързи от развитието на технологиите темпове. Мащабно проучване (LHH, 2025) на над 200 000 изгубили работата си и започнали програма за кариерен преход души в страни от Западна Европа и САЩ в периода януари 2024 г. - март 2025 г. на LHH (глобален доставчик на услуги за подпомагане при прекратяване на работа, кариерна трансформация и развитие на лидерски умения) например показва, че хората, чиито работни места са били изместени от решение с изкуствен интелект, остават без работа за повече от 12 месеца над два пъти по-често в сравнение с останалите кандидати. Парадоксален извод в цитираното изследване е най-силната засегнатост от подобни процеси на технологичните индустрии.

Специфичните **недостатъци** на образователните системи с изкуствен интелект са следните:

- 1) Вероятност от загуба у обучаемите на съсредоточеност и ирационалност (каквото не е поведението на интелигентните агенти) при решаване на задачите. Макар социалните работи да дават информация със същото качество както своите човешки аналози, засега те се използват предимно като инструменти за учене, а не като учители.
- 2) Ограничаване на „живото“ взаимодействие с обучителите в класните стаи, предоставящи менторство и социална подкрепа на обучаемите. Инструментите с изкуствен интелект не могат да заместят емпатията, креативността и интуитивността на човешките същества.
- 3) Проява на неетично поведение от страна на обучаемите, които разчитат изцяло или частично на генеративни приложения за да изпълняват учебни задачи и проекти. Подобно отклонение от моралните образователни принципи води до неточно измерване и оценяване на резултатите на учащите в краткосрочен план, до приемане за норма на тези форми на неетично образование и до подкопаване на уважението и усилията на обучаващите в средносрочен план и до трайни последици за развитието на знаеща, отговорна, квалифицирана и компетентна работна сила в дългосрочен план.
- 4) Загуба на способност за разбиране на изучавания материал, за аналитично мислене, за вземане на самостоятелни решения и проява на мързел у обучаемите и обучителите поради прекомерно доверяване във функционалните възможности и в резултатите, връщани от приложенията с изкуствен интелект. Делегирането на все повече учебни задачи на генеративни инструменти отслабва рефлекс за проверка на предложеното съдържание, притъпява уменията

за творческо съзидание и влошава способността за справяне със сложни проблеми в условия на напрежение и при двете страни в учебния процес.

5) Закърняване на когнитивните функции, атрофия на мозъчните клетки и общо намаление на когнитивния капацитет на човешкия мозък. Редица изследвания на ефектите от взаимодействието на хора с интелигентни системи регистрират намаляване на мозъчна активност, особено в алфа и тета вълновите диапазони (Kosmyna, 2025), и отслабване на човешките познавателни умения (т. нар. „когнитивно разтоварване“, причинено от склонността да се разчита прекалено на системи с изкуствен интелект за генериране на съдържание).

6) Предоставяне на голям обем лични и чувствителни данни на учащите в „ръцете“ на доставчиците на генеративни услуги с изкуствен интелект. За нуждите на предоставяне на персонализирано обучение системите с изкуствен интелект събират, съхраняват и анализират информация от частен характер, злоупотребата с която може да застраши поверителността и доверието на учащите.

7) Пристрастяване към общуването с чат-ботове, водещо до промяна в поведението на потребителите – повишена тревожност, дигитално „зомбиране“, промени в настроението, емоционална привързаност и дори суицидни прояви с фатален край. Според проучване на OpenAI и MIT Media Lab сред хиляди потребители на ChatGPT например (Fang, 2025) дългосрочната употреба на тази програма с изкуствен интелект увеличава риска от пристрастяване и поражда на дълбока парасоциална свързаност с нея, особено при самотните и хората без социални контакти.

8) Задълбочаване на дигиталните неравенства и качеството на обучение между образователните институции с осигурено и с недостатъчно финансиране за внедряване на усъвършенствани системи с изкуствен интелект и за професионално квалифициране на работещите с тях преподаватели.

Преодоляването на изложените общи и специфични проблеми в използването на системи с изкуствен интелект изисква целенасочени и разностранни мерки от всички участници в образователния процес, пълното описание на които надхвърля обемните рамки на настоящия доклад. Отговорното използване на технологията на изкуствения интелект в образованието изисква стратегическо планиране и спазване на регулативни правила, наложени от национални, наднационални и международни юрисдикции.

Последните ключови инициативи на Европейския съюз за повишаване степента на използване на системи с изкуствен интелект (през 2024 г. само 13.5% от европейските предприятия са прилагали такива в своята дейност) са наскоро анонсираните две стратегии за извеждане на Европа на челно място в този технологичен сектор. В Стратегията за прилагане на изкуствен интелект (European Commission, 2025) за създаване на фабрики, гигафабрики, инвестиционни съоръжения и авангардни центрове за интегриране на технологията в малките и средните предприятия от 11 ключови и публични сектори на европейската икономика е отделено специално място на обучението на експерти и привличането на таланти в областта на изкуствения интелект в Стария континент. Стратегията за изкуствен интелект в науката (European Commission, 2025) цели постигане на напредък чрез инструмента RAISE (виртуален европейски институт за обединяване и координиране на ресурси за разработване и прилагане на изкуствен интелект в науката) и предприемане на мерки за привличане на световни научни таланти и висококвалифицирани специалисти в ЕС чрез инициативата „Изберете Европа“.

България все още изостава в процеса на внедряване на системи с изкуствен интелект в ключови икономически сектори, в т. ч. този на образованието, от останалите страни членки на Европейския съюз, от САЩ и от Китай поради липсата на достатъчно специалисти, слабо развитие на науката и малкото инвестиции в проучвания и иновации. Конкретни стъпки, политики и действия за преодоляване на това изоставане могат да се открият в два изчерпателни документа. Насоки към учителите, учениците, родителите и училищата за целесъобразно, подходящо и отговорно използване на изкуствен интелект в различни педагогически сценарии и операции в учебния процес са разписани в проект на Министерството на образованието на

науката за използване на изкуствен интелект в образователната система у нас (Министерство на образованието и науката, 2024). Анализ на възможностите за инвестиране в изследвания и разработки с международна стойност и увеличена патентна дейност у нас е направен от Българска стопанска камера (Каспарян & Недева, 2025 г.).

Заклучение

Образователните организации създават интелектуална собственост под формата на открития, изобретения и технологии и ги прехвърлят към пазара, за да се превърнат в подобряващи човешкия живот продукти и услуги. Водещата роля на САЩ и на Китай в процеса на технологичен трансфер и разработване на системи с изкуствен интелект предизвиква Европейският съюз да предприеме стратегически мерки за намаляване на това изоставане. На Стария континент обаче екосистемата на разработване и приложение на технологии с изкуствен разум все още не е толкова зряла, а стартиращите производители са задушавани от бюрокрация и регулации. Въпреки че реалните ползи от интеграцията на решения с изкуствен интелект в образованието са значителни, пътят към широкото им приложение е изпълнен с етични съображения и практически препятствия.

Използвана литература

1. European Commission. (2025, October 8). *European AI in Science Strategy*. Retrieved October 21, 2025, from https://ec.europa.eu/research-and-innovation/ec.europa.eu/strategy/strategy-research-and-innovation/our-digital-future/european-ai-science-strategy_en#:~:text=The%20Commission%20has%20launched%20a%20European%20Strategy%20for,support%20scientists%20to%20adopt%20them%20
2. European Commission. (2025, October 8). *Shaping Europe's leadership in artificial intelligence with the AI continent action plan*. Retrieved from https://commission.europa.eu/topics/eu-competitiveness/ai-continent_en
3. Fang, M. (. (2025, March 21). *Early methods for studying affective use and emotional wellbeing in ChatGPT: An OpenAI and MIT Media Lab Research collaboration*. Retrieved October 20, 2025, from <https://www.media.mit.edu/posts/openai-mit-research-collaboration-affective-use-and-emotional-wellbeing-in-ChatGPT/>
4. Kosmyna, N. (2025, June 20). CNN: AI's effects on the brain. (O. Jimenez, Interviewer) Retrieved October 20, 2025, from <https://www.media.mit.edu/articles/a-i-s-effects-on-the-brain/>
5. LHH. (2025). *The reinvention imperative: How AI is reshaping jobs, individual careers and talent strategies*. Retrieved from <https://lhh.com/info.lhh.com/us/en/the-reinvention-imperative>
6. Ouyang, F., & Jiao, P. (2021). Artificial intelligence in education: The three paradigms. *Computers and Education: Artificial Intelligence*, 2(100020). doi:<https://doi.org/10.1016/j.caeai.2021.100020>
7. Sears, A. (2018, april 14). *The role of artificial intelligence in the classroom*. Retrieved october 5, 2020, from <https://elearningindustry.com/artificial-intelligence-in-the-classroom-role>
8. Каспарян, Х., & Недева, М. (2025 г., Април). *Изкуственият интелектТ: тенденции, регулации и сценарии за развитие*. Retrieved from <https://www.bia-bg.com/analyses/view/34190/>

9. Министерство на образованието и науката. (Януари 2024 г.). *Насоки за използване на изкуствен интелект в образователната система*. Изтеглено на 20 Октомври 2025 г. от <https://danybon.com/>: https://danybon.com/wp-content/uploads/2024/02/Nasoki-izpolzvane-II_190224-v30.01.2024.pdf
10. Попов, В. (2023). Приложение на изкуствения интелект в обучението. *Национална конференция "Дигитална трансформация на образованието – проблеми и решения, оценяване и акредитация"* (pp. 270-274). Русе: Русенски университет "Ангел Кънчев".

ТЕХНОЛОГИЧНИ РИСКОВЕ ОТ ДИГИТАЛИЗАЦИЯТА В НЕАКАДЕМИЧНА СРЕДА

TECHNOLOGICAL RISKS OF DIGITALIZATION IN NON-ACADEMIC ENVIRONMENT

Мариана Ковачева¹
e-mail: mkovacheva@unwe.bg¹

Абстракт

Дигитализацията се превръща в определяща характеристика на организационната трансформация, простирайки се отвъд академичните институции в бизнеса, правителството и секторите на обществените услуги. Макар че повишава ефективността, иновациите и свързаността, едновременно с това въвежда нови категории технологични рискове, които заплашват оперативната стабилност, целостта на данните и етичното управление. Този доклад разглежда технологичните рискове от дигитализацията в неакадемична среда чрез структурирана класификация на осем категории: заплахи за киберсигурността, рискове за поверителността на данните, системни повреди и прекъсвания, загуба на контрол над цифровите умения, етични и алгоритмични рискове, зависимост от доставчици и платформи, съкращения на работни места и правни и регулаторни предизвикателства. Въз основа на рамката за управление на риска ISO 31000:2018, изследването предоставя систематичен анализ на това как тези рискове могат да бъдат идентифицирани, анализирани и смекчени. Заключениета подчертават, че управлението на технологичните рискове изисква не само технически решения, но и социална, етична и организационна адаптация. Докладът предлага всеобхватна концептуална рамка, с което допринася за подобряване на осведомеността за риска, устойчивостта и отговорното цифрово управление в неакадемичните институции.

Abstract

Digitalization has become a defining characteristic of organizational transformation, extending beyond academic institutions into business, government, and public service sectors. While it enhances efficiency, innovation, and connectivity, it simultaneously introduces new categories of technological risks that threaten operational stability, data integrity, and ethical governance. This paper examines the technological risks of digitalization in non-academic environments through a structured classification of eight categories: cybersecurity threats, data privacy risks, system failures and downtime, loss of digital skills control, ethical and algorithmic risks, dependency on vendors and platforms, job displacement, and legal and regulatory challenges. Drawing upon the ISO 31000:2018 risk management framework, the study provides a systematic analysis of how these risks can be identified, analyzed, and mitigated. The findings underscore that managing technological risks requires not only technical solutions but also social, ethical, and organizational adaptation. By offering a comprehensive conceptual framework, the paper contributes to improving risk awareness, resilience, and responsible digital governance in non-academic institutions.

Ключови думи: digitalization, technological risks, non-academic environment

JEL: C88, L86

Introduction

In today's fast developing world where new technologies emerge every day and integrate in more aspects of our lives, the risks connected to them are also expanding. Non-academic institutions

¹ Гл. ас. д-р към катедра ИТК, УНСС, email: mkovacheva@unwe.bg

are such that do not focus on education, research or some scholarly activities, but they operate in other important areas of society, business and public life.

The digitalization is happening in every sphere around us and helps immensely in developing and integrating digital technologies into all areas of a business. There is a great importance in studying the technological risks in non-academic environments because even that digitalization is offering innovation, efficiency and speed, also introduces threats, dependency and vulnerability that can threaten the stability of business organizations. Understanding these emerging risks will help balance the innovation with security, and will ensure that digital system will remain reliable, ethical and sustainable.

This paper explores eight categories of technological risks which are arising from digitalization outside of academia. The defined groups are as follows: cybersecurity threats, data privacy risks, system failures and downtime, loss of digital skills control, ethical and algorithmic risks, dependency on vendors and platforms, job displacement, and legal and regulatory challenges.

Conceptual background

Digital transformation has become a defining force in modern organizations, yet it is important to distinguish between digitization and digitalization. Digitization refers to the technical conversion of analog information into digital formats - for example, scanning documents or automating data entry. [1] Digitalization, on the other hand, involves a broader social and organizational process that integrates digital technologies into everyday operations, fundamentally altering how value is created, communicated, and managed.[2] In non-academic environments - such as business enterprises, government institutions, and healthcare systems - digitalization extends beyond technology adoption to include process redesign, data-driven decision-making, and cultural adaptation. While it offers efficiency and innovation, it simultaneously introduces new categories of technological risk that can disrupt operations, compromise data integrity, and create ethical or regulatory challenges.

On the other hand, technological risk refers to the potential for adverse outcomes resulting from the design, implementation, or use of technology. It encompasses failures of systems or infrastructures as well as unintended human, ethical, or organizational consequences. These risks often stem from the uncertainty that accompanies innovation and from the interdependence of digital systems within complex socio-technical environments. In non-academic contexts, technological risks include technical failures, cybersecurity incidents, data privacy breaches, and ethical issues arising from algorithmic decision-making. They also encompass human factors such as skill obsolescence, organizational dependency on external vendors, and social effects like job displacement. Recognizing these diverse dimensions underscores that technological risk is not purely a technical matter but a multidimensional issue combining social, ethical, and operational factors.

When digital transformation progresses more rapidly than organizational structures or human competencies can adjust, this equilibrium is disrupted, leading to a higher probability of system failures, human errors, and the decomposition of digital skills. Consequently, managing technological risk extends beyond implementing technical safeguards - it also requires social adaptation, continuous training, and inclusive design practices that ensure people and technology evolve together. [11]

The ISO 31000:2018 framework offers a globally recognized standard for systematic risk management. It defines risk as the effect of uncertainty on objectives and outlines a continuous process of identification, analysis, evaluation, and treatment. Applying ISO 31000 principles to technological risks ensures a structured and comparable approach to their management. Within non-academic environments, this process enables organizations to identify digital vulnerabilities, assess their potential impact, and implement appropriate mitigation strategies - ranging from cybersecurity protocols to vendor oversight and employee training. [3]

In non-academic environments, digitalization transforms how organizations operate, interact with stakeholders, and deliver services. Businesses leverage digital tools for analytics and automation, governments deploy digital platforms for public service delivery, and healthcare institutions integrate electronic systems for patient management. However, these benefits also amplify exposure to

technological uncertainty - manifesting as cybersecurity threats, data privacy violations, or system dependencies.

Unlike academic institutions, non-academic organizations face direct market pressures, competitive dynamics, and regulatory obligations. These contextual factors heighten the significance of managing technological risk not only to protect assets but also to ensure operational resilience and legal compliance. Effective risk governance thus requires integrating technological foresight with ethical, legal, and social considerations, supported by international frameworks such as ISO 31000.

Classification of technological risks

All types of organizations increasingly rely on technologies and are becoming more digitalized. The eight categories of technological risks were defined through a literature-based classification review. The method of definition involved a review of academic publications – papers, articles, industry reports, etc. which addressed the technological risks associated with digitalization in non-academic environment. Similar groups were grouped under broader conceptual categories. The classification reflects both theoretical insights and practical observations taken from various non-academic sectors and organizations in them.

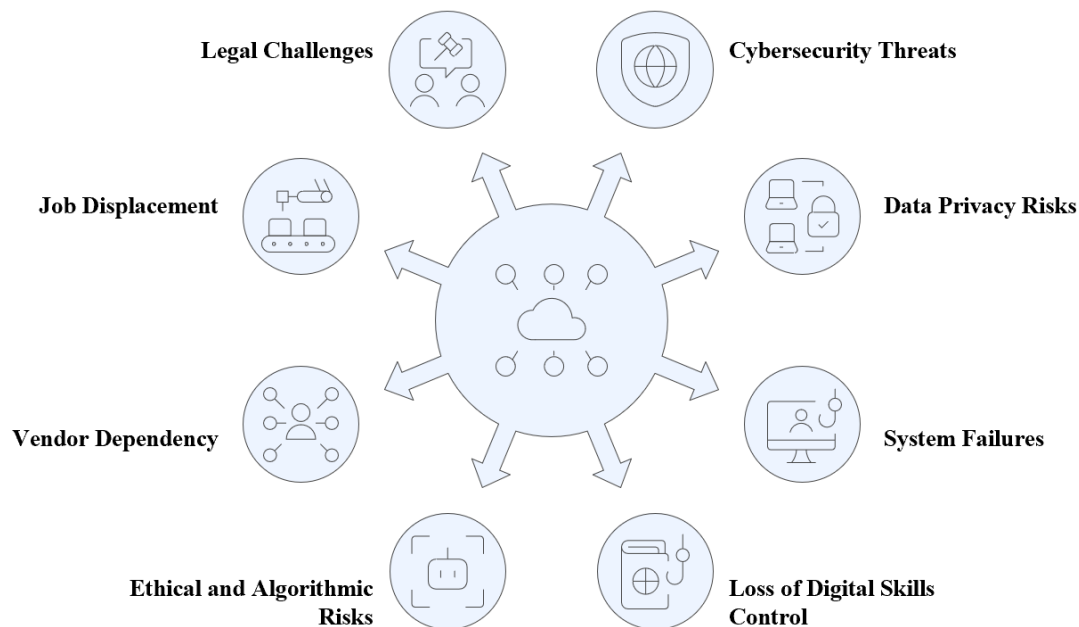


Figure 3 Categories of technological risks

The eight categories of technological risks in non-academia shown on Figure 1 are the following: cybersecurity threats, data privacy risks, system failures and downtime, loss of digital skills control, ethical and algorithmic risks, dependency on vendors and platforms, job displacement, and legal and regulatory challenges.

The classification of technological risks in this paper aligns with the ISO 31000:2018 risk management process, which emphasizes the systematic steps of risk identification, analysis, and evaluation. Each of the eight categories - such as cybersecurity, data privacy, and technological dependency, etc. - represents a key dimension of uncertainty that can affect organizational objectives in non-academic environments. By structuring the classification according to ISO 31000 principles, the framework ensures that risks are recognized in a comprehensive and comparable manner, supporting both academic analysis and practical application within established international risk management standards. [3]

Table 1 illustrates how the eight identified technological risk categories align with the stages of the ISO 31000:2018 risk management process. This integration ensures that the classification not only organizes risks conceptually but also provides a structured approach for identifying, analyzing, and managing them within non-academic environments. [3]

ISO 31000:2018 Process Stage	Description (according to ISO 31000)	Application to the Eight Technological Risk Categories
1.Risk Identification	Recognize and describe risks that could affect the achievement of objectives.	Identify the main sources of technological risk: cybersecurity threats, data privacy risks, system failures and downtime, loss of digital skills control, ethical and algorithmic risks, dependency on vendors and platforms, job displacement, and legal and regulatory changes.
2. Risk Analysis	Understand the nature, sources, likelihood, and potential consequences of each risk.	Analyse how each risk operates and interacts - for instance, the likelihood of cyberattacks , the severity of data breaches , the operational impact of system failures , or the social implications of algorithmic bias and job displacement.
3.Risk Evaluation	Compare analysed risks against defined criteria to determine their significance.	Prioritize which risks require the most urgent response. For example, cybersecurity and data privacy may pose immediate high-impact threats, while dependency on vendors and legal changes may represent longer-term strategic risks.
4.Risk Treatment	Develop and implement measures to mitigate, transfer, accept, or avoid risks.	Apply appropriate mitigation strategies: enhance cybersecurity protocols , strengthen data governance , establish redundancy for system continuity , provide continuous digital skills training , ensure ethical AI oversight , and plan for regulatory compliance and workforce transition.
5.Monitoring and Review	Continually assess the effectiveness of risk controls and adapt to change.	Regularly review and update controls as technologies evolve - for example, reassessing vendor dependencies , AI ethics compliance , and cyber defence mechanisms.
6.Communication and Consultation	Maintain stakeholder engagement and transparency throughout the risk process.	Ensure clear communication across technical teams, management, legal advisors, and employees to build a shared understanding of each technological risk and its mitigation.

Table 4 Technological risks aligned with ISO 31000:2018

- **Cybersecurity Threats**

Definition:

Cybersecurity threats refer to malicious digital actions that compromise the confidentiality, integrity, or availability of data and systems. These include cyberattacks such as phishing, ransomware, and unauthorized access that exploit vulnerabilities in networks and digital infrastructures.

Causes and Contributing Factors:

Key contributors include increased connectivity, inadequate security protocols, outdated software, and limited employee awareness. The rise of remote work in non-academic environment and the Internet of Things (IoT) has expanded attack surfaces, making organizations more exposed. [4] [5]

Impacts and Consequences:

Cyber incidents can cause data breaches, financial losses, operational disruptions, and reputational damage. In critical sectors, they may endanger service continuity or public safety, and often result in legal liabilities and regulatory penalties. [4]

- **Data Privacy Risks**

Definition:

Data privacy risks arise from the improper collection, storage, sharing, or use of personal or sensitive information. They concern both compliance with privacy laws and the ethical handling of digital data.

Causes and Contributing Factors:

Contributors include weak data governance, lack of transparency, and excessive data collection by organizations. Inadequate encryption, third-party data sharing, and poor user consent practices also elevate the risk. [5]

Impacts and Consequences:

Privacy violations lead to loss of individual trust, legal sanctions under regulations such as GDPR, and reputational harm. For organizations, breaches may disrupt operations and erode customer relationships. [5]

- **System Failures and Downtime**

Definition:

System failures and downtime refer to interruptions in digital infrastructure that halt or degrade the performance of critical systems, applications, or services.

Causes and Contributing Factors:

Common causes include software bugs, power outages, hardware malfunctions, or human error. Overreliance on automated systems and lack of redundancy amplify the likelihood and duration of failures. [6]

Impacts and Consequences:

Consequences include productivity loss, revenue reduction, and service unavailability. In sectors like healthcare or finance, system downtime can have severe operational or safety implications. It may also expose weaknesses in business continuity planning. [6]

- **Loss of Digital Skills Control**

Definition:

Loss of digital skills control refers to the decline or mismatch of employees' technological competencies relative to rapidly changing digital tools and systems.

Causes and Contributing Factors:

Rapid digital transformation, insufficient training, and reliance on outsourced IT functions reduce in-house skill retention. Frequent software updates or platform shifts can outpace employee adaptation. [7]

Impacts and Consequences:

A digital skills gap limits innovation increases dependency on vendors, and heightens vulnerability to misuse or error. It can lower productivity, employee confidence, and the overall resilience of the organization. [7]

- **Ethical and Algorithmic Risks**

Definition:

Ethical and algorithmic risks emerge from the design and use of automated systems, especially those employing artificial intelligence or machine learning, which may produce biased or opaque outcomes. [8]

Causes and Contributing Factors:

These risks arise from biased training data, lack of ethical oversight, and insufficient transparency in algorithmic decision-making. Pressure for efficiency can override ethical considerations.

Impacts and Consequences:

Consequences include unfair treatment, discrimination, and loss of public trust. Biased algorithms can damage reputations, violate regulations, and produce social inequalities or ethical controversies.[8]

- **Dependency on Vendors and Platforms**

Definition:

Dependency on vendors and platforms occurs when organizations rely heavily on external technology providers for essential services such as cloud storage, software, or maintenance.

Causes and Contributing Factors:

Outsourcing IT operations, using proprietary software, and engaging in long-term vendor contracts reduce flexibility. Limited interoperability and switching costs reinforce dependency. [9]

Impacts and Consequences:

Vendor dependency can cause disruptions if providers experience failures, price increases, or data breaches. It restricts control over data and technology decisions, posing operational and strategic risks. [9]

- **Job displacement**

Definition:

Job displacement refers to the reduction or transformation of employment caused by automation, artificial intelligence, or other digital technologies that replace or redefine human tasks. [10]

Causes and Contributing Factors:

Advances in robotics, AI, and process automation reduce demand for routine tasks. Lack of upskilling initiatives and resistance to organizational change further amplify the problem. [7]

Impacts and Consequences:

Displacement can lead to workforce instability, social inequality, and resistance to technological change. It may harm employee morale and increase public criticism of digitalization practices. [10]

- **Legal and Regulatory Challenges**

Definition:

Legal and regulatory changes represent the evolving laws and compliance requirements governing technology use, data protection, and digital transactions.

Causes and Contributing Factors:

Rapid innovation often outpaces legislation. Differences in international regulations, unclear compliance standards, and inconsistent enforcement create uncertainty for organizations. [12]

Impacts and Consequences:

Non-compliance can lead to penalties, operational restrictions, or loss of licenses. Constant regulatory evolution increases compliance costs and strategic uncertainty, particularly in data-driven sectors. [12]

Conclusion

Digitalization represents both an opportunity and a challenge for non-academic environments. As organizations adopt new technologies to enhance efficiency and competitiveness, they also become increasingly exposed to a wide range of technological risks. This paper has identified and classified eight major categories of such risks - ranging from cybersecurity and data privacy concerns to ethical, social, and regulatory issues - providing a structured approach to understanding their causes, consequences, and interconnections. Applying the ISO 31000:2018 risk management framework has demonstrated that technological risk management must be continuous, systematic, and multidimensional. Beyond deploying technical safeguards, organizations must foster digital literacy, maintain ethical oversight, and ensure compliance with evolving legal standards. Building resilience requires a balance between innovation and control, where human, organizational, and technological elements are integrated within a comprehensive governance structure. Ultimately, effective management of technological risks in non-academic settings enables organizations to harness the benefits of digital transformation while safeguarding their integrity, reliability, and public trust. As digital ecosystems continue to evolve, proactive risk governance will remain essential for achieving sustainable and secure digital growth.

Acknowledgement

The paper is prepared as a part of the research activities under the project Models and Instruments for Transforming Higher Education System through Transnational Multi-Sector Links - STUDIES – DIG (Project N: 101131544), funded by the European Union.

References

1. ‘Digitization’, Gartner. [Online]. Available: <https://www.gartner.com/en/information-technology/glossary/digitization>
2. ‘Digitalization’, Gartner. [Online]. Available: <https://www.gartner.com/en/information-technology/glossary/digitalization>
3. *ISO 31000:2018 Risk management — Guidelines*. [Online]. Available: <https://www.iso.org/standard/65694.html>
4. S. Saeed, S. Altamimi, N. Alkayyal, E. Alshehri, and D. Alabbad, ‘Digital Transformation and Cybersecurity Challenges for Businesses Resilience: Issues and Recommendations’, vol. 23, no. 15, July 2023.
5. R. Malik, ‘Data Privacy & Cybersecurity: A Governance Imperative’, Mar. 2025, [Online]. Available: <https://www.icsi.edu/media/webmodules/CSJ/March-2025/18.pdf>
6. S. Jagarlamudi, ‘LESSONS LEARNED FROM CRITICAL SYSTEM FAILURES: A TECHNICAL OVERVIEW’, vol. 9, no. 2, pp. 358–367, Oct. 2024, doi: <https://doi.org/10.5281/zenodo.13843406>.
7. L. Patrick Willcocks, ‘Automation, digitalization and the future of work: A critical review’, vol. 3, no. 2, pp. 184–199, Feb. 2024, doi: <https://doi.org/10.1108/JEBDE-09-2023-0018>.
8. M. Bhatia and S. Kumar, ‘Ethical Implications of Biased Algorithms in Business Analytics’, Dec. 2024, [Online]. Available: https://www.researchgate.net/publication/394086631_Ethical_Implications_of_Biased_Algorithms_in_Business_Analytics
9. M. Vedenkannas, ‘Many ways of avoiding dependence on information system vendors’. [Online]. Available: <https://www.vtv.fi/en/blog/many-ways-of-avoiding-dependence-on-information-system-vendors/>
10. C. Peiwen, N. Sulaiman, and S. Zhenglong, ‘The Impact of Artificial Intelligence Application on Job Displacement and Creation: A Systematic Review’, vol. 9, no. 4, pp. 2495–2517, May 2025, doi: <https://dx.doi.org/10.47772/IJRISS.2025.90400185>.
11. Y. Luo, ‘A general framework of digitization risks in international business’, vol. 53, pp. 344–361, May 2021, doi: [10.1057/s41267-021-00448-9](https://doi.org/10.1057/s41267-021-00448-9).
12. J. Paul Onoja, O. Hamza, A. Collins, U. Bright Chibunna, A. Eweja, and A. Ifesinachi Daraojimba, ‘Digital Transformation and Data Governance: Strategies for Regulatory Compliance and Secure AI-Driven Business Operations’, vol. 2, no. 1, pp. 43–55, Jan. 2021, doi: [10.54660/IJFMR.2021.2.1.43-55](https://doi.org/10.54660/IJFMR.2021.2.1.43-55).

TECHNOLOGICAL FEATURES OF MODERN VIDEO-SHARING PLATFORMS

Yavor Tabov¹,

e-mail: jtabov@unwe.bg¹

Abstract

This study explores the fundamental components of modern video-sharing platforms, focusing on metadata, video coding and encoding formats, and user interface elements. Metadata, which provides structured information about video and audio files, enhances discoverability, enables search engine optimization, and supports efficient organization and management of digital content. It is generally classified into descriptive, structural, and administrative categories, each serving specific functions in identification, navigation, and preservation. Video coding transforms visual data into a compact binary format to enable storage and transmission, relying on lossless, perceptual, and lossy compression principles. Video encoding formats define how content is packaged and delivered, affecting streaming performance, storage efficiency, and cross-platform compatibility. Finally, the user interface—comprising the video player, search engine, and recommender system—facilitates intuitive navigation, personalized content interaction, and efficient media delivery. Together, these technical and informational components underpin the effective functioning of contemporary digital media ecosystems, ensuring high-quality content access, discoverability, and user engagement.

Key words: metadata, video coding, video encoding, recommender system

JEL: C88, L86.

Metadata standards and interoperability

Metadata refers to information that describes other data, such as a video file, audio recording, or web page. It plays a key role in improving the discoverability of video content for both individual users and organizations that host videos on their websites or on platforms such as Google, iTunes, YouTube, MySpace, Youku, and Vimeo. Metadata does not alter the content of a video but serves as an informative label. When metadata is created or extracted from a video, it allows for the selection of relevant keywords and descriptions that support search engine optimization (SEO). Search engines use metadata to determine the ranking of content within search results and to generate brief descriptions displayed in those results. This process enhances both the visibility and the quality of traffic directed to websites. Metadata is particularly important for video materials, as video files lack inherent textual elements that can be used for keyword-based searches. While video producers may include keywords in file names, additional descriptive metadata is necessary for computers and search engines to accurately interpret and locate video content [8].

In most metadata standards, there is no explicit mechanism for representing multiple views of a single scene, particularly when these views differ in their temporal duration. In contrast, support for multi-channel audio is generally more advanced due to its longer history of standardization. Although most standards allow the representation of multi-view video content, such representations usually depend on application-specific semantics, and multiple alternative approaches are often possible [1].

In general, metadata are data describing the nature and characteristics of other data.. It can be compared to the index of a book, which contains concise information summarizing the book's content. By reading the index, a reader can gain an overview of the book without going through the entire text, although the index itself does not include detailed information. Similarly, metadata offers a structured

¹ Chief Assistant, PhD. Department of Information Technologies and Communications, Faculty of Applied Informatics and Statistics, University of National and World Economy, Sofia, Bulgaria, 0000-0002-8940-097X

summary that supports data organization and retrieval. In the context of search engines, metadata plays an essential role in locating specific web pages that match a user's query, thereby improving the accuracy and efficiency of information retrieval [8].

To clarify the concept of metadata, it can be divided into several categories that reflect their distinct roles within information systems. The classification of metadata depends on the specific domain, user requirements, and the nature of the represented resources or entities [4]. Metadata is generally classified into three main categories: descriptive, structural, and administrative metadata [2]. Each of these categories plays a specific role in describing, organizing, and managing digital resources, as shown below:

- **Descriptive Metadata** – provides information that enables the identification, discovery, and understanding of a digital resource. Common elements include title, author, and subject. Its main purposes are discovery, presentation, and interoperability.

- **Structural Metadata** – describes the internal organization of a digital resource and the hierarchical relationships among its components. It typically includes information about sequence and position within a hierarchy, supporting navigation and presentation.

- **Administrative Metadata** – contains information necessary for managing the lifecycle of digital resources, including their creation, maintenance, and preservation. It may include attributes such as file type, size, creation date, preservation activities, copyright status, and licensing information. The main functions are interoperability, digital object management, and preservation. This category is further divided into:

- **Technical Metadata** – specifies the technical characteristics and dependencies required for decoding and rendering a digital file.
- **Preservation Metadata** – records information necessary for the long-term maintenance and accessibility of digital materials, such as hardware and software dependencies.
- **Copyright Metadata** – documents details related to intellectual property rights and licensing associated with the digital content [4].

Metadata offers detailed information about the contents of a video library and supports the organization and management of such collections. It consists of descriptive information associated with a file, which may include details such as copyright data, author, keywords, file size, media format, people involved, location, and information related to captions or subtitles. This information is embedded within the file and remains attached when the file is copied, shared, or published. In the context of video, metadata is generally categorized into two main types:

- **Automatic metadata** – generated by devices or software such as cameras. It records technical details like date, time, camera model, aperture, shutter speed, GPS location, and other settings, usually in the Exchangeable Image File Format (EXIF). Most EXIF data cannot be changed after recording.

- **Manual metadata** – created by people to describe the video's content. It includes titles, descriptions, and transcripts that make videos easier to search and categorize. Manually added metadata improves visibility in search engines and supports marketing by helping viewers find relevant content.

Both types of metadata are essential for video sharing and discovery. Platforms like YouTube rely on metadata to classify videos and enhance search and recommendation functions [8]. In general, video metadata is structured information embedded within a digital video file that describes its content and technical characteristics. It can be extracted from different parts of the file's hierarchical structure, such as headers or specific data segments. This metadata may include details related to video, audio, and subtitles, and is essential for organizing, analyzing, or verifying video files across different encoding standards [10].

Video coding principles and encoding formats

In the digital media landscape, video data must be efficiently represented, transmitted, and stored without significant loss of quality. This efficiency is achieved through video encoding, a process that compresses raw footage while maintaining visual fidelity and ensuring compatibility across devices and platforms. To support this process, various encoding standards and formats have been developed, each balancing compression performance with image quality and playback flexibility.

Video coding transforms visual data into a compact binary format to enable efficient storage and transmission. Over the past two decades, the main international video coding standards have been jointly developed by two organizations: ISO/IEC's Moving Picture Experts Group (MPEG) and ITU-T's Video Coding Experts Group (VCEG). These collaborations produced a series of widely adopted standards: H.262/MPEG-2 (1995), H.264/AVC (2003), H.265/HEVC (2013), and H.266/VVC (2020). All these standards rely on three core principles:

- Lossless compression, which removes spatial and temporal redundancy through techniques such as motion estimation and entropy coding.
- Perceptual compression, which discards information not noticeable to the human visual system.
- Lossy compression, achieved through quantization to further reduce data size [3].

While video coding establishes the theoretical framework and algorithms governing video compression, video encoding represents the practical implementation of these principles. It determines how a video stream is processed, packaged, and stored in specific file formats for playback or distribution.

Video encoding format refers to the technical structure and compression method used to represent a film or video file. It defines the file's quality, adaptability, and usability across different platforms and devices. Numerous encoding formats exist, including AVI, DV, HDV, MOV, WMV, MPEG-1, MPEG-2, MPEG-4, MJPEG, MKV, and 4K, among others.

Raw video recordings occupy substantial storage space, often several gigabytes per minute depending on resolution and frame rate. Therefore, compression is essential to reduce file size for efficient storage, distribution, or streaming. Different compression methods vary in efficiency, quality retention, and compatibility, influencing the choice of format based on production needs and technical requirements. Video encoding formats determine how a film is stored, accessed, and distributed—whether on physical media such as VHS or VCD, on digital storage devices, or via online streaming platforms. The format also affects the level of compression possible, the resulting image quality, and the ease of modification or playback. Over time, multiple standards have been developed to address evolving technological demands and media applications [6].

Within the technological framework of modern video-sharing platforms, video coding and video encoding formats play a central role in optimizing content delivery and playback. Platforms rely on standardized coding algorithms such as H.264/AVC and H.265/HEVC to compress and transcode uploaded videos into multiple formats and bitrates. This enables adaptive streaming, ensuring consistent quality across devices and varying network conditions. The efficiency of the selected coding standard and encoding format directly influences streaming performance, storage requirements, and overall user experience, making them essential components of contemporary digital media infrastructure.

Fig. 1 illustrates the principles of video coding and the corresponding encoding formats.

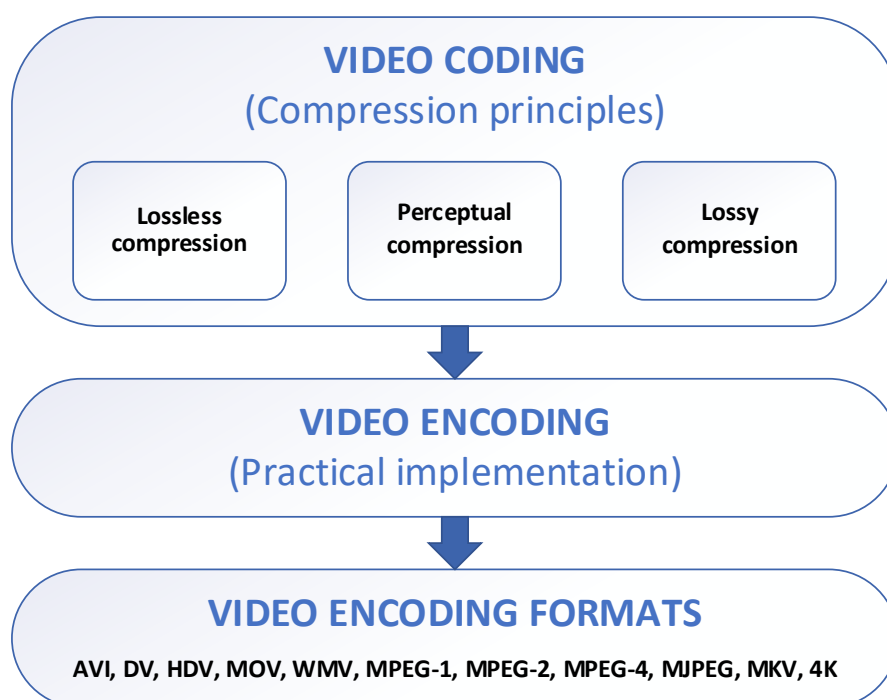


Figure 1: Overview of video coding principles and encoding formats

User interface main components of video-sharing platforms

In the context of modern video-sharing platforms, several components of the user interface play a decisive role in shaping user engagement and content accessibility. The main technical components of the User Interface (UI) in modern video-sharing platforms are designed to facilitate efficient media delivery, intuitive navigation, and personalized content interaction. These components integrate playback technologies, search functionalities, and recommendation algorithms to create a cohesive and responsive user experience across devices. In such user interfaces, three core elements can be identified:

- **Video Player.** As a central module of video-sharing platforms, the video player provides users with direct access and control over audiovisual content. The adoption of HTML5 eliminated the need for external plug-ins or proprietary software by introducing native support for open multimedia formats in web browsers. This standard enables smooth streaming, adaptive playback, and consistent performance across devices such as smartphones, tablets, computers, and smart TVs. Through HTML5 integration, video delivery has become more stable, secure, and widely accessible, establishing a technological foundation for modern online media playback [9].

- **Search Engine.** Search functionality allows users to efficiently discover videos or channels from vast platform databases. Technically, a search engine retrieves information based on user-specified keywords or phrases, using indexed metadata—such as titles, descriptions, and tags—to return relevant results. In platforms like YouTube and Vimeo, these engines are optimized for multimedia retrieval and often employ personalization algorithms to adapt search results to individual viewing patterns and interests [7].

- **Recommender System.** Recommendation mechanisms personalize content access by suggesting videos that match a viewer’s preferences and past interactions. Evolving from earlier domains such as book or movie recommendation, these systems now underpin major platforms like YouTube and Netflix. By analyzing user behavior and video attributes, they predict content likely to sustain engagement. Beyond improving user experience, recommendation algorithms serve a strategic role in increasing retention, visibility, and viewing time, making them a defining element of contemporary video-sharing ecosystems [5].

Conclusion

From the conducted research, the following key points can be summarized:

- Metadata represents information that provides details about other data, such as video files, audio recordings, or web pages. It is classified into three main categories: Descriptive Metadata, Structural Metadata, and Administrative Metadata. Each category serves a distinct purpose in organizing, identifying, and managing digital resources.
- Video coding involves converting visual data into a compact binary format to enable efficient storage and transmission. It is based on the principles of lossless, perceptual, and lossy compression,
- In modern video-sharing platforms, video coding and encoding formats play a vital role in ensuring efficient content delivery and smooth playback. Their effectiveness directly impacts streaming performance, storage efficiency, and user satisfaction, making them key components of the digital media ecosystem.
- Key components of the user interface in modern video-sharing platforms are the video player, search engine, and recommender system. These elements collectively ensure efficient content delivery, intuitive navigation, and personalized user experiences.

Acknowledgement

This work was financially supported by the UNWE Research Programme (Research Grant No. 22/2024/A).

References

1. Bailer, W., & Höffernig, M. (2009). Metadata for creation and distribution of multi-view video content. In Proceedings of the 10th International Workshop of the Multimedia Metadata Community on Semantic Multimedia Database Technologies (SeMuDaTe'09) (Vol. 539, Paper 4). CEUR-WS.org. https://ceur-ws.org/Vol-539/paper_4.pdf
2. Calibo, D. I., & Niguidula, J. D. (2019). Metadata extraction analysis: A review of video data in effect to social media compression. JOIV: International Journal on Informatics Visualization, 3(1), 54-58.
3. Casella, A. (2025). Evolution of video codecs: a comparative analysis of HEVC, VVC, AV1, and AVM (Doctoral dissertation, Politecnico di Torino).
4. Formenton, D., & de Souza Gracioso, L. (2022). Metadata standards in web archiving technological resources for ensuring the digital preservation of archived websites/Padroes de metadados no arquivamento da web: recursos tecnologicos para a garantia da preservacao digital de websites arquivados. Revista Digital de Biblioteconomia e Ciencia da Informacao, 20(1), NA-NA.
5. Lubos, S., Felfernig, A., & Tautschnig, M. (2023). An overview of video recommender systems: state-of-the-art and research issues. Frontiers in big Data, 6, 1281614.
6. Onyejelem, T. E. Film Production: Video Encoding Formats, Theories and Aesthetics. University of Lagos Press and Bookshop Ltd.
7. Reiter, M., Cupka, O., & Miklosik, A. (2021). Search Engine Optimization of Video Content. Research Gate, November.
8. Sathish Kumar, S., & Pushpa, P. (2019). Preparing metadata for making video file searchable on online. JETIR, 6(3), March 2019.

9. Shende, V. S. (2024). HTML5 in Web Development. Recent Advancements in Science and Technology, 206.
10. Xiang, Z., Horváth, J., Baireddy, S., Bestagini, P., Tubaro, S., & Delp, E. J. (2021). Forensic analysis of video files using metadata. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (pp. 1042-1051).

ПОДОБРЯВАНЕ НА ОБУЧЕНИЕТО И ПОВИШАВАНЕ НА КВАЛИФИКАЦИЯТА ЧРЕЗ СИМУЛАЦИИ С РАЗШИРЕНА РЕАЛНОСТ

The Enhancement of Learning and Upskills through Extended Reality Simulations

Kaloyan Dimitrov¹

e-mail: kdimitrov@unwe.bg¹

Абстракт

Интегрирането на технологиите за разширена реалност (XR) оказва трансформиращо влияние върху съвременните среди за обучение и професионално обучение. Настоящото проучване разглежда ефективността на XR симулациите за подобряване на резултатите от обучението и повишаване на квалификацията в различни образователни и професионални контексти. Целта на тази статия е да установи основните предимства на използването на XR симулации в индустриалния контекст, с особен акцент върху областите на обучението и повишаването на квалификацията. Като част от изследването беше проведен систематичен преглед на литературата, за да се идентифицират и теоретизират посочените ползи и приложения на XR симулациите. Тези констатации подчертават трансформационния потенциал на разширената реалност като катализатор за иновативна педагогика и развитие на работната сила. Настоящото проучване представя аргументи за увеличаване на инвестициите и изследванията в областта на XR инструментите. В заключението се твърди, че тази технология може да улесни моделите на обучение и учене, като по този начин се установи по-добра връзка между теоретичните концепции и тяхното практическо приложение в реалния свят. Следователно XR симулациите имат потенциал да служат като трансформативен инструмент в областта на обучението и повишаването на квалификацията.

Abstract

The integration of Extended Reality (XR) technologies is having a transformative effect on contemporary learning and professional training environments. The present study investigates the effectiveness of XR simulations in enhancing learning outcomes and upskilling across diverse educational and vocational contexts. The aim of this paper is to ascertain the primary advantages of utilizing XR simulation within industry contexts, with a particular focus on the domains of learning and upskilling. A systematic literature review was conducted as part of the research to identify and theorize over the stated benefits and applications from XR simulations. These findings emphasize the transformative potential of Extended Reality as a catalyst for innovative pedagogy and workforce development. This study posits the argument for augmented investment and research in the domain of XR tools. The conclusion argues that this technology can facilitate training and learning models, thereby establishing a better connection between theoretical concepts and their practical applications in the real world. Hence, XR simulations have the potential to serve as a transformative instrument in the domains of learning and upskilling.

Key words: immersive technologies, extended reality, immersive education, immersive learning, immersive training, skills.

JEL: A22, I21, I25, O31, O35.

¹ Assoc. Prof. Dr., Department Industrial Business, University of National and World Economy – Sofia, Bulgaria, ORCID: 0009-0006-3752-3413, e-mail: kdimitrov@unwe.bg

Introduction

A multitude of industries are undergoing rapid transformation due to a combination of geopolitical shifts, climate goals and demographic changes, thereby driving a growing demand for skilled labour. Conventional training methodologies are found to be deficient in effectively integrating practical and theoretical learning, thus underscoring the necessity for innovative, scalable, and immersive training solutions to meet the future demands of the workforce.

Advancements in Extended Reality (XR) technology have created an alternative to traditional training, offering the potential of safe, efficient and scalable training with a high degree of realism and practical learning. Despite recent technological advances and reduced hardware costs, there has been a paucity of large-scale industrial implementations of XR simulations.

The aim of this paper is to ascertain the primary advantages of utilizing XR simulation within industrial contexts, with a particular focus on the domains of learning and upskilling. A systematic literature review was conducted as part of the research to identify and theorize the benefits and applications from XR simulations.

As demonstrated in the extant literature on XR simulation, the design of the XR training environment is heavily dependent on the applied learning style and the industry use case.

Conceptual Foundations of Extended Reality Technology: Brief overview

The term “extended reality” is an umbrella term for “virtual Reality” (VR), “augmented reality” (AR) and “mixed reality” (MR). These immersive technologies blend the real and digital worlds, with the objective of enhancing productivity. First, VR is a technology that employs a headset to create a fully immersive experience, transporting the user into a simulated environment that is presented in three dimensions. Second, AR can be defined as the overlaying of computer-generated graphics onto a real-world view. Third, MR encompasses the ability to observe and engage with holograms that are integrated within one's physical environment [1].

The Emerging Immersive Technology is no longer a futuristic concept. The advent of novel digital technologies, including but not limited to Extended reality, artificial intelligence (AI), and the Metaverse, has led to a marked increase in their utilization within the industry. This development is indicative of a broader shift towards the integration of technology into various facets of production and consumer behavior, with the objective of optimizing efficiency and enhancing the shopping experience. This technology has been demonstrated to have the capacity to enhance contemporary business performance. Forbes cited ABI Research's forecast that the XR market will grow from \$44.7 billion in 2024 to \$299.3 billion by 2030 [2].

Extended reality technology plays a key role in addressing critical business challenges. The use of XR solutions for core business requirements occurs through the provision of effective actions in four key areas [3]:

- *Learning and Development*: This approach to education is characterized by its immersive and hands-on nature, with a focus on the acquisition of practical skills in a risk-free environment;
- *Operations and Workflow*: The utilization of real-time guidance and simulation has been demonstrated to enhance the efficiency of tasks;
- *Sales and Marketing*: Engaging demonstrations and virtual showcases have been shown to be effective in exciting potential buyers;
- *Cost Reduction and Efficiency*: A reduction in errors, travel, and wasted resources will result in savings.

Extended Reality Simulations as a Strategic L&D Solution

In the contemporary business context, the industry is undergoing a significant transformation with respect to the manner in which its human capital is being developed and capacitated. Conventional training and development methodologies are frequently characterized by their protracted nature, substantial financial expense and challenging scalability. Digital tools are currently transforming the landscape of workforce training, enhancing the efficiency, accessibility, and safety of learning processes.

Table 1: Reality-Virtuality Continuum of XR Technologies

	Type of Reality	Description	Best Use Cases
1	Virtual Reality	Fully immersive digital environment	Complex simulations, high-risk training
2	Augmented Reality	Digital overlay on real world	On-the-job guidance, equipment maintenance
3	Mixed Reality	Interaction between real and digital objects	Collaborative training, design reviews

Source: Author's summary along Milgram's continuum

In the contemporary business context, organizations are increasingly utilizing technological solutions to enhance the competencies of their workforce. This technological adoption encompasses a wide spectrum of applications, ranging from extended reality simulations to AI-driven skill assessments. The primary objective of these initiatives is to facilitate the rapid and effective upskilling of employees, thereby ensuring their readiness to face the challenges and opportunities of the modern business environment.

- Virtual Reality Training

The employment of VR simulations enables immersive, hands-on training experiences, circumventing the risks associated with real-world scenarios. Employees have the opportunity to refine their skills in operating technologies and techniques, respond to hazardous situations, and develop competencies in a controlled, risk-free environment [4]. The VR utilization has been demonstrated to facilitate the emulation of perilous on-site scenarios. This, in turn, has been shown to enhance the efficacy of safety training and facilitate the retention of knowledge.

- Augmented Reality for On-Site Learning

AR technology has been developed to superimpose digital instructions over physical equipment, thereby enabling personnel to access real-time guidance while undertaking tasks. This enhancement in accuracy is particularly beneficial in complex work activities, where it has been shown to reduce errors [4]. In summary, it is recommended that employees be equipped with augmented reality headsets with a view to enhancing on-site learning and productivity.

- AI-Powered Skill Assessment

The utilization of artificial intelligence in the evaluation of employee performance entails the implementation of digital testing methodologies and real-time monitoring systems. The employment of these technologies facilitates the identification of skill gaps within the workforce, thus enabling the recommendation of targeted training programmes. This enables business organizations to personalize training, thereby ensuring that employees develop the requisite skills for specific tasks [5]. In summary, it is recommended that organizations implement AI-driven analytics to optimize their workforce development strategies.

- Mobile Learning Platforms

Mobile-based training facilitates on-the-go learning, thereby enabling employees to access educational content from any location. This is a particularly beneficial feature for employees working at remote locations or those with busy schedules, as it allows them to learn at their own pace [6]. In summary, the utilization of microlearning modules within mobile applications facilitates the delivery of concise and efficacious training sessions.

- Gamification in Training Programs

The incorporation of game-like challenges, leaderboards, and rewards has been demonstrated to enhance the engagement and competitiveness of training, thereby increasing participation and motivation. This approach has been demonstrated to enhance skill mastery while simultaneously rendering the learning process more enjoyable [7]. In summary, it is recommended that rewards-based gamification techniques be employed in order to enhance worker engagement.

- Digital Certification and Compliance Tracking

Automated systems have been developed to track worker certifications, thus ensuring compliance with safety and industry regulations. Digital platforms have been shown to reduce the amount of paperwork involved in the process, and to assist companies in the maintenance of records pertaining to mandatory training requirements [8]. In summary, the utilization of AI-powered compliance tracking tools facilitates the automation of certification renewals, thereby minimizing administrative workload.

The Potential Benefits of Learning and Upskills through XR Simulations

As demonstrated by extant research and industry evidence, the most significant benefits of XR training are those that provide a clear, measurable advantage over traditional methods such as classroom learning, e-learning modules or physical simulation [9]. The following figure outlines the most relevant benefits associated with XR training.



Source: Author's summary

Figure 1: Most relevant benefits of XR Training

The aforementioned benefits are the driving factors behind the rapid adoption of XR by industries as diverse as healthcare, manufacturing, corporate retail, the military and many others. The ability to provide a quantifiable solution to the most critical training challenges is the primary incentive for these industries.

A synthesis of current academic research and industry implementation has identified several core areas in which XR technology can be employed to enhance learning and upskilling [10]. This technology is moving from a niche tool to a mainstream pedagogical solution, particularly where traditional training

is too costly, dangerous, or logistically difficult to scale [11, 12]. The following table provides a synopsis of the identified key application domains.

Table 2: Key Application Domains of Learning and Upskilling via Extended Reality (XR)

Benefit Category	Key Takeaways & Impact	Industry Implementation
Complex Psychomotor & High-Risk Skills	XR is leveraging its ability to create “safe-to-fail” environments for high-stakes professions. Engaging with the concept of embodied cognition – the notion that the human brain acquires knowledge through the process of physical interaction – enables trainees to develop procedural “muscle memory” in a risk-free environment. This is facilitated by the use of virtual reality technology.	Healthcare; Aviation & Aerospace; Manufacturing; Emergency Services.
Technical, Procedural & Maintenance Skills	This area primarily uses augmented and mixed reality to overlay digital information onto the real world, providing on-the-job performance support rather than off-site simulation. The utilization of “just-in-time” visual prompts and step-by-step instructions by AR systems has been demonstrated to engender a reduction in the mental effort (extraneous cognitive load) demanded for the recollection of intricate procedures. This, in turn, has been shown to result in a significant decline in error rates and expedited task completion times.	Field Service; Maintenance; Manufacturing; Assembly; Logistics; Tourism.
Interpersonal & "Soft" Skill Development	XR technology is used to simulate nuanced, high-stakes human interactions, which are notoriously difficult to practice. These simulations frequently incorporate AI-driven virtual humans that are capable of reacting to the learner's tone of voice and word choice; highlights VR's effectiveness for skills that require emotional regulation and behavioral change. Studies on empathy training (e.g., having a doctor experience a scenario from the patient's perspective) show marked increases in patient-centric communication.	Leadership; Business Management; Sales; Customer Service; Diversity, Equity & Inclusion.
Foundational Knowledge & Conceptual Learning	The utilization of XR technology facilitates the conceptualization of abstract concepts and the standardization of fundamental training, including the orientation of new employees.	Employee Onboarding; Product Knowledge; Design & Visualization.

Source: Author's summary

Conclusion

In summary, Extended Reality simulations, incorporating virtual, augmented, and mixed reality, have demonstrated considerable potential in enhancing learning and upskilling across diverse domains, particularly in the fields of healthcare and education. XR provides immersive and interactive environments that facilitate the understanding of complex concepts, improve skill acquisition, and increase learner satisfaction. Nevertheless, the efficacy of XR is contingent upon its integration into well-structured learning frameworks and the overcoming of challenges such as accessibility and cost.

While XR offers promising opportunities for enhancing learning and upskilling, its effectiveness is not universally guaranteed. The success of XR applications is contingent upon the context, the target audience, and the specific learning objectives. Further research is required to comprehensively ascertain the long-term impact of XR on learning outcomes and to address the challenges associated with its implementation.

It can thus be concluded that XR is not merely a new method of delivering content; rather, it signifies a fundamental shift towards experiential, data-driven, and scalable learning. It has been demonstrated that this medium is more effective than almost any other, thus rendering it a powerful tool for the future of education and professional development across all industries.

This study concludes that Extended Reality has the potential to serve as a transformative instrument in the domains of upskilling and education. It is argued that this technology can facilitate experiential

learning models, thereby establishing a connection between theoretical concepts and their practical applications in the real world.

Acknowledge

The author acknowledges with gratitude funding provided by the Bulgarian National Research Fund under administrative contract KP-06-H65/5 of 12.12.2022.

References

1. Denizova, V. (2024). Virtual Reality (VR) and Augmented reality (AR) as tools for tourism development: Opportunities and challenges. *Innovative Information Technologies for Economy Digitalization (IITED)*, (1), 282-292.
2. Walter Van Uytven, V.W. (2024). The State and Challenges of Enterprise XR: Hype and Reality. Available at: <https://www.forbes.com/councils/forbestechcouncil/2024/10/18/the-state-and-challenges-of-enterprise-xr-hype-and-reality/>
3. See, Z. S., Ledger, S., Goodman, L., Matthews, B., Jones, D., Fealy, S., & Amin, M. (2023). Playable experiences through technologies: Opportunities and challenges for teaching simulation learning and extended reality solution creation. *The Journal of Information Technology Education: Innovations in Practice*, 22, 67-90.
4. Evans-Uzosike, I. O., Okatta, C. G., Otokiti, B. O., Ejike, O. G., & Kufile, O. T. (2022). Extended Reality in Human Capital Development: A Review of VR/AR-Based Immersive Learning Architectures for Enterprise-Scale Employee Training. *International Scientific Refereed Research Journal*, 5(5), 111-124.
5. Milev, P. (2023). The Role of Data Visualization in Enhancing Textual Analysis. *Trakia Journal of Sciences*, 21(1), 248-253.
6. Kolyandov, S. (2021). The rising popularity of digital transaction platforms. *Trakia Journal of Sciences*, 19(1), 122-129.
7. Hancko, D., Majlingova, A., & Kačíková, D. (2025). Integrating Virtual Reality, Augmented Reality, Mixed Reality, Extended Reality, and Simulation-Based Systems into Fire and Rescue Service Training: Current Practices and Future Directions. *Fire*, 8(6), 228.
8. Becattini, N., Cascini, G., & Rodriguez, M. C. (20257). Learning in eXtended Reality-A. In *Design Tools and Methods in Industrial Engineering IV: Proceedings of the Fourth International Conference on Design Tools and Methods in Industrial Engineering*, ADM 2024, September 11–13, 2024, Palermo, Italy, Volume 2 (p. 411). Springer Nature.
9. Chaudhary, P., Raut, N., Patil, H., & Kate, N. (2024). Up-skilling in fashion retail. *Transforming Education with Virtual Reality*, 323-336.
10. Padmaja, V., & Mukul, K. (2021). Upskilling and reskilling in the digital age: the way forward for higher educational institutions. In *Transforming higher education through digitalization* (pp. 253-275). CRC Press.
11. Boland, B. (2023). Extended reality vocational training's ability to improve soft skills development and increase equity in the workforce. *AI, Computer Science and Robotics Technology*, (13).
12. Chotrov, D., Uzunova, Z., Maleshkov, S., Stoyanova, P., Stoyanova, T., Lindgren, P., ... & Poulkov, V. (March). Virtual Collaboration in Mixed Reality for Enhancing

Business Model Innovation. In *2023 Joint International Conference on Digital Arts, Media and Technology with ECTI Northern Section Conference on Electrical, Electronics, Computer and Telecommunications Engineering (ECTI DAMT & NCON)* (pp. 528-532). IEEE.

ОТВОРЕНИТЕ ДАННИ В КОНТЕКСТА НА УСТОЙЧИВИТЕ ХРАНИТЕЛНИ СИСТЕМИ

Open Data in the Context of Sustainable Food Systems

Асен Божилов¹

e-mail: a.bozhikov@uni-svishtov.bg¹

Абстракт

Докладът изследва значението на отворените данни за подобряване на прозрачността, информираността и ефективността на устойчивите хранителни системи. Отворените данни се разглеждат като стратегически ресурс, който подпомага научните изследвания, иновациите и обществените политики, насочени към устойчиво производство и потребление на храни. Те имат потенциала да бъдат катализатор за по-устойчиви хранителни практики, по-добро управление на ресурсите и засилено сътрудничество между научните, институционалните и гражданските общности. Разгледани са различни източници на отворени данни и тяхната роля за проследяване на социални, икономически и екологични индикатори, свързани с хранителните системи.

Abstract

The report discusses the importance of open data for improving transparency, awareness, and effectiveness of sustainable food systems. Open data is seen as a strategic resource that support scientific research, innovation, and public policies aimed at sustainable food production and consumption. It has the potential to be a catalyst for more sustainable food practices, better resource management, and enhanced collaboration between scientific, governmental, and civic communities. Various sources of open data and their role in tracking social, economic, and environmental indicators related to food systems are examined.

Ключови думи: отворени данни, платформи за отворени данни, прозрачност, устойчиви хранителни системи.

JEL: D83, Q57

Значение и характеристики на отворените данни

Развитието на информационните и комуникационни технологии (ИКТ) и тяхното използване в организации от различен тип и сфера на функциониране неминуемо води след себе си до оставяне на дигитални следи и генериране на огромни количества данни. Данните имат потенциала да трансформират начина на вземане на решения, както и начина на управление и справяне със социални и икономически феномени, появяващи се в съвременното общество. Това важи с изключително голяма сила за хранителните системи, чието нормално функциониране е заплашено от непрекъснато нарастващото население на планетата, намаляването на обработваемите площи земя и настъпващите климатични промени (Ling, Hu, Li, Zhao, & Ye, 2025). В тази връзка дигиталната революция с нейните задвижвани от данни технологии може да осигури възможност за оптимизиране на използването на ресурсите, подобряване на добивите от културите и изграждане на по-устойчиви системи за производство и потребление на храни, с което да се осигури плавното преминаване към устойчиви хранителни системи и постигане на продоволствена сигурност. Ако в голямата си част тези данни се предоставят под формата на

¹ Head Assistant, PhD. Department “Business Informatics”, D. A. Tsenov Academy of Economics,
ORCID: 0000-0002-0680-2800

отворени данни, това би довело до насърчаване на прозрачността, стимулиране на иновациите и повишаване на обществената ползва (Wirtz, Weyerer, & Rösch, 2022). Редица изследвания доказват ползността на отворените данни за намаляване на разходите за набавяне на информация и подобряване на ефективността на работа на отделите за научноизследователска и развойна дейност в организациите (Yan, Xing, Chen, & Zhao, 2023), а също така и за стимулиране на зелените иновации чрез повишена прозрачност на данните за околната среда и създаване на по-отговорни политики в тази област (Peng & Xiao, 2024).

Отворените данни се определят като: „цифрови данни, които са достъпни с техническите и правните характеристики, необходими за свободното им използване, повторна употреба и разпространение от всеки, по всяко време и навсякъде“ (Open Data Charter, 2015). Основните характеристики на отворените данни включват:

- *свободен достъп* – няма технически или финансови бариери (напр. платени лицензи или регистрации);
- *машинно четими* – данните са в структуриран формат (CSV, JSON, XML, RDF), който позволява автоматична обработка;
- *отворен лиценз* – ясно дефинирани условия за използване (напр. Creative Commons BY, Open Data Commons ODbL);
- *метаданни и документация* – придружаваща информация, която описва структурата, съдържанието и контекста на данните.

Отворените данни се базират на ключови концепции, като отвореност и съвместно сътрудничество. Чрез тях може значително да се подобрят процесите по вземане на решения и съвместно създаване на знания. Подобни данни могат да бъдат открити за различни тематични области – околна среда, селско стопанство, демографска статистика, заетост на населението, образование и др. Тяхната подготовка обаче изисква спазването на различни принципи, запазване на личната неприкосновеност и наличие на подходяща дигитална инфраструктура.

Отворените данни обикновено са достъпни за потребителите чрез портали или платформи за отворени данни. В тях, на базата на задаване на различни филтри, като критерии за търсене, потребителят може да ограничи набора от резултатни данни. Често порталите за отворени данни не са атрактивни за техните ползватели и в тази връзка някои автори предлагат прилагане на елементи на геймификация в тези платформи, за да се постигне по-високо ниво на ангажираност (Satrio, Kleiman, & Janssen, 2026).

Повечето държавни, частни, неправителствени, научни и други организации, които предоставят отворени данни следват FAIR (Findable, Accessible, Interoperable, Reusable) ръководните принципи при тяхното подготвяне. Те представляват набор от насоки и съображения при проучване, споделяне и повторна употреба в среди за публикуване на данни и дигитални ресурси (Wilkinson, Dumontier, & Aalbersberg, 2016), като фокуса е поставен върху това данните да са:

- *откриваеми (Findable)* - да се намират лесно чрез уникални идентификатори и подробни метаданни, както от хора, така и от машини;
- *достъпни (Accessible)* - да се съхраняват на сигурни места, които гарантират, че ще бъдат достъпни за бъдещи използвания;
- *оперативно съвместими (Interoperable)* - да се организират по такъв начин, че да могат лесно да се комбинират и обработват с други данни или в други системи;
- *повторно използваеми (Reusable)* - да се документират, така че да може да се разбере как и защо са били събрани.

Следва да се спомене, че спазването на FAIR принципите не се отнася само до отворените данни, а по принцип за данните. От своя страна важно значение по отношение на отворените данни има 5 звездния модел, предложен през 2010 година от създателя на World Wide Web Сър Тим Бърнърс-Лий. Целта на модела е да помогне на организациите да разберат и подобрят процеса

на публикуване на отворени данни в уеб среда (5 Star open data, 2012). Той разглежда 5 нива на отвореност и полезност при публикуване на данните в уеб, като всяка ниво надгражда постигнатото на предходното и по този начин последователно се достига до наистина отворени и оперативно съвместими данни. Първите три нива (1-3) се фокусират върху достъпността на данните и предоставянето им в отворени, машинно четими формати. Последните две нива (4-5) наблягат върху свързаността и семантичната стойност на данните. Достигането на петото ниво означава, че данните са се превърнали в част от семантичната мрежа, в която информацията от различни източници може да бъде свързана и разбрана от машините.

Разгледаните принципи и модели очертават техническите и етичните измерения на отворените данни, но тяхното реално прилагане е неразривно свързано с правната и политическата рамка, в която функционират организациите или общностите, които създават и подготвят за публикуване съответните данни. Следователно, за да се разбере контекстът на отворените данни в хранителните системи, е необходимо да се анализира нормативната рамка, която определя границите между прозрачност, достъпност и защита на личната информация.

Политики и инициативи за отворени данни

Отворените данни са изключително полезен инструмент за постигане на прозрачност и ефикасност, но от съществено значение е те да се използват отговорно и етично. В тази връзка, за да се осигури неприкосновеност на личните данни, през годините се разработиха регулаторни мерки, чрез които да се предпазва изтичането на подобна информация в масиви с отворени данни. Като такива може да посочим действащите в момента Общ регламент за защита на данните (General Data Protection Regulation, GDPR) в Европейския съюз (ЕС), както и Калифорнийски закон за поверителност на потребителите (California Consumer Privacy Act, CCPA) в САЩ. Те предпазват съответно гражданите на ЕС и потребителите в Калифорния, като ограничават споделянето на лични данни в масиви с отворени данни. По този начин чрез прилагане на анонимизация, агрегиране и минимизация на данните се избягват потенциални рискове за излагане на показ на личността на потребителя или неговите онлайн действия. Друг важен документ в ЕС, касаещ преизползването на данни е Директивата 2019/1024 относно отворените данни и повторното използване на информацията от общественния сектор. Тя указва, че данни, които се създават или събират от общественния сектор (метеорологични, географски, статистически, икономически) трябва да са отворени, достъпни и преизползваеми с изключение на случаите, когато те включват лични данни или носят някакво ниво на конфиденциалност (Европейски Съвет, 2019).

На международно равнище прилагането на принципите на отворените данни в сферата на селското стопанство и хранителните системи се осъществява чрез редица стратегически инициативи и партньорства. Те целят не само да осигурят свободен достъп до информация, но и да изградят устойчиви екосистеми от данни, които подпомагат научните изследвания, иновациите и разработването на релевантни политики. В този контекст особено показателни са усилията на глобални и национални организации, които развиват политики, инфраструктури и стандарти за споделяне и използване на данни. Сред най-значимите примери са инициативите Global Open Data for Agriculture and Nutrition (GODAN) и Food Data Transparency Partnership (FDTP), както и стратегическата рамка на Европейския съюз „От фермата до вилницата“ (Farm to Fork), които поставят отворените данни в центъра на трансформацията към устойчиви хранителни системи.

GODAN е най-известната международна инициатива, която промотира създаването на отворени данни за селско стопанство и храни (Transforming agriculture through Open Data, 2025). Създадена е през 2013 година, като резултат от новосформирания в G8 Алианс за продоволствена сигурност и хранене. GODAN представлява глобална организация, която се състои от над 1000 члена, в това число правителствени и неправителствени организации, научни институции и компании от частния сектор (Sánchez-Padial, 2017). Фокусът на работата на GODAN е върху създаване на необходимите политики в областта на отворените данни на най-високо ниво, както и

подпомагане на обществени и частни организации по отношение на подготвяне и предоставяне на отворени данни.

Друга подобна инициатива, която се реализира в Обединеното кралство е FDTP. Тя обхваща множество заинтересовани страни – различни министерства, индустриални организации, академични институции и гражданското общество. Целта на FDTP е да се подобри достъпността, качеството и съпоставимостта на данните във веригата за доставки на храни, за да се осигури производството и продажбата на по-екологично устойчиви и по-здравословни храни и напитки (Food Data Transparency Partnership, 2025). Въпреки че FDTP не е изцяло инициатива за отворени данни, нейната същност се преплита значително с принципите за отворени данни. Предвижда се постигане на по-голяма прозрачност чрез предоставяне на значителен обем от данни като отворени или с определен лиценз за свободно използване. Набляга се на разработването на методи, метрики и източници на данни, които да позволяват данните да са машинно четими и съпоставими между различни организации. Говори се за изграждане на инфраструктура, която да гарантира използваемостта на тези данни т.е. те да бъдат структурирани, актуални и стандартизирани.

В ЕС беше приета стратегията „От фермата до вилницата“, която се явява водеща политика в областта на хранителната система в рамките на Европейската Зелена сделка, чиято цел е да превърне Европа в първия климатично неутрален континент до 2050 г. Стратегията „От фермата до вилницата“ има за цел да създаде хранителна система, която е справедлива, здравословна и екологична, като отчита дълбокото въздействие, което селското стопанство и консумацията на храни оказват върху климата, биологичното разнообразие, общественото здраве и икономическата устойчивост (Таков, 2024). В основата на тази политика е залегнало разбирането, че устойчивостта в хранителните системи не може да бъде постигната без вземане на решения, основани на данни, прозрачност и използване на иновативни цифрови технологии. Възприемането на стратегическото значение на отворените данни от страна на ЕС, в съчетание с постижения като Общото европейско пространство за селскостопански данни (CEADS), прави решенията, основани на данни, централни както за разработването на политики, така и за трансформацията на територията. Отворените данни са основен фактор за реализиране на практически всички аспекти на стратегията: подкрепя на прозрачността и основаното на факти политическо вземане на решения, насърчаване на цифровите иновации, овластяване на потребителите и гражданското общество и осигуряване на ефективен мониторинг и адаптивно управление на сложни системи.

Приложение на отворените данни при прехода към устойчиви хранителни системи

Хранителната система представлява сложен набор от взаимодействия между множество участници и извършваните от тях дейности, които се отнасят до производството, обработката, разпределянето, подготовката и консумирането на храна, както и резултатите от тези дейности (HLPE, 2014). Устойчивата хранителна система е залегнала в основата на устойчивите цели за развитие на ООН. Нейната основна цел е да осигури продоволствена сигурност и храна за настоящото и бъдеще население на планетата, което означава чрез нея да се гарантира (FAO, 2018):

- *икономическа устойчивост* (носи икономическа изгода през цялото време);
- *социална устойчивост* (създава ползи за обществото);
- *екологична устойчивост* (няма негативно влияние върху околната среда).

В този контекст отворените данни имат стратегическо значение, тъй като осигуряват прозрачност по цялата хранителна верига, подпомагат вземането на решения в политиките за храните и климата, позволяват мониторинг на напредъка по Целите за устойчиво развитие (по-специално Премахване на глада [№2], Отговорно потребление и производство [№12], Климатични действия [№13]) и насърчават сътрудничество между институции и научни организации.

Данните, които се създават и използват в хранителните системи се генерират от най-различни източници, като:

- *сензори* (измерване на влажност на почвата; метеорологични станции; датчици за хранителни вещества; детектори, монтирани на машини и др.);
- *сателити*;
- *системи за управление на стопанства*;
- *платформи за пазарна информация*.

Всеки един от посочените източници генерира данни, които са полезни на някакъв етап от производството и консумирането на храни. Наличието на големи обеми с данни от сензори е в основата на иновативни управленски подходи, като прецизното земеделие, където се разчита на прилагането на различни технологични иновации (Branzova, 2024). Чрез тях се реализиране събиране, обработка и анализ на времеви, пространствени и индивидуални данни, за да се подобри добива, качеството на реколтата, да се използват ефективно ресурсите и по този начин да се премине към устойчиви операции (Soussi, Zero, Sacile, Trincherro, & Fossa, 2024). Комбинирайки ги с различни масиви отворени данни с по-широк контекст и дълги динамични редове се дава възможност да се разширят анализите и да се подобрят моделите и базираните на тях последващи решения. Същевременно, проучване относно използваните масиви с данни (отворени или споделени данни) в научни публикации в областта на селското стопанство за периода 2012 – 2022 година, установява, че близо 70% от цитираните източници са свързани със сателитни данни (Chamorro-Padial, García, & Gil, 2024). На следващо място, с 12%, са масиви с изображения (на растения, плевели, насекоми) и едва след това се нареждат данните от географски информационни системи, източници на метеорологични данни или някакви статистически данни (всеки от тях с около 3%).

На международна ниво с изключителна значимост се открояват следните няколко портала за отворени данни:

- *FAOSTAT* (<https://www.fao.org/faostat/en/>) – поддържан от Организацията по прехрана и земеделие (FAO), порталът предоставя стандартизирани данни за дълъг период от време за производство, търговия, потребление, цени, ползване на земя, риболов, горско стопанство и свързани теми за държави и територии от целия свят, позволявайки да се сравняват показатели, да се изграждат или тестват модели, да се следят индикатори за продоволствена сигурност и устойчивост и да се правят възпроизводими анализи, като същевременно е налице добра документация от гледна точка на метаданни и методологични ограничения;
- *Eurostat* (<https://ec.europa.eu/eurostat/data/database>) – предоставя статистически данни в стандартизиран вид за страните от Европейския съюз за производството и добивите по култури и животни, търговията и цените на храните, потребление и хранителни навици, загуби и отпадъци от храни, използване на земи и води, емисии и отпечатъци от земеделие, здравословност на диетите и сигурност на снабдяването, като тези данни се използват за оценка на напредъка, разработване на политики и изграждане на мониторинг инструменти за преход към по-устойчиви хранителни системи;
- *World Bank Open Data* (<https://data360.worldbank.org/en/search>) - Световната банка поддържа отворена платформа за данни, която предлага широк набор от икономически, екологични и социални индикатори, включително тематични набори за селско стопанство и развитие на селските райони (площи, производство, добиви, напояване, използване на торове и заетост в селското стопанство), както и специализирани инициативи за трансформация на хранителните системи и агрегирани показатели за продоволствена сигурност и устойчивост, които позволяват сравняване между държави, анализ на тенденции, моделиране и ангажимент чрез инструменти и колекции като World Development Indicators и програмата Food Systems 2030.

Наред с посочените портали за отворени данни на международно ниво съществуват и такива на национално равнище (като например data.egov.bg в България, gov.uk в Обединеното кралство), в които също могат да се открият полезни данни и индикатори, свързани с функционирането на устойчивите хранителни системи. Чрез тяхното съдържание може да се направи връзката „локални институции - глобална екосистема“. Заедно с това, редица държави поддържат национални отворени бази данни за хранителния състав на множество хранителни продукти. В това отношение добри примери са базата USDA FoodData Central (<https://fdc.nal.usda.gov/>) на САЩ, Fineli (<https://fineli.fi/fineli/en/index>) на Финландия и Food and Nutrient Database (<https://www.data.go.kr/index.do>) на Южна Корея. Посочените портали и платформи за отворени данни представляват често използваните източници за анализ на данни, но далеч не отразяват всички налични.

Интересен пример за глобална отворена база данни за хранителни продукти е Open Food Fact (<https://world.openfoodfacts.org/>). Това е инициатива, създадена през 2012 година във Франция, която представлява граждански проект с нестопанска цел. Сайтът функционира на принципа на Wikipedia. Всеки потребител може да се включи като доброволец и да добавя данни за хранителни продукти (снимки, състав, хранителна стойност, баркод, екологична значимост). По този начин базата се разширява с най-различни продукти, продавани в различни държави по света. Всички данни се публикуват под лиценз ODbL и са свободно достъпни онлайн, като също могат да се изтеглят в CSV формат или посредством предоставен програмен интерфейс. Open Food Facts свързва отворените данни с устойчивите хранителни системи, като насърчава устойчивото потребление чрез предоставяне на прозрачна, достъпна и използваема за всеки информацията за храните - от обикновените потребители през изследователи до представители на законодателните органи.

За по-доброто разбиране на целия набор от източници на отворени данни за производството и консумирането на устойчиви храни е направен опит за извеждане на класификация на данните и съответните източници, в които могат да бъдат открити изцяло или частично. Следващата таблица представя два възможни варианта за изготвяне на подобна групировка.

Таблица 1: Класификация на отворените данни за устойчиви хранителни системи

	Група	Значение	Източници
Вариант 1	Данни за селскостопанско производство и употреба на ресурси	Фундаментални за оценка на устойчивостта на хранителната система.	FAOSTAT Eurostat Национални статистически институти
	Данни за състава на храните	Необходими за оценка на хранителната стойност на продуктите.	Open Food Facts USDA FoodData Central Национални бази данни за състава на храните
	Данни за потребление на храни и хранителните навици	Предоставят директен поглед върху консумацията на храна от хората.	World Bank Open Data Проучвания за бюджетите на домакинствата
	Данни за въздействието върху околната среда	Отбелязват екологичния отпечатък на храните през целия им жизнен цикъл.	Agri-footprint Ecoinvent (не са изцяло отворени)
Вариант 2	Данни от производство	Данни за площи, добиви, употреба на торове и пестициди, емисии на парникови газове, потребление на вода. Това са предимно агрегирани, макро данни. Често са с висока надеждност, но с по-ниска детайлност.	FAOSTAT Eurostat Национални статистически институти

Данни за продуктите и преработката	Данни за хранителен състав, съставки, алергени, опаковка, производител, екологичен отпечатък на продукта (Eco-Score). Данните на ниво отделен продукт са изключително детайлни, но качеството може да варира, особено при потребителски генерираните бази данни.	Open Food Facts USDA FoodData Central Food and Nutrient Database
Данни за потреблението и отпадъците	Разходи на домакинствата за храна, консумирани количества (често косвено изчислени), данни за хранителни отпадъци. Това често е най-трудната за намиране и най-малко стандартизирана категория данни.	Проучвания за бюджетите на домакинствата Доклади на Световната банка Научни публикации

Source: Авторско проучване

В първия случай разделението на данните е направено, като се следва логическата последователност от дейности и процеси, които протичат в хранителната система. При втория вариант класификацията на данните е направена на основата на мястото им във веригата на стойност на храните. Както се вижда от таблицата, описаните вече портали и платформи за отворени данни намират място и в двете класификации.

Заклучение

Настоящият доклад очертава значението на отворените данни като стратегически ресурс за прехода към устойчиви хранителни системи. Анализът потвърждава, че отворените данни вече не са просто технически набор от информация, а фундаментален катализатор за иновации, научни изследвания и разработване на политики, основани на доказателства. Чрез инициативи като GODAN, FDTP и наличието на утвърдени платформи като FAOSTAT и Eurostat, екосистемата на отворените данни в селското стопанство и производството и консумацията на храни непрекъснато се разширява. Въпреки безспорния напредък, остават предизвикателства, свързани с качеството на данните, тяхната оперативна съвместимост, навременност и достъпност. Успехът на прехода към устойчиви хранителни системи ще зависи не само от генерирането на повече отворени данни, но и от възможността им да бъдат интегрирани, анализирани и превърнати в приложими знания. Продължаващото сътрудничество между правителства, частния сектор, научната общност и гражданското общество за насърчаване на отвореността на данните е от съществено значение за осигуряването на продоволствена сигурност и екологична устойчивост за бъдещите поколения.

References

1. Европейски Съвет. (2019, Юни 20). *Директива - 2019/1024 - EN*. Retrieved from EUR-Lex: <https://eur-lex.europa.eu/legal-content/BG/ALL/?uri=CELEX:32019L1024>
2. Таков, Д. (2024). Устойчивост на хранителните системи - за оцеляване на хората и планетата. *Хранителна индустрия и търговия*, 6-8.
3. 5 Star open data. (2012, January 22). Retrieved from 5 Star open data: <https://5stardata.info/en/>
4. Branzova, P. (2024). Pretsiznoto zemedelie: Tehnologichni inovatsii za ustoychi-vo selsko stopanstvo (Precision Agriculture: Technological Innovations for Sustainable Agriculture). *Economic Thought Journal*, 69(1), 24-36. doi:<https://doi.org/10.56497/etj2469102>

5. Chamorro-Padial, J., García, R., & Gil, R. (2024). A systematic review of open data in agriculture. *Computers and Electronics in Agricultur*, 108775.
6. FAO. (2018). *Sustainable food systems: Concept and framework*. Food and Agriculture Organization of United Nations. Retrieved October 17, 2025, from <https://openknowledge.fao.org/server/api/core/bitstreams/b620989c-407b-4caf-a152-f790f55fec71/content>
7. *Food Data Transparency Partnership*. (2025, 10 12). Retrieved from GOV UK: <https://www.gov.uk/government/groups/food-data-transparency-partnership>
8. HLPE. (2014). *Food Losses and Waste in the Context of Sustainable Food Systems, A Report by the High Level Panel of Experts on Food Security and Nutrition of the Committee on World Food Security of Sustainable Food Systems*. Retrieved September 17, 2025, from www.fao.org/cfs/cfs-hlpe
9. Ling, L., Hu, L., Li, S., Zhao, X., & Ye, X. (2025). How does open public data affect enterprise green transformation? *Socio-Economic Planning Sciences*, 102(Desember), 102342.
10. Open Data Charter. (2015). *International Open Data Charter*. Retrieved from https://opendatacharter.org/wp-content/uploads/2023/12/opendatacharter-charter_F.pdf
11. Peng, X., & Xiao, D. (2024). Can Open Government Data Improve City Green Land-Use Efficiency? Evidence from China. *Land*, 1813, 1-17.
12. Sánchez-Padial, A. (2017). Open data initiatives facing the challenges of global food security. *agriRxiv*, Wallingford.
13. Satrio, B., Kleiman, F., & Janssen, M. (2026). Open Data Portals Engagement: A Cross-Country Analysis of Game Elements. *Business Modeling and Software Design. Lecture Notes in Business Information Processing*, 241-250.
14. Soussi, A., Zero, E., Sacile, R., Trincherro, D., & Fossa, M. (2024). Smart Sensors and Smart Data for Precision Agriculture: A Review. *Sensors*, 24(8), 2647. doi:<https://doi.org/10.3390/s24082647>
15. Transforming agriculture through Open Data. (2025, Септември 11). Retrieved from Godan: <https://www.godan.info/about>
16. Wilkinson, M., Dumontier, M., & Aalbersberg, I. e. (2016). The FAIR Guiding Principles for scientific data management and stewardship. *Scientific Data*, 3(1). doi:<https://doi.org/10.1038/sdata.2016.18>
17. Wirtz, B., Weyerer, J., & Rösch, M. (2022). Open government data: A systematic literature review of empirical research. *Electronic Markets*, 32(4), 2381-2404.
18. Yan, R., Xing, C., Chen, X., & Zhao, Y. (2023). Is it real or illusory? an empirical examination of the impact of open government data on innovation capability in the case of China. *Technology in Society*, 75, 102396.

МЕТОДИ ЗА АВТОМАТИЗАЦИЯ НА БИЗНЕС АНАЛИЗИ ЧРЕЗ ИНТЕГРИРАНЕ НА НЕЕДНОРОДНИ ДАННИ

Methods for Automating Business Analysis by Integrating Heterogeneous Data

Теодор Тодоров¹,
e-mail: ttodorov131@unwe.bg¹

Абстракт

В съвременната бизнес среда няма компания, на която да не се налага да обработва и анализира данни. Независимо дали става дума за структура в частния или общественния сектор, с двама или десетки служители, обемът информация нараства, но с него расте и потенциалът бизнес единиците да извличат качествено знание. Най-разпространеното „решение“ за съхранение и обработка на първична информация дори и през 2025 остава Microsoft Excel. Наличието на напреднали инструменти с изкуствен интелект не винаги решава първичния проблем – как данните да бъдат интегрирани и представени, така че да могат да се използват занапред за задълбочен анализ и проследяване на тенденции. Настоящият доклад следва да предложи методология за автоматизация обработката на информация, която обединява процесите по събиране, уеднаквяване и трансформация на разнородни данни, извлечени от Excel източници. Ще бъдат предложени добри практики при първичното съхраняване на данни, ще бъдат разгледани съвременни low-code/no-code решения и ще се предложат идеи за приложение в малки и средни предприятия с ограничен ресурс.

Abstract

In the current business environment, there is no company that does not have to process and analyze data. Regardless of whether it is a structure in the private or public sector, with two or dozens of employees, the volume of information is growing, but with it the potential for business units to extract qualitative knowledge is also growing. The most common “solution” for storing and processing primary information even in 2025 remains Microsoft Excel. The availability of advanced artificial intelligence tools does not always solve the primary problem – how to integrate and present data so that it can be used in the future for in-depth analysis and trend tracking. This report should propose a methodology for automating information processing, which unites the processes of collecting, unifying and transforming heterogeneous data extracted from Excel sources. Good practices in primary data storage will be proposed, modern low-code/no-code solutions will be considered, and ideas for application in small and medium-sized enterprises with limited resources will be proposed.

Ключови думи: обработка на данни, изкуствен интелект, извличане на информация
JEL: C810

Увод

Годината е 1985 г. и компанията Microsoft пуска в продажба първата версия на системата за управление на електронни таблици MS Excel за компютри Apple Macintosh. Макап Lotus 123 – доминиращата платформа за изчисления по това време, да предлага възможности за изчисления, създаване на графики и макроси, Excel успява да се наложи с простотата на своето използване. На потребителите не се налага да имат сериозни познания по програмиране, за да въвеждат, пресмятат и визуализират информация. Excel надгражда идеята с по-интуитивен интерфейс и интеграция с графичен интерфейс – нещо ново за средата на 80-те. В този момент, едва ли някой предполага, че MS Excel ще се превърне в най-универсалния инструмент за управление на данни в света. През 1987 г. излиза версия и за новата MS Windows, а през 1993 г. Excel надминава по

¹ Докторант, катедра „Информационни технологии и комуникации“, факултет „Приложна информатика и статистика“, УНСС, гр. София, ORCID 0009-0002-5119-8873

пазарен дял Lotus 1-2-3, превръщайки се в синоним на електронна таблица (Walkenbach, 2013). С годините Excel се развива от инструмент за прости изчисления в универсална работна среда за данни. Финансисти, счетоводители, маркетингози и инженери започват да го използват не само за събиране и изваждане, но и за съхранение на големи обеми от информация, проследяване на запаси, проекти и бюджети. Така възниква парадокс - продукт, създаден като изчислителен инструмент, постепенно се превръща в *де факто* база от данни. Както посочва Реймънд Панко (2013), „повечето организации разчитат повече на електронни таблици, отколкото на официалните си бази данни“ – не защото това е най-доброто решение, а защото Excel е познат, достъпен и не изисква специална ИТ инфраструктура.

С разпространението на Microsoft Office през 90-те Excel става стандартен компонент на всеки офис компютър. В много компании той се превръща в първото и последното звено на анализа: данните се въвеждат ръчно, трансформират се с формули и макроси, а резултатите често се използват за управленски решения. Тази практика, макар и удобна, поставя сериозни ограничения: липса на контрол върху версиите, несъгласуваност на данните и трудности при обединяване на множество файлове. Kruck и Maher (2012) описват това като „илюзията на базата данни“ – Excel дава усещане за структура, но не осигурява нито устойчивост, нито връзки между данните, каквито предлагат релационните системи.

Днес, въпреки наличието на мощни BI и ERP решения, Excel остава повсеместен. Данните на Gartner (2024) показват, че над 80% от малките предприятия в световен мащаб продължават да използват електронни таблици като основен инструмент за съхранение и анализ на информация. Това е резултат от културна и икономическа реалност – Excel е познат, гъвкав и „*достатъчно добър*“ за ежедневните нужди. Именно тук възниква и ключовият въпрос на настоящия доклад: как можем да запазим познатата среда, но да автоматизираме обработката, уеднаквяването и интегрирането на нейните нееднородни данни, така че тя да служи като основа за съвременен бизнес анализ.

Excel срещу базите данни

На теория, Excel и базите данни решават един и същ проблем — съхраняване и обработка на данни. На практика, обаче, те принадлежат към два напълно различни свята. Електронната таблица е създадена, за да бъде *гъвкав инструмент за изчисления*, докато базата от данни е *структурирана система за съхранение и взаимовръзки* (Gartner). Разликата може да се назове и още по-просто - в Excel редовете „живеят“ самостоятелно, докато в една база от данни те имат взаимоотношения и зависимости. Excel не разполага с механизми за **референтна цялост** – няма първични и външни ключове, които да предотвратят дублиране или несъответствия. Една и съща стойност, например името на клиент или идентификатор на продукт, може да се появи по десет различни начина в различни файлове. В релационна база данни това би било невъзможно: всяка таблица е свързана чрез ключове, които гарантират логическа последователност и уникалност на данните (Inmon, 2005; Kimball & Ross, 2013). Това структурно различие обяснява защо толкова често бизнес анализът в Excel започва с часове или дни ръчно почистване – премахване на дубликати, корекция на грешни формати, уеднаквяване на имена и дати. Така възниква феноменът на т.нар. **нееднородни данни (heterogeneous data)** – множество файлове с различни колони, различни формати на датите, несъответстващи типове стойности и липсващи данни. Обикновено всеки файл отразява различен процес, екип или период, но без обща структура, която да ги обедини в единна картина. Според Power (2020), „в корпоративната практика нееднородните данни са по-скоро правило, отколкото изключение“, а най-често срещаните източници на такава нееднородност са именно електронни таблици и CSV файлове, създадени ръчно.

Най-простият пример за структурен конфликт е при форматирането на дати: в един файл датите са записани като „01/02/2025“, в друг като „1 февр. 2025“, а в трети – просто като пореден номер (числото, което Excel използва за броене на дни от 1900 година, например 43551). Подобни

разминавания водят до грешки при обединяване на файлове и – естествено – могат да доведат до погрешни изводи при анализ. В SQL база данни подобно разминаване е невъзможно – типът на данните се определя предварително (DATE, VARCHAR, INTEGER и др.), което предотвратява въвеждането на несъвместими стойности. Същият проблем възниква и при липсата на **метаданни** – информация за това какво означава всяка колона. В Excel често единствената „документация“ е заглавният ред, а той може да бъде променен, съкратен или преведен различно в различни файлове. Без стандарт за именуване и без контрол на версиите, обединяването на множество таблици се превръща в предизвикателство, което често надвишава времето за самия анализ. Както отбелязва Kruck и Maher (2012), електронната таблица „придава илюзия за структура“, но реално оставя потребителя да решава сам какво означава всяка колона и как трябва да се тълкуват данните. Всичко това води до необходимостта от **интеграционен слой** – процес или инструмент, който да обедини нееднородните източници и да наложи минимална форма на хармонизация. В миналото тази роля се е изпълнявала от релационните бази данни, но днес все по-често се използват гъвкави ETL инструменти и low-code решения, които позволяват да се създаде подобна структура *върху* Excel, без да се изисква преминаване към напълно нова платформа. Така Excel продължава да бъде неофициалната „входна точка“ към данните, но съвременните технологии дават възможност те да бъдат интегрирани, почистени и анализирани автоматично, без да се жертва достъпността на познатата среда.

Автоматизацията – от ръчна обработка към интелигентни потоци

В продължение на десетилетия работата с данни в предприятията следва една и съща логика — човекът е в центъра на процеса. Той копира, пренася, сравнява, филтрира и обединява файлове, за да изведе смисъл от информацията. Excel, уви, все още може да бъде основен инструмент на тази ръчна обработка, но именно тя днес е сериозно ограничение. Според проучване на Deloitte (2023), над 40% от времето на бизнес анализаторите в малките и средни предприятия се отделя за техническа подготовка на данни, а не за самия анализ. В тази среда автоматизацията не е просто удобство, а **методологична необходимост**. Тя означава прехвърляне на повторемите и податливи на грешки стъпки – копиране, съвпадение на колони, проверка на формати – към система, която може да ги извършва сама, по предварително зададени правила. Най-простата форма на това е позната на всеки потребител на Excel: макросите. Те автоматизират рутинни операции чрез запис на последователни действия. Но макар и полезни, макросите остават „вътрешно решение“ – те работят само в рамките на един файл и не решават проблема с интегрирането на данни от множество източници.

С навлизането на **Power Query** (2010 г.) и по-късно на **Power Automate** и **Power BI**, Microsoft постепенно превърна Excel от изолирана електронна таблица в елемент от по-широка екосистема за обработка на данни. Power Query позволи потребителите да извличат, трансформират и зареждат информация (т.нар. ETL процес) от десетки източници — други Excel файлове, бази данни, уеб страници или API услуги. Тези инструменти, въпреки че все още изискват известна техническа настройка, вече въвеждат идеята за „**интелигентен поток на данни**“ (**data pipeline**) – процес, който може да се изпълнява автоматично и многократно при постъпване на нова информация. Следващата еволюционна стъпка е появата на **low-code** и **no-code** платформи – визуални среди, в които бизнес потребителите могат да изграждат автоматизирани потоци без програмиране. Примери за такива решения са **Make (Integromat)**, **Zapier**, **Airtable**, **KNIME** и **Talend Open Studio**. Чрез тях една компания може например:

- автоматично да извлича нови Excel файлове от споделена папка или имейл;
- да уеднаквява формати (дата, валута, мерни единици);
- да премахва дублирани записи;
- и накрая да записва резултата в централен отчетен файл или BI платформа.

Подобни сценарии обикновено вече не изискват програмиране, а могат да бъдат създадени от самия бизнес анализатор чрез „визуални блокове“. Gartner (2024) нарича това явление

„демократизация на интеграцията“ – процес, при който достъпът до автоматизацията напуска ИТ отдела и се премества към оперативните звена на организацията. По този начин малките и средните предприятия получават възможност да изграждат собствени потоци за обработка на данни без значителни инвестиции. Различните low-code инструменти прокламират на своите продуктови страници, че могат да съкращават времето за интеграция на данни с до 60% и да намалят риска от човешки грешки. Въпреки това, автоматизацията не е универсално решение на казуса. Нейната ефективност зависи от качеството на първичните данни и от яснотата на логиката, която системата трябва да следва. Ако във входните файлове липсват основни идентификатори, стойностите са непоследователни или форматите варират твърде драстично, нито един инструмент – колкото и интелигентен да е – не може напълно да елиминира нуждата от човешка намеса. Затова днешната цел не е пълна автоматизация, а **създаване на „интелигентна инфраструктура“**, която да комбинира повторемостта на машината с експертизата на анализатора.

Така автоматизацията се превръща не просто в техническо решение, а в **нов тип организационно знание** – начин компаниите да мислят за данните си като за жив процес, който може да се поддържа, обновява и развива. Именно тази концепция стои в основата на следващата глава, в която ще бъде предложен методологичен модел за интеграция на нееднородни данни, адаптиран за нуждите на малките и средни предприятия.

Модел за автоматизация на данни

Съвременният бизнес анализ изисква не просто инструменти, а **методология** — последователен начин за организиране на данните така, че те да бъдат достъпни, съпоставими и полезни. За да се постигне това, в настоящия доклад се предлага опростен, но универсален модел за автоматизация на анализа на нееднородни данни, базиран на четири основни фази: **Extract (Извличане) -> Harmonize (хармонизиране, превръщане на данните в еднородни) -> Validate (валидиране на хармонизираните данни) -> Analyze (анализ; EHVA)**. Този процес на практика е базиран почти изцяло на класическата ETL парадигма, но е адаптиран за реалността на малките и средни предприятия, които работят с ограничен ресурс и разчитат предимно на Excel и сродни инструменти. EHVA е предназначен за интеграция и анализ на нееднородни бизнес данни в среда с ограничен технологичен ресурс. За разлика от традиционния ETL, EHVA не включва етапа на „зареждане“ на данни (Load), тъй като фокусът е върху уеднаквяването и проверката на данните, а не върху тяхното съхранение.

Извличане на данни

Първият етап от процеса включва системно събиране на всички файлове и таблици, които съдържат релевантна информация – обикновено Excel документи, CSV файлове или експорти от различни системи. Тук ключовият момент не е техническият инструмент, а **последователността на организацията**. Файловете трябва да бъдат събирани по ясен стандарт:

- Уеднаквени имена на колони (напр. „Дата“, „Приход“, „Разход“ вместо произволни наименования);
- Групиране според общ критерий – напр. период, обект или процес;
- Запазване на оригиналните файлове за проследимост.

Дори при ръчно събиране, този етап може да бъде частично автоматизиран чрез инструменти като **Power Automate, Zapier** или **Make**, които могат да извличат и подреждат файлове от имейл, облак или вътрешна система (Gartner, 2024).

Уеднаквяване на структурата

След извличането следва най-съществената стъпка – хармонизацията. Тук целта не е пълно уеднаквяване на всички данни, а постигане на **съвместимост**, достатъчна за анализ. Low-code решения като **Power Query**, **KNIME** или **Talend** позволяват на потребителите визуално да зададат как да се обединяват таблици с различни структури:

- Съпоставяне на колони (*schema matching*);
- Конвертиране на формати (дата, валута, мерни единици);
- Обединяване на таблици с различна дълбочина на детайла (напр. дневни и месечни отчети).

Този процес може да се „запише“ като шаблон и да се прилага автоматично при всеки нов набор от данни. Според Kim и Park (2022), подобна визуална настройка може да съкрати времето за интеграция на хетерогенни файлове с до 60% спрямо ръчната обработка.

Проверка и осигуряване на достоверност

След уеднаквяването идва моментът на валидиране — автоматична проверка за липсващи стойности, дублирани редове, несъвпадащи формати и логически грешки.

Дори без машинно обучение, low-code платформи могат да използват **вградени правила за валидация**:

- Проверка за празни клетки в ключови колони (ID, дати, стойности);
- Откриване на дублирани редове чрез сравняване на комбинации от полета;
- Известяване при отклонения от предварително дефинирани граници (например, необичайно висока стойност на разход).

Тези проверки могат да бъдат автоматично документирани, така че при следващ анализ системата да „знае“ какво се счита за допустимо.

Анализ и визуализация

След като данните са извлечени, уеднаквени и валидирани, те могат да бъдат визуализирани и анализирани в BI среда (напр. Power BI, Google Data Studio, Tableau). При всяка нова итерация на процеса — например добавяне на нов месец, нов файл или нова категория — системата може автоматично да опреснява анализите, без нужда от повторна ръчна обработка.

Така се постига **самообновяващ се аналитичен цикъл**, при който усилието е концентрирано в еднократната настройка, а не в постоянната поддръжка. Според Deloitte (2023), подобна архитектура може да намали оперативното време за изготвяне на периодични бизнес отчети с над 50%, като същевременно увеличава точността на данните.

ЕНВА моделът не изисква скъпа инфраструктура или специализирани екипи. Той може да бъде реализиран с достъпни и познати инструменти, а ключовото му предимство е **принципът на повторяемост** – веднъж настроен, процесът работи автономно и може да се прилага върху неограничен брой нови файлове. В контекста на малките и средни предприятия това означава, че усилието се измества от „правене на анализ“ към „изграждане на процес“, който сам подготвя данните за анализ.

Заклучение

В динамичната бизнес среда през 2025 година, данните са не само в стратегически ресурс, но в ежедневен предизвикателен обект за управление. За малките и средни предприятия това предизвикателство често се изразява в липса на хармонизирани източници, ограничен човешки и технически капацитет и зависимост от инструменти като Excel, които, макар и гъвкави, не са проектирани за обработка на големи и нееднородни масиви от информация. Настоящият доклад показва, че автоматизацията на бизнес анализа невинаги изисква непременно сложна инфраструктура или високоспециализиран персонал – тя може да бъде постигната чрез комбинация от съществуващи low-code решения и ясна методология. Предложеният модел по своята същност не е нещо ново и нечувано, но представлява адаптирана рамка, ориентирана към реалните условия на българските малки и средни фирми, които боравят с разнообразни източници и периодически натрупващи се данни. Тази методология извежда фокуса от традиционната трансформация и съхранение (ETL) към процес на **хармонизиране и повторям анализ**, който може да се реализира в достъпна среда. По този начин EHVA не просто описва технически етапи, а предлага нов организационен подход към обработката на данни – при който усилието е концентрирано в еднократната настройка, а не в рутинното ръчно изпълнение. Внедряването на подобен модел носи няколко съществени предимства: съкращава времето за подготовка на информация, намалява риска от човешки грешки и осигурява по-висока повторямост на резултатите. Най-същественото обаче е, че то позволява на предприятията да **превърнат данните си в инструмент за придобиване на знание**, а не просто в отчетен ресурс. Съчетаването на машинното обучение, low-code технологиите и ясно дефинирана логика на обработка създава възможност за устойчива трансформация – такава, която не замества човека, а го освобождава да мисли стратегически.

Автоматизацията на бизнес анализите чрез интегриране на нееднородни данни не е само технологична тенденция, а форма на бизнес интелигентност. В свят, в който скоростта на промяната надвишава способността ни да реагираме, именно структурираното, хармонизирано знание ще бъде най-ценният капитал на малкия и средния бизнес.

Източници

1. Berthold, M.R., Cebron, N., Dill, F., et al. (2022). *KNIME: the Konstanz information miner*, KNIME Press.
2. Deloitte (2023). *The Rise of Low-Code/No-Code Development: Democratizing Application Development*. Deloitte Insights.
3. Gartner (2024). *Hype Cycle for Data Integration Tools*. Gartner Inc. Available at: <https://www.gartner.com/en/articles/data-integration> (Accessed: October 2025).
4. Gartner (2024). *Market Guide for Data Quality Tools*. Gartner Inc. Available at: <https://www.gartner.com/en/data-analytics/topics/data-quality> (Accessed: October 2025).
5. Gartner (2024). *Enterprise Low-Code Application Platforms*. Gartner Peer Insights. Available at: <https://www.gartner.com/reviews/market/enterprise-low-code-application-platform> (Accessed: October 2025).
6. Inmon, W.H. (2005). *Building the Data Warehouse* (4th ed.). Wiley.
7. Elshan, E., Dickhaut, E. and Ebel, P. (2023). *An Investigation of Why Low Code Platforms Provide Answers and New Challenges*
8. Kimball, R. and Ross, M. (2013). *The Data Warehouse Toolkit: The Definitive Guide to Dimensional Modeling* (3rd ed.). Wiley.
9. Panko, R.R. and Ordway, N. (2005). *Sarbanes–Oxley: What About All the Spreadsheets? Proceedings of the European Spreadsheet Risks Interest Group (EuSpRIG 2005)*

10. Microsoft (2023). *Power Query and Power Automate Documentation*. Microsoft Learn.
Available at: <https://learn.microsoft.com/power-query> (Accessed: October 2025).
11. Panko, R.R. (2008). *Spreadsheet Errors: What We Know. What We Think We Can Do*.
12. Power, D.J. (2015). *Data Science Supporting Decision Making*.
Available at: <https://www.tandfonline.com/doi/full/10.1080/12460125.2016.1171610>
(Accessed: October 2025)
13. Walkenbach, J. (2013). *Excel Bible* (10th Anniversary ed.). Wiley.

ОСНОВНИ АСПЕКТИ НА ВЪВЕЖДАНЕТО НА ЦИФРОВО ЕВРО

Key aspects of the introduction of the digital euro

Аглика Кънева¹

e-mail: akaneva@unwe.bg¹

Абстракт

В доклада са представени същността и основните характеристики на цифровото евро. Изброяват се целите на цифровото евро. Разгледани са предимствата на цифровото евро.

Abstract

In this paper, the nature and main characteristics of the digital euro have been presented. The objectives of the digital euro have been enumerated. The advantages of the digital euro have been examined.

Ключови думи: цифрово евро, ЕЦБ, плащания.

JEL: G21, O3

Увод

Еврот е символ на единството и силата на Европа от създаването си (European Commission, n.d.b). След двугодишна фаза на проучване на цифровото евро Управителният съвет на ЕЦБ реши да премине към подготвителна фаза, която започна на 1 ноември 2023 г. Подготвителната фаза е насочена към по-нататъшно разработване и тестване на цифровото евро в съответствие с проектните решения и техническите изисквания, определени през фазата на проучване. В този контекст Евросистемата провежда обширни анализи, тестове, експерименти и ангажира заинтересованите страни, за да гарантира, че цифровото евро отговаря на най-високите стандарти за качество, сигурност и използваемост. Настоящата фаза ще продължи две години и ще приключи в края на 2025 г., когато Управителният съвет ще реши дали да премине към следващата фаза на подготовката и, ако това стане, ще определи нейния обхват и продължителност. Решението за емитирането на дигитално евро ще бъде разгледано от Управителния съвет на ЕЦБ едва след приемането на съответното законодателство.

На 28 юни 2023 г. Европейската комисия представи проект на законодателно предложение за евентуално въвеждане на цифрово евро. Целта на законодателството е да се гарантира, че в бъдеще цифровото евро ще даде на гражданите и фирмите допълнителна възможност да плащат по електронен път, като използват широко приета, евтина, сигурна и устойчива форма на публични пари навсякъде в еврозоната.

Цифровото евро ще позволи на потребителите да извършват сигурни незабавни плащания във физически и онлайн магазини, както и между физически лица, независимо от страната от еврозоната, в която се намират, или от техния доставчик на платежни услуги. ЕЦБ в момента проучва как това би могло да се осъществи на практика (ECB, 2024).

Основната цел на доклада е да се изследват основните аспекти на въвеждането на цифрово евро.

За да се реализира целта, в доклада се очертават следните конкретни задачи:

1. Представяне на същността на цифровото евро.
2. Изброяване на основните характеристики на цифровото евро.

¹ Chief Assist. Prof., Ph.D. Department of Finance, Faculty of Finance and Accountancy, UNWE, ORCID ID: 0000-0002-6903-9523.

3. Изследване на предимствата от въвеждане на цифрово евро.

Обект на изследването е цифровото евро. Предмет на изследването са основните характеристики на цифровото евро.

Основната теза на изследването е, че въвеждането на цифрово евро от ЕЦБ ще създаде ползи за икономиката на еврозоната.

Същност на цифровото евро

Цифровото евро ще бъде „цифрова валута на централната банка“, емитирана от ЕЦБ и достъпна за широката публика. То ще бъде точно като пари в брой, но в цифрова форма. Подобно на паричните средства всяко цифрово евро, притежавано от потребителите, ще бъде пряко обезпечено от ЕЦБ. То ще бъде предоставяно на гражданите и фирмите от банките и други доставчици на платежни услуги.

За разлика от криптоактивите цифровото евро ще бъде пари на централната банка. ЕЦБ ще гарантира, че то е сигурно, запазва стабилна стойност, и че може да бъде обменено по номинална стойност за евро в брой. За разлика от него стойността на криптоактивите може да варира значително и тяхната обмяна в евро в брой или дори в пари на търговска банка не може да бъде гарантирана.

Въпреки че парите в брой все още са широко разпространени и ще останат широко достъпни и приемани, все повече граждани и фирми избират да плащат по електронен път. В този контекст цифровото евро има няколко цели:

- да се гарантира, че хората, фирмите и държавните институции продължават да имат достъп до публична форма на цифрово парично средство за плащания, което е достъпно и се приема навсякъде в еврозоната по всяко време вместо да разчитат само на частни решения;
- да се предостави форма на дигитални пари, която гарантира същото ниво на защита на личните данни като наличните пари в брой за разлика от съществуващите решения за дигитални плащания, и е достъпна за всички граждани, включително и за тези, които нямат банкови сметки;
- да насърчава иновациите и конкуренцията в областта на плащанията на дребно, включително като дава възможност на банките и другите доставчици на платежни услуги да разработват нови решения за своите клиенти;
- да подкрепи отворената стратегическа автономност на Европа и да засили международната роля на еврото.

Цифровото евро ще бъде достъпно за физически лица, фирми и държавни органи, които пребивават или са със седалище временно или постоянно в държава-членка на еврозоната.

В определени случаи цифровото евро може да бъде достъпно и за лица, фирми и държавни институции, които не пребивават или не са със седалище в страна-членка на еврозоната като например:

- потребители и фирми, които са открили сметка в цифрово евро, когато са пребивавали или са били със седалище в държава-членка от еврозоната;
- потребители, които пътуват до еврозоната с лични или професионални цели;
- потребители, които пребивават или са установени в страна-членка извън еврозоната или в трета страна, при спазване на предварително определени условия, договорени с националните органи и/или централните банки.

На потребителите, посочени в първите два случая, може да бъде предоставен само временен достъп до сметки в цифрово евро при спазване на изискванията, определени от ЕЦБ.

Цифровото евро ще бъде достъпно за хората само по тяхно искане. Притежаването и използването на цифрово евро няма да бъде задължително.

Цифровото евро ще допълва евро банкнотите и монетите, а няма да ги замести. Физическите лица, фирмите и държавните институции ще имат възможност да избират дали да плащат с евро банкноти и монети, други частни електронни средства за плащане или с цифрово евро (European Commission, n.d.b).

Цифровото евро ще се различава от стабилните криптовалути и криптоактивите. Цифровото евро ще бъде пари на централната банка. Това означава, че то ще бъде обезпечено от централната банка и ще бъде проектирано да отговаря на потребностите на хората, които го използват. Поради това то би било безрисково. В допълнение то ще гарантира защитата на личните данни. Централните банки имат мандат да поддържат стойността на парите независимо дали са във физическа или цифрова форма.

Стабилността и надеждността на стабилните криптовалути всъщност зависят от лицето, което ги издава, и от доверието и изпълнимостта на тяхната гаранция, че ще поддържат стойността на парите във времето. Частните емитенти могат също да използват лични данни за търговски цели.

Няма лице, което да носи отговорност за криптоактивите, като това означава, че искове не могат да бъдат изпълнени (ECB, 2024).

Основни характеристики на цифровото евро

Добавената стойност на цифровото евро в сравнение със съществуващите частни решения за цифрови плащания като карти и мобилни плащания се състои в:

- потребителите могат да извършват дигитални плащания независимо от това къде се намират на територията на еврозоната - цифровото евро ще бъде единно платежно средство, което може да се използва в цялата еврозона, независимо от това къде се намират потребителите и коя търговска банка или доставчик на платежни услуги използват. Потребителите ще могат да плащат по всяко време и навсякъде в еврозоната, а плащанията ще се изпращат и получават незабавно 24 часа в денонощието, 7 дни в седмицата, 365 дни в годината. В момента не всички частни цифрови решения функционират безпроблемно в целия ЕС;
- възможност за дигитално плащане дори без достъп до интернет - цифровото евро може да се използва за извършване на плащания и при липса на интернет връзка, ако потребителят се намира физически близо до другата страна по транзакцията независимо дали това е друго лице или магазин или т.нар. „офлайн цифрово евро“. Потребителите, фирмите и държавните органи ще могат да извършват и получават плащания дори в отдалечени райони с нестабилна интернет връзка, както и в случай на липса на комуникационни мрежи или енергийна инфраструктура;
- по-голям избор за потребителите - цифровото евро ще допълни съществуващите частни решения за цифрови плащания. Това предоставя по-голям избор на потребителите, които ще имат възможност да избират между всички съществуващи възможности за плащане в зависимост от своите потребности, предпочитания и обстоятелства;
- възможност за дигитално плащане, дори ако потребителите не притежават банкова сметка - цифровото евро ще насърчи цифровата и финансова интеграция и по този начин ще допринесе за намаляване на цифровото разделение, като позволи на лицата без банкови сметки да извършват или получават цифрови плащания и да разполагат с безплатен достъп до основни функционалности. Такива функционалности ще включват конвертиране на пари в брой в цифрово евро и обратното;
- по-добра защита на личните данни на потребителите - цифровото евро ще даде възможност на потребителите да извършват цифрови плащания, като същевременно гарантира защитата на техните данни. При използване на цифровото евро офлайн защитата на личните данни на потребителя е същата като при използване на пари в брой.

Основните услуги за крайните потребители като откриване и закриване на сметка в цифрово евро, справки за салда, внасяне и теглене на средства от сметката в цифрово евро, както и извършване на преводи и плащания ще се предоставят безплатно.

Подобно на извършваните в момента платежни услуги потребителите на дигитално евро няма да заплащат такси при покупки в дигитално евро независимо дали са в страната или в чужбина. Банките могат да начисляват такси на своите клиенти само за банковите сметки, към които може да бъде свързано цифровото евро, и за доброволни неосновни услуги като например условни плащания.

За да подкрепи финансовата интеграция, цифровото евро ще бъде лесно за използване, достъпно навсякъде и по всяко време и безплатно за основна употреба.

Всички търговски банки, предоставящи услуги по разплащателни сметки, ще бъдат задължени да предоставят основни услуги за разплащания в цифрово евро при поискване. Някои държавни институции като местни и регионални органи на властта и пощенски клонове също ще предоставят цифровото евро на потребители, които не желаят да откриват сметка в цифрово евро, свързана с банка или друг доставчик на платежни услуги. Това ще позволи на хората без банкови сметки да имат достъп до цифровото евро.

Цифровото евро ще бъде лесно за употреба включително за хора с увреждания, функционални ограничения или ограничени цифрови умения, както и за възрастни хора в съответствие с Директива (ЕС) 2019/882 (Европейски акт за достъпността).

Също така ще бъде възможно да се извършват трансакции в цифрово евро без интернет връзка или „офлайн употреба“, като това е важна функция в райони с лош достъп до онлайн услуги или например в случаи на прекъсване на електрозахранването. Ще бъде възможно да се съхранява цифрово евро на местно ниво на електронни устройства т.е. „офлайн“ цифрово евро.

Нивото на защита на личните данни, което ще се въведе с цифровото евро, ще бъде безпрецедентно за електронните плащания.

Онлайн плащанията с цифрово евро са онлайн незабавни преводи, които могат да се извършват от разстояние. За да се извършват такива плащания, е необходима интернет връзка. По отношение на използването им от потребителите онлайн плащанията с цифрово евро няма да се различават от съществуващите електронни системи за незабавни плащания.

Офлайн плащанията с цифрово евро са незабавни преводи, които могат да се извършват без интернет връзка, ако има физическа близост между устройствата на платеща и получателя, както е при плащанията в брой в момента. Потребителите ще могат да съхраняват цифрово евро в своите устройства за офлайн употреба под определен лимит, точно както го правят с парите в портфейлите си. Офлайн плащанията с цифрово евро ще се валидират на принципа „peer-to-peer“ т.е. платещът и получателят ще проверяват директно дали прехвърлянето на стойност между тях действително е осъществено. Такива плащания биха се използвали основно за малки суми. Подобно на парите в брой подробностите за офлайн плащанията в цифрово евро няма да бъдат видими нито за банката на потребителите, нито за ЕЦБ.

Когато потребителите зареждат или теглят цифрово евро от портфейлите си, те трябва да са свързани към интернет.

Като парите в брой дигиталното евро ще се емитира пряко от ЕЦБ и националните централни банки на страните-членки на ЕС и ще представлява задължение на тези институции. Това означава, че то ще бъде публични пари или пари на централната банка. Цифровото евро ще има статут на законно платежно средство, което означава, че ще бъде достъпно за всички европейски граждани и резиденти и ще се приема навсякъде в еврозоната. Това не се отнася за съществуващите електронни средства за разплащане, предоставяни от търговските банки.

Потребителите ще могат да откриват сметка в цифрово евро във всяка търговска банка или друг доставчик на платежни услуги като платежни институции и институции за електронни пари. Ако потребителите нямат сметка в търговска банка или не искат да открият сметка в цифрово евро в

банка или доставчик на платежни услуги, те могат да ползват услугите на държавна институция, определена от тяхната държава-членка като пощенски клон. Клиентите могат да сменят доставчика си, когато пожелаят. В случай че техният посредник вече не е в състояние да им осигури достъп до тяхната сметка в цифрово евро, ще има автоматичен механизъм, който гарантира, че те могат да получат достъп до своето цифрово евро с помощта на друг частен посредник.

Цифровото евро ще бъде достъпно за плащания чрез онлайн банкиране и за мобилни плащания онлайн и офлайн в цялата еврозона. Това би било особено полезно, като се имат предвид затрудненията, с които понякога се сблъскват хората, когато се опитват да използват банковите си карти в чужбина. Клиентите ще могат да извършват плащания в цифрово евро чрез стандартния интерфейс за онлайн банкиране на своята банка или доставчик на платежни услуги, чрез специално приложение за цифрово евро или чрез други средства като карти, и да плащат, докато пазаруват в уебсайтове за електронна търговия, точно както при други електронни парични преводи и плащания.

Местата и ситуациите, в които потребителите ще могат да плащат с цифровото евро, ще се увеличават с времето. В даден момент цифровото евро ще бъде достъпно за всички основни ситуации, свързани с плащания, например за превод на пари („от лице на лице“), за плащане в магазини и ресторанти, за онлайн плащания или за плащания към държавни органи.

При мобилните плащания потребителите ще могат да извършват или получават плащания чрез мобилните приложения на своите доставчици на платежни услуги, точно както в момента. ЕЦБ може да реши да предложи специално приложение за цифровото евро, което потребителите да могат да избират да използват (European Commission, n.d.b).

В момента потребителите рядко имат възможност да използват незабавни плащания, когато плащат в магазините, което означава, че търговците не получават парите си веднага. Цифровото евро ще промени това, тъй като всички цифрови плащания в евро ще бъдат незабавни.

Единният пакет от правила, стандарти и процедури, разработен и, ако бъде одобрен, приложен за цифровото евро, би означавал, че незабавните плащания биха могли да бъдат доразвити, за да обхванат всички страни от еврозоната. Това би намалило зависимостта от частни неевропейски компании, които в момента доминират в сектора на плащанията (ECB, 2024).

Подобно на паричните средства сметките в цифрово евро няма да носят лихва т.е. потребителите не могат да получават лихва по депозитите си в цифрово евро.

Потребителите могат да имат една или няколко сметки в цифрово евро като при сметките в търговски банки. Потребителите могат също така да имат съвместна сметка с друго лице например член на семейството или роднина в допълнение към индивидуалните си сметки в цифрово евро. Потребителите ще могат да променят броя на сметките си в цифрово евро, ако желаят.

Като при парите в брой потребителите ще могат да използват цифровото си евро по начина, по който решат. Същевременно, въпреки че държавните органи не могат да програмират цифровото евро, хората ще могат да извършват текущи преводи и плащания, както и в момента.

Потребителите ще могат да използват цифровото евро за плащания на всякаква сума. Въпреки това, за да се гарантира паричната и финансовата стабилност, могат да бъдат наложени ограничения върху сумата, която потребителите могат да държат в цифрово евро. За да се гарантира, че търговските банки ще продължат да изпълняват полезна роля във функционирането на икономиката включително чрез отпускане на кредити, ЕЦБ ще има право да определя лимити за държане на цифрово евро. Тези ограничения ще се хармонизират на равнище еврозона и ще са стриктно подчинени на принципа на пропорционалност. Законодателното предложение, представено от Европейската комисия, установява конкретни критерии за възможни ограничения на функцията на цифровото евро като средство за съхранение на стойността като индивидуални лимити за държане.

В допълнение Европейската комисия ще определи лимити за използването на цифровото евро онлайн с цел да ограничи рисковете от пране на пари и финансиране на тероризъм.

Подобно на парите в брой цифровото евро ще получи статут на законно платежно средство, което означава, че търговците в еврозоната ще бъдат задължени да приемат плащания в цифрово евро от потребителите.

Въпреки това, ще се прилагат изключения за търговците в ситуации, в които налагането на такова задължение би било несъразмерно, например от гледна точка на разходите, които търговците извършват за изграждане и поддържане на платежната инфраструктура, необходима за приемане на плащания в цифрово евро. По-конкретно микропредприятията няма да бъдат задължени да приемат плащания в цифрово евро, освен ако вече не приемат сравними цифрови средства за плащане от потребителите като дебитни карти.

Търговците ще могат да обработват транзакции в цифрово евро като другите електронни средства за плащане чрез сметка в цифрово евро в банка или друг доставчик на платежни услуги, както и със съответните хардуерни и софтуерни компоненти. Доколкото е възможно, целта ще бъде да се гарантира, че търговците могат да използват устройствата, с които вече разполагат, за обработка на частни цифрови плащания.

Европейската комисия и ЕЦБ извършиха подробен анализ на потенциалното въздействие на цифровото евро върху банковия сектор. Този анализ показва, че само в случаи на голямо намаление на депозитите в полза на цифровото евро, банките могат да изпитат ликвидни затруднения и повишени разходи за финансиране, което може да доведе до намаляване на предоставяните кредити в икономиката. За да се противодейства на потенциалните рискове, предложението за регламент относно цифровото евро предоставя на ЕЦБ определени инструменти като лимити за държане, с които да ограничи функцията на цифровото евро като средство за съхраняване на стойността, ако ЕЦБ прецени, че това е необходимо за защита на финансовата стабилност. Тези инструменти няма да ограничават ежедневните плащания на потребителите, като същевременно ще смекчат въздействието върху банките и икономиката като цяло (European Commission, n.d.b).

Докато бъде прието окончателното законодателство, Евросистемата предлага модел за компенсации, който би създал справедливи икономически стимули за доставчиците на платежни услуги като банките да поемат оперативните разходи за разпространение на цифровото евро.

Като при другите платежни системи в момента доставчиците на платежни услуги, разпространяващи цифровото евро, ще могат да начисляват такси на търговците за тези услуги. Определянето на цените за търговците и доставчиците на платежни услуги ще бъде обект на ограничение, както е предложено от Европейската комисия в законодателното предложение за цифровото евро.

Като при производството и емитирането на банкноти Евросистемата ще поеме разходите за създаването на схема и инфраструктура за цифровото евро. Евросистемата ще се стреми да сведе до минимум допълнителните инвестиционни разходи за посредниците, като използва максимално съществуващата инфраструктура (ECB, 2024).

Физическите лица, фирми и държавни институции, намиращи се в еврозоната, ще могат да плащат с цифрово евро или да изпращат цифрово евро извън еврозоната в случаите, когато:

- получателите на средства, намиращи се извън еврозоната, имат сметка в цифрово евро и могат да обработват транзакции в цифрово евро или
- те използват междувалутни плащания, при които изпратеното от тях цифрово евро се конвертира в местна валута, когато плащането достигне получателите на средствата, намиращи се извън еврозоната. Това ще бъде полезно за туристи, които посещават страна извън еврозоната.

Междувалутните плащания с централни банки на трети страни например са плащания в цифрово евро срещу друга цифрова валута на централна банка.

Физическите лица, фирми и държавни органи, които пребивават или са със седалище извън еврозоната, могат да получат достъп до цифровото евро, като си открият сметки в цифрово евро при доставчици на платежни услуги, които са със седалище или извършват дейност в страна-членка на Европейското икономическо споразумение или в трета страна, при условие че е сключено предварително споразумение между ЕС и третите страни и/или споразумения между ЕЦБ и националните централни банки на страните-членки извън еврозоната и на третите страни (European Commission, n.d.b).

Механизъм на осъществяване на плащания с цифрово евро

Първата стъпка при извършване на плащания с цифрово евро ще бъде да се създаде портфейл с цифрово евро чрез банка или пощенски клон.

След като се създаде портфейл с цифрово евро, потребителите могат да внасят пари в него чрез свързана банкова сметка, или като депозират пари в брой. Тогава потребителите ще могат да започнат да извършват плащания, използвайки цифровото евро в портфейла си.

Когато получават пари в цифрово евро, потребителите ще ги съхраняват в своя портфейл с цифрово евро до определен лимит или ще ги депозират в банковата си сметка. Потребителите могат да направят това ръчно или да го настроят автоматично (ECB, n.d.).

Правна рамка на въвеждането на цифровото евро

Регламентът за цифровото евро ще уреди създаването на цифровото евро като нова форма на пари на централната банка, ще регулира неговите основни елементи и ще предостави възможност, но не и задължение, на ЕЦБ да емитира цифрово евро. ЕЦБ ще вземе решение дали да пусне в обращение цифровото евро в съответствие с мандата и задачите си.

Предложението за регламент относно цифровото евро, ако бъде прието от Европейския парламент и Съвета, ще въведе цифровото евро и ще определи необходимите правила за него по-специално по отношение на неговия статут на законно платежно средство, защита на личните данни, борба с прането на пари, разпространение, употреба (ограничения за използването му като средство за съхранение на стойността и условия за неговата употреба извън еврозоната) и основни технически характеристики. Основните технически характеристики включват главните функционалности на цифровото евро - офлайн, онлайн и условни плащания.

Съгласно Регламента за цифровото евро ЕЦБ може да приема детайлни мерки, правила и стандарти в съответствие със своите правомощия включително с цел да осигури гладкото и ефективно функциониране на платежната система за цифрово евро в съответствие с член 22 от Устава си.

ЕЦБ ще отговаря за създаването и проектирането на цифровото евро в съответствие с изискванията, определени в Регламента за цифровото евро (European Commission, n.d.b).

Предимства от въвеждането на цифрово евро

Цифровото евро ще бъде въведено преди всичко за използване от резидентите и фирмите в еврозоната и евентуално в целия ЕС. Въпреки това, използването на цифрово евро в международните разплащания на дребно може да донесе ползи за еврозоната и/или други икономики по отношение на търговията и паричните преводи, като улесни трансграничните плащания извън еврозоната. Това от своя страна би довело до конкретни ползи като улесняване на търговските отношения и намаляване на валутните рискове.

Цифровото евро би могло да засили конкуренцията и иновациите на европейския пазар за плащания на дребно, като улесни разработването на пълна гама от решения за крайни потребители в цялата еврозона, и като подкрепи цифровите финансови услуги. Например условните плащания, които представляват възможността да се даде указание за автоматично плащане при изпълнение на предварително определени условия, също биха могли да насърчат

развитието на иновативни услуги в сектора в ЕС, основани на цифровото евро, като допълнение към частните платежни решения.

Цифровото евро може да служи като резервна или допълнителна възможност в кризисни ситуации, или когато при частните средства за плащане възникнат оперативни проблеми. Това би засилило оперативната устойчивост на икономиката на ЕС.

Офлайн цифровото евро би могло също така да повиши устойчивостта на европейската платежна система, като гарантира непрекъснатото извършване на офлайн плащания с публични средства при спиране на интернет връзката (European Commission, n.d.b).

Също така цифровото евро ще засили стратегическата автономност и паричния суверенитет на еврозоната, като повиши ефективността на европейската платежна екосистема като цяло, насърчи иновациите и увеличи нейната устойчивост на възможни кибератаки или технически проблеми като прекъсвания на електрозахранването (ECB, 2024).

Въвеждане на цифрови валути от други централни банки

Много централни банки по целия свят започнаха проучвания относно евентуалното въвеждане на цифрови валути на централните банки (CBDCs). Те провеждат изследвания и пилотни програми, за да проучат потенциалните ползи и последици, подобно на това, което прави в момента ЕЦБ.

В развитите страни тази мотивация произтича предимно от намаляващото използване на пари в брой и необходимостта да се предложи електронна алтернатива за плащания с публични средства. В по-слабо развитите страни акцентът се поставя върху засилване на финансовата интеграция и подобряване на системата за разплащания на дребно.

В рамките на ЕС Швеция започна проучване на възможността за въвеждане на електронна крона.

Извън ЕС Великобритания провежда консултации и започна проучване за въвеждането на дигитална лира подобно на техническите проучвания на ЕЦБ за въвеждането на дигитално евро.

Извън Европа Китай вече пусна в обращение дигитален юан. Цифровият юан вече е достъпен за плащания в нарастващ брой региони и употребата му се улеснява от големите банки и доставчици на платежни услуги. САЩ проучват възможността за въвеждане на дигитален долар, но все още не са взели решение дали това е необходимо (European Commission, n.d.b).

Заклучение

Много централни банки по света в момента проучват възможността за емитиране на цифрови валути на централните банки и се увеличава броят на страните, които вече емитират такива валути (European Commission, n.d.b).

Като се има предвид, че други държави разработват свои собствени дигитални валути и се увеличава популярността на криптовалутите, от гледна точка на паричния суверенитет създаването на дигитална версия на еврото придобива все по-голямо значение (European Commission, n.d.b).

Стабилните криптовалути и други криптоактиви, които не са деноминирани в евро, ако се използват широко за плащания, могат да застрашат стабилността на паричната система на еврозоната. По тази причина е важно да се създаде цифрова форма на еврото, за да се гарантира, че гражданите, фирмите и държавните институции продължават да имат достъп до публична форма на пари в евро, която е достъпна и се приема навсякъде в еврозоната и по всяко време. Цифровото евро също така ще улесни плащанията на територията на еврозоната. То ще позволи на потребителите да плащат и да превеждат пари с висока степен на защита на личните данни и за разлика от много други решения за дигитални плащания дори без интернет връзка (European Commission, n.d.b).

Ако бъде прието, законодателното предложение, представено от Европейската комисия, ще регулира основните елементи на цифровото евро. След приемането на предложението за Регламент относно цифровото евро от Европейския парламент и Съвета, ЕЦБ трябва да вземе окончателното решение относно емитирането на цифрово евро. Пускането на цифровото евро в обращение може да отнеме още няколко години. Потребителите ще получават цифрово евро от своите търговски банки или доставчици на платежни услуги, или от публични органи, определени от държавите-членки, в замяна на депозити или евро в брой. Цифровото евро ще се емитира от ЕЦБ и националните централни банки на страните-членки на еврозоната.

Законодателното предложение относно законното платежно средство ще гарантира, че еврото в брой ще продължи да бъде широко приемано за плащания и лесно достъпно за гражданите, фирми и държавните институции в цялата еврозона (European Commission, n.d.b).

Потребителите, фирмите и държавните органи ще могат да извършват и получават плащания дори в отдалечени райони с ненадеждна интернет връзка и в случай на липса на комуникационни мрежи или енергийна инфраструктура. Също така в ежедневни ситуации хората биха могли да се възползват от цифровото евро, достъпно офлайн, когато плащат в ситуации, в които не са свързани с интернет (European Commission, n.d.b).

Разплащанията с дигитално евро ще бъдат сигурни и незабавни, както във физическите и в онлайн магазините, така и между физически лица (ECB, n.d.).

Литература

1. Bank for International Settlements, (2021). Fintech and the digital transformation of financial services: implications for market structure and public policy, BIS Papers, No 117, 13 July.
2. ECB, (2022). The digital transformation of the European banking sector: the supervisor's perspective, speech by Pentti Hakkarainen, Member of the Supervisory Board of the ECB, at the Institute for Financial Integrity and Sustainability, Frankfurt am Main, 13 January.
3. ECB, (2024). FAQs on a digital euro.
4. ECB, (n.d.). How would a digital euro work?
5. European Commission, (n.d.a). Digital euro, Brussels.
6. European Commission, (n.d.b). Frequently asked questions on the digital euro and the legal tender of cash, Brussels.
7. European Commission, (n.d.c). Overview of digital finance, Brussels.
8. Hakkarainen, P., (2018). The digitalisation of banking – supervisory implications, speech at the Lisbon Research Centre on Regulation and Supervision of the Financial Sector Conference, Lisbon, 6 June.
9. Kearney, (2022). European banks are balancing the push toward digitalization with continued demand for in-person interactions, European Retail Banking Radar, Frankfurt am Main.
10. McKinsey, (2021). McKinsey's Global Banking Annual Review, 1 December.
11. McKinsey Digital, (2020). Europe's digital migration during COVID-19: Getting past the broad trends and averages, 24 July.

МОДЕЛИ И СТРАТЕГИИ ЗА ДИГИТАЛНА ТРАНСФОРМАЦИЯ: НОВОТО НОРМАЛНО

MODELS AND STRATEGIES FOR DIGITAL TRANSFORMATION: THE NEW NORMAL

Светла Ценова¹

email: svetla.tsenova@unwe.bg¹

Абстракт

Дигитализацията е важна тема за компаниите във всички отрасли на икономиката. Съвкупността от дигитални решения се счита за четвърти колосален етап от технологичната и икономическа еволюция на човечеството, която касае по-добри подходи за боравене с нарастващите потоци от информация. Съществен елемент от тази трансформация е изкуственият интелект (ИИ), който е тясно свързан с целенасочената обработка на големи масиви от разнородни данни, използването на усъвършенствани компютърни алгоритми и основана на данни взаимовръзка между устройства и хора. Целта на настоящата статия е а) да анализира актуални модели и стратегии на дигитални пионери, б) да илюстрира тяхното приложение и ефекти върху ключови отрасли и в) да очертае препоръки за практическа рамка за системно и трайно внедряване и от отрасли – дигитални последователи. Резултатите от анализа на специализираната литература и практически казуси към момента показват, че устойчивата трансформация може да успее само ако **технологиите, стратегията и културата** се третират като **«триединство от равностойни измерения»**. Тази взаимосвързаност намира израз в редица модели за дигитална зрялост.

Abstract

Digitalization is a critical topic for companies across all sectors of the economy. The collective advancement of digital solutions is considered the fourth major stage in the technological and economic evolution of humanity, addressing the need for improved approaches to managing ever-growing information flows. A key element of this transformation is Artificial Intelligence (AI), which is closely linked to the targeted processing of large volumes of heterogeneous data, the use of advanced computational algorithms, and data-driven interactions between devices and people. The aim of this paper is to: (a) analyze current models and strategies employed by digital pioneers; (b) illustrate their application and impact across key industries; and (c) outline recommendations for a practical framework supporting systematic and sustainable implementation in industries that are digital followers. The results of the analysis of specialized literature and practical case studies indicate that sustainable transformation can succeed only when technology, strategy, and culture are treated as a “trinity of equally important dimensions.” This interdependence is reflected in a number of established digital maturity models.

Ключови думи: Дигитална трансформация, Дигитални стратегии, дигитално-ориентирана култура, Модели на дигитална зрялост, Трансформация на бизнес модели, Дигитална парадигма

JEL: O30, O32, O33, M14, M15, M16, M21, L21, L86, D83, Q01, Q56

Въведение

Дигиталната трансформация описва дълбоката интеграция на дигиталните технологии във всички области на организацията. Тя не само променя **процесите и структурите**, но и

¹ Редовен докторант и хоноруван преподавател в катедра „Управление“ на УНСС, email: svetla.tsenova@unwe.bg

предефинира *бизнес моделите*, създаването на *стойност и взаимоотношенията с клиентите и останалите заинтересовани страни* (Vaska et al., 2021). Според проучване на списание „MIT Sloan Management Review“, 90% от анкетирания мениджъри очакват дигитализацията да играе основна или много важна роля за тяхната компания (Kane et al., 2016). Вследствие на пандемията от COVID-19 и възхода на изкуствения интелект (ИИ), облачните изчисления и анализа на големи данни, интеграцията на данни се превръща от просто „мода“ в стратегическа необходимост за оцеляване. Процесът на дигиталната трансформация (ДТ) оказва значително въздействие (Schwer et al. 2018) върху начина, по който организациите функционират, гледат към бъдещето и се развиват устойчиво. Ефектите по линия на дигиталната трансформация се смятат за сравними с възникналите по време на първата индустриална революция (Westerman, Bonnet и McAfee 2014).

Преглед на актуалната специализирана литература

Академичният дебат относно дигиталната трансформация се засилва значително през последните години. От една страна, Виал (2019) определя «дигиталната трансформация» като „процес, който има за цел да подобри даден стопански обект при задействане на значителни промени, въздействайки върху комбинации от технически, социални и бизнес променливи“. Това разбиране разглежда дигиталната трансформация не просто като ИТ проект, а и като социо-културна трансформация. Няколко години преди Виал авторският колектив Бхарадвадж, Ел Сауи, Павлу и Венкатраман (2013) въвеждат концепцията за **дигитална бизнес стратегия (DBS)**, която подчертава стратегическата насоченост и дългосрочното въздействие на резултатите от цифровата трансформация върху бъдещето на организацията. Въпросните автори застават уверено зад убеждението, че дигиталните технологии се превръщат в *неразделна част от корпоративната стратегия* и насърчават стратегическата гъвкавост в дългосрочна перспектива. Уестърман, Боне и МакАфи (2014) показват още, че успешните компании за дигитална трансформация съчетават **ясна дигитална визия, силно лидерство и инвестиции в дигитални решения**. Кейн и кол. (2015) пък установяват, че организациите с **дигитално ориентирана култура** реагират по-гъвкаво на пазарните промени и внедряват иновации по-бързо, което се отразява върху ключов фактор като времето за достигане на пазарите. Gartner (2023) допълва тази перспектива с емпирични данни относно *Дигиталната зрялост*, която според него зависи не толкова от тех-инвестициите колкото от „подготвеността на културата“ и „лидерския синхрон“. Изследователският тандем Уорнър и Вагер (2019) разширяват този подход с **рамката за динамични възможности**, която обяснява как организациите разработват, интегрират в бизнес моделите си и преконфигурират цифрови решения, за да създават непрекъснато дигитална добавена стойност за клиентите, обществото и околната среда.

Подходи за управление на дигиталната трансформация

Дигиталната трансформация (ДТ) се счита за цялостен процес на промяна, който интегрира и изисква синхронизирано съчетание от (1) технологични иновации, (2) организационни промени и (3) културна адаптация. Успешното преминаване през цифрова трансформация се оценява не само като ключов двигател на конкурентоспособността, иновациите и издръжливостта на организациите, но и като фактор за по-адекватен отговор на изискванията на обществото за екологичност, устойчивост, социалност и съпричастност към глобалната еко-система. В общ план, могат да бъдат изведени *три основни направления на изследване на дигиталната трансформация*:

1. **Модели на дигитална зрялост** - използват се за пазарно позициониране с дигитална подготвеност и стратегическо планиране.
2. **Подходи за иновации и прекалибриране на бизнес модели** - акцент върху дигиталното препроектиране на организациите и създаването на дигитална стойност.

3. **Стратегически и организационни модели** - фокус върху развиването и утвърждаването на дигитална култура, дигиталното лидерство и управлението на дигиталните промени.
4. **Пътна карта на Кагерман** - Кагерман дава предписание за следване на следната пътна карта по пътя на дигитална трансформация на организациите: **дигитализация → интеграция → автоматизация → автономност** (Kagermann et al., 2018).

Тези подходи се допълват взаимно: докато *моделите за цифрова зрялост* осигуряват ориентация, *обновяването на бизнес моделите* осигуряват съществената насока, а *организационните модели* гарантират, че промяната е устойчиво закотвена (Teichert, 2019; Cosa, 2024).

Внедряването на **стратегии за децентрализирано управление на дигиталната трансформация** се подкрепя все повече от интегрирани софтуерни платформи, които опростяват планирането, внедряването и мониторинга.

Таблица 1: Инструменти и платформи в подкрепа на дигиталната трансформация

КАТЕГОРИЯ НА ИНСТРУМЕНТА	ПРИМЕРИ	ПРИЛОЖЕНИЕ
Оценка на дигиталната зрялост	Deloitte DigitalMaturity™, SAP Digital Transformation Navigator	Оценка на състоянието на трансформацията, извеждане на препоръки за действие
Управление на промени и проекти	Джира, Asana, Monday.com	Гъвкаво управление на портфолио, приоритизиране, прозрачност
Платформи за данни и анализи	Microsoft Fabric, Google Cloud BigQuery, Snowflake	Централно съхранение и анализ на данни, основа за решения, основани на данни
Автоматизация на бизнес процесите	UiPath, Power Automate, Camunda	Повишена ефективност, дигитализация на процесите
Инструменти за иновации и сътрудничество	Miro, Notion, SharePoint, Slack	Насърчаване на гъвкави и съвместни методи на работа
Киберсигурност и съответствие	IBM Guardium, Microsoft Purview	Осигуряване на защита на данните и съответствие по време на трансформация

Такива инструменти не са самоцел. Те следва да бъдат вградени в **стратегически модел на дигитално управление в съчетание с висока дигитална компетентност**, за да постигнат дългосрочен ефект и да съпътстват прехода на технологичната иновация към успешен и устойчив бизнес модел в синхрон с дигиталната бизнес стратегия на организацията.

От технологична иновация към бизнес модел

Дигиталната трансформация поставя света нащрек и го залива с все по-нови формати на бизнес модели, които поставят адаптацията на човешкия вид пред огромни предизвикателства за адаптация към дигитална работна и житейска среда (Brynjolfsson & McAfee, 2014). По отношение на нови бизнес модели се очаква те да бъдат насочени към автономни AI агенти, които да могат да планират и предприемат действия за постигане на цели, поставени от управленските кадри и/или предприемачи. Оформя се виртуална работна сила от агенти за подпомагане на експертните кадри и повишаване на тяхната ефективност. Предизвикателството е постигане на 100%-во съответствие на човешките намерения по веригата на доставки и изпълнението от страна на агента (Kaplan & Haenlein, 2019). Друга очаквана AI иновация са платформите за управление на AI. Предизвикателствата тук са свързани с гарантиране на отговорното използване на AI, разясняването на начина на работа на AI, както и осигуряването на прозрачност при изграждането на доверие и отчетност (Floridi & Cowls, 2019). Предпазването от дезинформация и защита от компроментиране на репутацията на марката ще изискват действия в областта на киберзащитата, обучението и повишаването на кибер-грамотността и развиване на дигитални компетентности (European Commission, 2024). В ход са възможностите

на постквантовата криптография и невидимата заобикаляща интелигентност, естествено интегрирана в околната среда и допринасяща за удобство и реалистични изживявания. Това ниво на интеграция между хора и нови технологии прави връзката двупосочна. Обектите докладват своята идентичност в реално време, своето поведение и свойства на околната среда. Тук доставчиците на този тип услуги ще се наложи да отговарят за отсъствие на пропуски по отношение на поверителността. Потребителите ще е необходимо да деактивират етикетите и да включват различни нива на поверителност и анонимност.

Енергоефективните изчисления са иновация и потенциален бизнес модел, който повишава съответствието с международния регламент за въглероден отпечатък. От организациите ще се изисква да актуализират облачната си инфраструктура и да повишат дела на използване на възобновяеми източници на енергия за работата на своите системи (NIST, 2023). Хибридните изчисления са друга технология и бизнес възможност, която предвижда производството на автономни модули, стигащи до интеграция на човешкото тяло като платформа. Предизвикателството тук са интеграцията и оркестрацията. Необходимо ще е и да се работи с нагласите и изграждане на дигитални навици и дигитална грамотност, както и промяна на възприятието на дигиталната идентичност на човека и „човешките“ свойства на продуктите на изкуствения интелект. Пространствените изчисления са технологично решение с потенциал за бизнес модел, които надграждат пространствено и времево и придават нови измерения на познатия досега физически свят, прилагайки разширена и виртуална реалност. Подпомага дистанционното виртуално и максимално реалистично вземане на решение (Milgram & Kishino, 1994). Намира приложение в отрасли като търговия на мебели, облекло и пр., в игровия сегмент, както и в областта на иновативното образование. В областта на многофункционалните роботи, като форма на изкуствен интелект в тесен смисъл, се заражда друго поле за бизнес растеж. Широко приложение могат да намерят в гастрономията, в здравеопазването, в социалните или частни домове на възрастни хора и хора с определена степен на инвалидност (Sharkey & Sharkey, 2012). Могат да заместват или да работят с хора. Имат висок социален ефект, но има и определени етични въпроси, които предстои да бъдат изяснени. И не на последно място, като обещаваща иновация и възможност за предприемаческа активност – неврологичните подобрения – технологии за усъвършенстване на когнитивните способности, четящи и дешифриращи мозъчната активност. Освен като възможност за подобряване на човешките умения, персонализирано обучение и подобрения в областта на безопасността, тази технология би могла да намери приложение за подкрепа на естествената издръжливост на някои професии, например лекари. Отново предизвикателно тук са етичните аспекти – автономия, неприкосновеност - и изграждането на законодателна основа, редом с изграждането на дигитални навици и задоволително ниво на дигитална грамотност (Yuste et al., 2017).

Дискусия

Представените тенденции показват, че дигиталната трансформация не само поражда технологични иновации, но и изисква дълбока социална и организационна промяна. Все по-широкото използване на автономни AI агенти променя ролята на човека в работния процес: от активен изпълнител към определящ цели, наблюдаващ и оценяващ резултатите (Kaplan & Haenlein, 2019). Това води до формирането на нови професионални профили, при които стратегическото мислене, етичната преценка и разбирането на системите са ключови компетентности. В същото време предизвикателство остава постигането на пълно съответствие между човешките намерения и алгоритмичните действия, особено в сложни и динамични среди. Необходимостта от платформи за отговорно управление на изкуствения интелект се усилва от рисковете, свързани с непрозрачност, алгоритмични пристрастия и дезинформация (Floridi & Cowls, 2019). Прозрачността и отчетността са критични фактори за изграждане на доверие между организациите, потребителите и обществото. Това предполага не само технологични решения, но и регулаторни рамки, като например Европейския акт за изкуствения интелект (European Commission, 2024), които да гарантират, че AI системите работят безопасно, справедливо и проследимо. Паралелно с това, стремежът към енергоефективни и хибридни

изчислителни архитектури подчертава растящото значение на устойчивостта в контекста на цифровата трансформация. Организациите са принудени да съчетават екологична отговорност с икономическа ефективност (NIST, 2023). Това изисква дългосрочни инвестиции в инфраструктура, обучение и нови подходи към стратегическото управление на информационните ресурси. Технологии като пространствените изчисления (XR) и многофункционалните роботи свидетелстват, че цифровите системи навлизат все по-дълбоко във физическата и социалната среда на човека. Макар да откриват нови възможности в образованието, здравеопазването и социалните грижи, те поставят и етични въпроси, свързани с автономията, човешките взаимоотношения и ролята на грижата (Sharkey & Sharkey, 2012). Особено чувствителни са невротехнологичните приложения, позволяващи директно взаимодействие между мозъчната активност и цифрови системи. Те предоставят възможности за подобряване на когнитивните способности или за подпомагане на хора в натоварени професионални дейности, но същевременно повдигат въпроси за личната автономия, собствеността върху данните и идентичността (Yuste et al., 2017).

Заклучение и перспективи

В обобщение може да се заключи, че технологичният прогрес може да бъде устойчив и обществено приемлив само ако технологичната иновация бъде съчетана с етична рефлексия, ясна регулация и развитие на дигитални компетентности. Дигиталната трансформация следователно не е единствено технологичен процес, а цялостна културна промяна. Дигиталната трансформация не се изчерпва с избора и внедряването на технологичните иновации. Тя се състои и от елементи като стратегическа и организационна промяна на ниво култура, нагласи, манталитет и организационно поведение. Различните отрасли показват различни нива на зрялост и са изправени пред различни предизвикателства, но общите фактори за успех са еднозначни: **ясна визия, управленски решения, основани на данни, гъвкави структури и дигитално мислене и действие**. Бъдещите изследвания трябва да се фокусират повече върху **междусекторни сравнителни изследвания, измеримост на нематериални фактори** (напр. култура, лидерство) и **трансформация, ориентирана към устойчивост**. Още днес се смята, че следващите поколения GPT¹ ще умеят да заместят интегрирани и сложни дейности, някои от тях за секунди. Компаниите, форсиращи GPT изтъкват, че GPT 5 ще умее да създава приложения за секунди, GPT 6 – да ръководи компании, а GPT 8 – да лекува рак. Това са част от предизвестените тенденции на бъдещето, които ще бъдат и във фокуса на академичните среди. На практика, **интеграцията на изкуствен интелект, облачните изчисления и платформите с ниско кодиране** ще определят следващите нива на цифрова трансформация – към самообучаващи се, адаптивни и автономни организации.

Литература

1. Brynjolfsson, E., & McAfee, A. (2014). *The second machine age: Work, progress, and prosperity in a time of brilliant technologies*. W. W. Norton & Company.
2. Bharadwaj, A., El Sawy, O. A., Pavlou, P. A., & Venkatraman, N. (2013). Digital business strategy: Toward a next generation of insights. *MIS Quarterly*, 37(2), 471–482.
3. Cosa, M. (2024). Business digital transformation: Strategy adaptation, communication and future agenda. *Journal of Strategy and Management*, 17(2), 244–259.
4. Davenport, T., & Ronanki, R. (2018). Artificial intelligence for the real world. *Harvard Business Review*, 96(1), 108–116.

¹ What is GPT - generative pretrained transformer? <https://www.ibm.com/think/topics/gpt>, 01.10.20025

5. European Commission. (2024). *EU Artificial Intelligence Act*. Publications Office of the European Union.
6. Floridi, L., & Cowls, J. (2019). A unified framework of five principles for AI in society. *Harvard Data Science Review*, 1(1), 1–19.
7. Gartner. (2023). Digital transformation and leadership readiness report 2023. Gartner Research.
8. Kagermann, H., Wahlster, W., & Helbig, J. (2018). Industrie 4.0: Mit dem Internet der Dinge auf dem Weg zur 4. industriellen Revolution. acatech.
9. Kane, G. C., Palmer, D., Phillips, A. N., Kiron, D., & Buckley, N. (2015). Strategy, not technology, drives digital transformation. MIT Sloan Management Review and Deloitte University Press.
10. Kane, G., Palmer, D., Phillips, A. N., Kiron, D., & Buckley, N. (2016). Aligning the Organization for its Digital Future. Retrieved from
11. Kaplan, A., & Haenlein, M. (2019). Siri, Siri, in my hand: Who's the fairest in the land? On the interpretations, illustrations, and implications of artificial intelligence. *Business Horizons*, 62(1), 15–25.
12. Milgram, P., & Kishino, F. (1994). A taxonomy of mixed reality visual displays. *IEICE Transactions on Information and Systems*, 77(12), 1321–1329.
13. Mitroulis, D., & Kitsios, F. (2023). Digital transformation and digital maturity models. In *Handbook of Research on Artificial Intelligence, Innovation and Entrepreneurship*. Edward Elgar Publishing.
14. NIST. (2023). Energy-efficient cloud computing: Guidance and metrics. National Institute of Standards and Technology.
15. OECD. (2023). Digital Government Review of Germany: Accelerating the Digital Transformation of the Public Sector. OECD Publishing.
16. Schwer, Karlheinz, Christian Hitz, Robin Wyss, Dominik Wirz, and Clemente Minonne. 2018. “Digital Maturity Variables and Their Impact on the Enterprise Architecture Layers.” *Problems and Perspectives in Management* 16 (4): 141–54.
17. Sharkey, A., & Sharkey, N. (2012). Granny and the robots: Ethical issues in robot care for the elderly. *Ethics and Information Technology*, 14(1), 27–40.
18. Teichert, R. (2019). Digital Transformation Maturity: A Systematic Review of Literature. *Acta Universitatis Agriculturae et Silviculturae Mendelianae Brunensis*, 67(6), 1673–1687.
19. Vaska, S., Massaro, M., Bagarotto, E. M., & Dal Mas, F. (2021). The digital transformation of business model innovation: A structured literature review. *Frontiers in Psychology*, 11.
20. Vial, G. (2019). Understanding digital transformation: A review and a research agenda. *The Journal of Strategic Information Systems*, 28(2), 118–144.
21. Warner, K. & Wäger, M. (2019). Building dynamic capabilities for digital transformation: An ongoing process of strategic renewal. *Long Range Planning*, 52(3), 326–349.
22. Westerman, G., Bonnet, D., & McAfee, A. (2014). Leading digital: Turning technology into business transformation. Harvard Business Review Press.
23. Xu, J., Naseer, H., Maynard, S., & Fillipou, J. (2022). Leveraging Data and Analytics for Digital Business Transformation through DataOps. arXiv preprint.
24. Yuste, R., Goering, S., Bi, G., Carmena, J. M., Carter, A., Fins, J. J., & Wolpaw, J. (2017). Four ethical priorities for neurotechnologies and AI. *Nature*, 551(7679), 159–163.

25. Zimmermann, A., Bogner, J., Jugel, D., & Schweda, C. (2016). Digital enterprise architecture with micro-granular systems and services. Retrieved from <https://publikationen.reutlingen-university.de/frontdoor/index/index/docId/1170>

ASSESSMENT AND MANAGEMENT OF TECHNOLOGICAL RISKS IN THE DIGITALIZATION PROCESS OF ACADEMIC INSTITUTIONS

Ivona Velkova¹,

e-mail: ivonavelkova@unwe.bg¹,

Абстракт

Дигитализацията в академичните институции предоставя възможности за по-голяма ефективност, подобрена достъпност и оптимизирано управление на образователните процеси. В същото време тя поражда редица технологични рискове, които могат да възпрепятстват институционалните дейности, да компрометират сигурността на данните и да нарушат непрекъснатостта на учебния и научноизследователския процес. Настоящото изследване представя систематичен обзор на основните технологични рискове в академичната среда, включително нарушения на сигурността и поверителността на данните, системни откази и прекъсвания на услуги, както и остаряване на технологиите. Предлага се методологична рамка за идентифициране, оценка и приоритизиране на тези рискове чрез използване на матрица „вероятност – въздействие – готовност“. Примери от университетската практика и международни стандарти (ENISA, NCSC, ISO 27001, GDPR) се използват за илюстриране на често срещани уязвимости и подходи за управление. Анализът подчертава повтарящи се модели като висока зависимост от външни доставчици, недостатъчна готовност за възстановяване при бедствия и рискове, свързани с остарели технологии. Резултатите акцентират върху необходимостта от систематична оценка и приоритизация на рисковете, което ще позволи на академичните институции да укрепят своята устойчивост и да осигурят сигурна и ефективна дигитализация.

Abstract

Digitalization in academic institutions provides opportunities for greater efficiency, improved accessibility, and optimized management of educational processes. At the same time, it generates a range of technological risks that may hinder institutional operations, compromise data security, and disrupt the continuity of teaching and research activities. This study presents a systematic overview of the main technological risks in the academic environment, including data security and privacy breaches, system failures and service interruptions, as well as technology obsolescence. A methodological framework is proposed for the identification, assessment, and prioritization of these risks through the use of a likelihood - impact - readiness matrix. Examples from university practice and international standards (ENISA, NCSC, ISO 27001, GDPR) are included to illustrate common vulnerabilities and management approaches. The analysis highlights recurring patterns such as high dependency on external providers, insufficient disaster recovery preparedness, and risks associated with outdated technologies. The results emphasize the need for systematic risk assessment and prioritization, enabling academic institutions to strengthen their resilience and support secure and effective digitalization.

Keywords: digitalization, technological risks, cybersecurity, academic institutions

JEL: I2, I23, I24, O33

Introduction

Digitalization has become one of the driving forces behind the transformation of higher education. Universities depend on digital systems in almost every area of their activity, from administration and teaching support to virtual classrooms and research management. Thanks to this shift, universities have gained broader access to information and more agile ways of working, while cooperation between

¹ Гл. ас. д-р в катедра ИТК, УНСС, e-mail: ivonavelkova@unwe.bg

disciplines and regions has become smoother and more natural. Still, the growing dependence on technology also carries new risks. System failures, data breaches, or poorly secured platforms can quickly disrupt academic life, threaten the continuity of education, and erode institutional trust and reputation [1].

At the core of this digital transformation are the technologies that keep academic life running - hardware and software, learning management systems (LMS), student information systems (SIS), research databases, cloud services, and a growing range of digital tools. These systems make modern universities possible, yet they also come with their own risks. Cyberattacks, system failures, outdated components, or poor integration can all interrupt daily operations and expose critical data [2].

The challenge has grown as new digital tools continue to appear, each requiring fresh investment, training, and technical adaptation. The links between systems make things even more complicated. When an LMS platform connects with identity management or national research networks, a single weak point can trigger a chain reaction across multiple services [3]. Many universities also face another constraint - limited funding and a shortage of skilled IT professionals - which makes it harder to maintain strong security and lasting technological resilience [2].

In recent years, cybersecurity agencies such as ENISA have pointed out that higher education remains one of the sectors most at risk. Large data breaches, ransomware incidents, and system outages during online exams have shown how quickly technical problems can escalate - from temporary disruptions to serious breaches of trust between institutions, students, and staff [4].

That's why a systematic look at technological risks in academia is so important. Understanding these challenges means not only identifying specific threats but also seeing how they connect. For instance, an outdated system may cause compatibility problems that make a network more vulnerable to attack. A comprehensive analysis should therefore bring together theory, real-world cases, and cross-institutional comparisons to reveal these relationships.

This paper focuses on how technological risks in higher education can be assessed and managed more effectively. It aims to outline the key risk factors, explain their causes and effects, and point to practical ways to address them. Using examples from university experience and established frameworks such as the GDPR, ISO/IEC 27001, and the NIST Cybersecurity Framework, the study links conceptual understanding with hands-on strategies for building stronger digital resilience in higher education.

Methodology

The paper adopts a structured and systematic methodology to identify, assess, and prioritize technological risks in the digitalization of higher education. The analysis focuses on the technological dimension of institutional risk, deliberately excluding organizational, financial, and human factors to maintain methodological precision. It follows an exploratory-descriptive design that combines qualitative analysis of secondary data such as incident reports, institutional policies, and international frameworks - with an analytical model classifying risks by frequency, impact, and institutional readiness. The resulting framework bridges conceptual understanding with practical evaluation and reflects the realities of digital transformation across academia.

The methodology draws on several internationally recognized sources, including the ENISA Threat Landscape 2023 [4], the EDUCAUSE Horizon Report 2022 [5], ISO/IEC 27001 [6], the NIST Cybersecurity Framework [7], and the principles of the General Data Protection Regulation (GDPR) [8]. These standards and reports were selected because they provide widely accepted criteria for understanding and managing technological risks in the public and educational sectors. They also ensure consistency and comparability with existing research and policy practices.

Risk Matrix Construction and Data Sources

To assess the technological risks affecting academic institutions, this study employed a likelihood - impact - readiness matrix as a core analytical tool. The matrix was designed to synthesize information from multiple qualitative and quantitative sources, enabling a balanced and transparent evaluation of risk exposure.

1. Data Sources

The construction of the matrix relied on three primary sources of evidence:

- International threat intelligence and sector reports – including ENISA Threat Landscape for Education 2023 [4], EDUCAUSE Horizon Report 2022 [5], and Jisc Cybersecurity Posture of UK Higher Education 2022 [9], which provided sector-wide statistics on incident frequency, attack types, and recurring vulnerabilities.
- Documented university case studies – such as the University of Manchester data breach (2023) [10], the University of Waterloo ransomware attack (2023) [11], and the cybersecurity incident at Mount Saint Mary College in 2022 [12]. These examples provide concrete evidence of how cybersecurity incidents in higher education can affect institutions financially, operationally, and reputationally.
- Regulatory and technical standards – namely ISO/IEC 27005 (Information Security Risk Management) [6], ISO 31000 (Risk Management Guidelines) [13], and NIST SP 800-30 (Guide for Conducting Risk Assessments) [14]. These frameworks provided the conceptual foundation for defining scales and evaluation criteria.

2. Evaluation Logic

Each risk category was assessed along three dimensions:

- Likelihood: determined from the frequency and recurrence of incidents reported in ENISA, EDUCAUSE, and Jisc data. Risks appearing in two or more independent sources were rated High; those mentioned sporadically were Medium; and isolated or hypothetical risks were Low.
- Impact: assessed using case study evidence on the extent of service disruption, cost of recovery, data loss magnitude, and reputational damage. Severe multi-system or regulatory consequences (GDPR fines) were rated Critical; moderate single-system disruptions High; and limited local effects Medium.
- Readiness: evaluated through the presence of preventive and response measures observed in university audits and policy documents (MFA enforcement, backup testing frequency, DPO appointment). Institutions exhibiting multiple active controls were considered High readiness; minimal or reactive measures Low.

3. Integration and Scoring

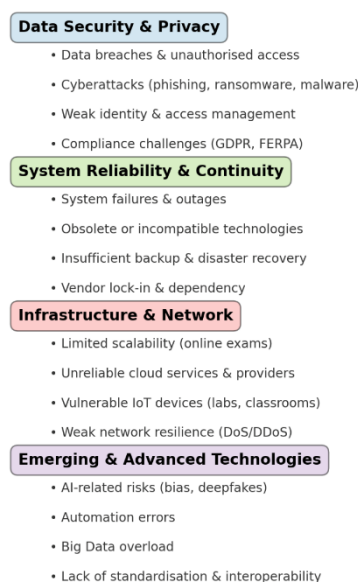
The three dimensions were combined using a qualitative matrix model, where risk priority equals the combined weight of likelihood and impact, adjusted by readiness. For instance, a High-likelihood and High-impact risk with Low readiness was classified as Critical, while the same risk with High readiness was downgraded to High. This weighting ensures that strong institutional preparedness can mitigate otherwise severe risks.

4. Validation

The resulting matrix was informed by international benchmarks published by ENISA and EDUCAUSE to maintain consistency with sector-level perspectives. While the assessment remains qualitative, drawing on multiple data sources supports its reliability and comparability across institutions.

The process of identifying risks began with an analytical review of documented cybersecurity incidents and sector reports related to higher education. Common patterns were derived from international analyses and institutional experiences described in publicly available sources. Each identified risk was then examined according to its technological characteristics and its relevance to the academic environment. From this analysis, four major categories were defined: data security and privacy, system reliability and continuity, infrastructure and network, and emerging and advanced technologies. These categories represent the main technological domains of higher education most frequently exposed to disruption. Selected institutional incidents, referenced in the accompanying matrix, illustrate how these risk categories manifest in practice across different universities.

The classification of these risk domains represents the author’s analytical synthesis, developed by comparing existing frameworks and academic sources. While ENISA, EDUCAUSE, and ISO/IEC 27001 outline various typologies of digital risk, this study integrates their principles into a unified model tailored to the context of higher education. This combined approach allows for a focused assessment of the technological dimension of risk while maintaining consistency with internationally recognized standards.



Source: Author's research

Figure 1. Classification of technological risks in Academia

Figure 1 gives an overview of how technological risks in higher education are grouped into four main areas. It also provides the basis for the next section, which explores each category through its key causes, resulting impacts, and ways these risks can be mitigated.

To evaluate and prioritize the identified risks, the analysis employs the previously described matrix. This framework consolidates the findings from the data review and enables a structured comparison across risk categories. By visualizing the combined dimensions of frequency, potential impact, and institutional preparedness, the matrix indicates which risks may warrant immediate strategic attention and which are more suitable for gradual capacity development. The resulting assessment provides the basis for the comparative and analytical discussion in the following sections.

Finally, the analysis compares findings from various universities and international studies to identify recurring challenges and interdependencies. The comparison reveals patterns such as strong dependence on external service providers, outdated technologies, and limited investment in disaster recovery. These insights form the foundation for the subsequent sections, which examine each risk category in greater detail and propose strategies for strengthening technological resilience in higher education.

Analysis of technological risks in Academia

The digital transformation of higher education has reshaped the way universities function, integrating teaching, research, and administration into a single interconnected digital ecosystem. While this transformation enables innovation and operational efficiency, it also introduces new layers of technological risk that can compromise data integrity, interrupt academic services, and undermine institutional reputation. According to recent international analyses of cybersecurity in higher education, the sector remains among the most frequently targeted by cyber threats due to the high value of academic data, the decentralization of IT infrastructures, and often constrained cybersecurity budgets.

Building on the methodological framework presented earlier, this section examines four main categories of technological risks that define the academic digital landscape: data security and privacy, system reliability and continuity, infrastructure and network, and emerging and advanced technologies.

Data security and privacy

Universities manage vast repositories of sensitive data - including student information, research outputs, and administrative records - making them attractive targets for cyberattacks. Data security and privacy violations are the most critical category of technological risk, given their potential to disrupt operations, violate regulations such as the GDPR, and erode institutional trust.

Table 1. Key indicators and evaluation of data security and privacy risks

Subcategory	Key Indicators	Likelihood	Impact	Priority
Data breaches and unauthorized access	Number of incidents per year; time to detect and recover; volume of exposed records	High	Critical	Critical
Cyberattacks (phishing, ransomware, malware)	Frequency of attacks; downtime duration; recovery cost	High	High	High
Weak identity & access management (IAM)	MFA adoption; inactive accounts; privilege misuse	Medium-High	High	High
Compliance challenges (GDPR, ISO/FERPA)	Number of data protection incidents; regulator notifications; vendor compliance rate	Medium	Medium-High	Moderate-High

Source: Author's synthesis based on multiple international frameworks and reports (ENISA, EDUCAUSE, NIST, ISO 27001, GDPR)

Data security breaches and ransomware attacks represent the most severe risks for academic institutions. Their likelihood is consistently high due to the openness of academic networks and the diversity of users, while the impact includes service disruption, reputational loss, and legal penalties. Weak identity management and fragmented compliance practices further increase exposure.

Effective mitigation requires a layered defense model combining technical controls (encryption, MFA), organizational measures (Data Protection Officer, incident response planning), and cultural interventions (awareness training). Integrating these measures under GDPR, ISO/IEC 27001, and the NIST Cybersecurity Framework significantly enhances institutional readiness.

System reliability and continuity

Academic operations depend on the continuous availability of digital systems supporting teaching, research, and administration. Failures in this area—caused by aging infrastructure, insufficient redundancy, or vendor dependency—can lead to widespread disruption. Table 2 outlines the main subcategories of system reliability risks, together with key indicators, likelihood, and relative priority, as identified in the analytical framework.

Table 2. System reliability and continuity risks

Subcategory	Key Indicators	Likelihood	Impact	Priority
System failures and outages	Number of critical service interruptions; uptime percentage	High	High	Critical
Obsolete or incompatible technologies	% of unsupported systems; integration failures	Medium-High	High	High
Insufficient backup and recovery	Backup frequency; test success rate; RPO/RTO metrics	Medium	High	High
Vendor dependency	% of outsourced critical services; strength of SLAs	Medium	Medium-High	Moderate-High

Source: Author's synthesis based on multiple international frameworks and reports (ENISA, EDUCAUSE, NIST, ISO 27001, GDPR)

In higher education, system reliability risks often surface only during major service disruptions. Evidence from institutional reports indicates that limited recovery testing and outdated technologies exacerbate the duration and cost of outages. To address these vulnerabilities, institutions should adopt redundant architectures, conduct structured recovery exercises, and align continuity objectives (RTO/RPO) with established standards such as ISO 27031 and ITIL Service Operations.

Infrastructure and network

The backbone of digital academia lies in its infrastructure: cloud platforms, network systems, IoT devices, and scalable connectivity. While these components enable efficiency and flexibility, they also introduce new vulnerabilities.

Table 3. Indicators and assessment of infrastructure & network risks

Subcategory	Key Indicators	Likelihood	Impact	Priority
Limited scalability (online exams)	Maximum concurrent capacity; response time under stress	High	High	Critical
Unreliable cloud services and providers	Frequency of outages; provider diversification	Medium-High	High	High
Vulnerable IoT devices	Unauthenticated devices; detected vulnerabilities	Medium	Medium-High	Moderate-High
Weak network resilience (DoS/DDoS)	Number of incidents per year; mitigation capability	High	Medium-High	High

Source: Author's synthesis based on international cybersecurity and infrastructure resilience frameworks (ENISA Cloud Security Guidelines, NIST Cybersecurity Framework, ISO/IEC 27036).

Infrastructure failures - especially during online examinations or registration peaks - have direct academic and operational consequences. Cloud dependency and the growing number of IoT devices expand the attack surface, while DDoS attacks remain among the most frequent disruptions.

Mitigation requires proactive capacity planning, multi-cloud diversification, IoT segmentation, and network redundancy. Documented incidents in higher education demonstrate that insufficient scalability or insecure endpoints can rapidly escalate into institution-wide service failures. International standards such as ISO/IEC 27036, the ENISA Cloud Security Guidelines, and the NIST Cybersecurity Framework offer structured guidance for strengthening technological resilience.

Emerging and advanced technologies

The rapid deployment of AI, automation, and big data systems in academia offers significant opportunities for innovation but also introduces new categories of technological risk. AI-driven integrity challenges, automation errors in assessment platforms, and inadequate data governance can undermine institutional integrity and operational efficiency.

Table 4. Indicators and assessment of emerging technology risks

Subcategory	Key Indicators	Likelihood	Impact	Priority
AI-related misuse	% of flagged submissions; integrity violations	High	High	Critical
Automation errors	Number of process failures; manual intervention rate	Medium-High	High	High
Big data overload	Data utilization rate; storage-to-processing ratio	Medium	Medium-High	Moderate-High
Lack of interoperability	Integration failures; % of systems using open standards	Medium	High	High

Source: Author's synthesis based on principles from the UNESCO Recommendation on the Ethics of AI (2021) [15], the EU AI Act (draft, 2024) [16], and ISO/IEC 38505 (Data Governance) [17].

The integration of emerging technologies in higher education introduces both opportunities and vulnerabilities. Institutions face growing challenges related to AI-driven academic integrity, automation reliability, and the governance of large and complex datasets. In many cases, these risks stem from rapid adoption without adequate oversight, ethical frameworks, or interoperability standards.

Effective mitigation requires a balanced approach that combines innovation with governance: establishing AI ethics policies, ensuring human oversight in automated processes, and implementing strong data lifecycle management practices.

Comparative evaluation and discussion

Looking across all four domains, it becomes clear that technological risks in academia rarely act in isolation. System failure can expose data vulnerabilities, while inadequate governance over emerging technologies can intensify both security and reliability concerns. These patterns show that universities face a web of interconnected risks rather than discrete problems. Table 5 summarizes this relationship by comparing likelihood, impact, and priority across each category.

Table 5. Consolidated comparative risk matrix for academic institutions

Risk Category	Likelihood	Impact	Priority	Justification
Data Security and Privacy (breaches, cyberattacks, IAM, compliance)	High	Critical	Critical	Real-world incidents, such as the 2023 data breach at the University of Manchester, the ransomware attack at the University of Waterloo (2023), and the cybersecurity incident at Mount Saint Mary College (2022), demonstrate the tangible consequences of cybersecurity failures in higher education. These events often result in GDPR penalties, reputational loss, and disruption of teaching and research, making this the most visible and damaging risk category.
System Reliability and Continuity (outages, backup failures, vendor dependency)	Medium–High	High	High	Failures during exams or enrolment disrupt critical operations. These are often caused by underfunded IT systems, weak continuity planning, and heavy reliance on external service providers.
Infrastructure and Networks (scalability, cloud outages, IoT insecurity, DDoS)	Medium	High	High	Attacks and failures propagate rapidly across interconnected systems. IoT vulnerabilities and cloud dependencies create systemic risks, particularly during high-demand periods.
Emerging and Advanced Technologies (AI misuse, automation errors, data overload, lack of standards)	Medium–High	High	High	Misuse of AI tools is increasingly common, while automation errors and interoperability gaps undermine academic credibility and long-term sustainability. The impact grows as adoption accelerates.

Source: Author's synthesis based on comparative analysis of institutional risk domains and international cybersecurity frameworks.

Among all categories, Data Security and Privacy remains the most critical due to its high frequency and severe consequences. However, the remaining domains: System Reliability, Infrastructure and Networks, and Emerging Technologies - are closely interconnected. For example, outdated infrastructure can exacerbate data breaches, while weak governance over AI systems can amplify both privacy and integrity risks.

These interconnections highlight the need for a holistic approach to technological risk management - one that integrates cybersecurity, infrastructure resilience, and responsible innovation governance. A reactive response to individual incidents is no longer sufficient. Instead, universities should cultivate institutional digital resilience grounded in prevention, continuous monitoring, and the ability to recover rapidly from technological disruptions.

Recommendations

Strengthening technological resilience in higher education requires a coordinated and long-term strategy that connects governance, infrastructure, and human capacity. Based on the findings of this study, several key directions can guide universities in managing technological risks more effectively.

A first priority is to establish an integrated risk management framework that unites institutional policies, technical standards, and operational practices. Universities should align their internal governance with international models such as ISO 31000, ISO/IEC 27001, and the NIST Cybersecurity Framework. This alignment ensures that risk identification, assessment, and response activities follow a consistent logic across departments and campuses. It also helps transform cybersecurity from a reactive function into a continuous management process.

Equally important is the need to enhance data governance and compliance mechanisms. Academic institutions handle sensitive research outputs and personal information that demand rigorous protection under regulations such as the GDPR. Implementing privacy-by-design principles, conducting periodic audits, and maintaining clear accountability structures are essential steps. Establishing dedicated Data Protection Officers and cross-departmental privacy committees can further embed compliance within the university culture rather than treating it as an external requirement.

A third strategic area involves system reliability and continuity planning. The study revealed that many universities still rely on legacy systems and ad-hoc recovery procedures. Institutions should define measurable recovery time and recovery point objectives (RTO/RPO), test them regularly through simulation exercises, and secure sufficient financial resources for infrastructure renewal. Emphasizing redundancy, automated backups, and cloud failover solutions can minimize downtime and preserve academic operations during crises.

Another crucial step is to invest in scalable and secure infrastructure. As demand for digital services expands, universities must ensure that their network capacity, cloud architecture, and IoT devices can adapt without compromising security. Multi-cloud diversification, continuous network monitoring, and segmentation of connected devices can substantially reduce the impact of targeted attacks and technical failures.

The rapid adoption of artificial intelligence, analytics, and automation also calls for ethical and responsible innovation practices. Universities should develop clear institutional policies on AI use in research and teaching, promote algorithmic transparency, and encourage human oversight in automated decision-making. Embedding AI literacy and digital ethics into academic curricula will help prepare students and staff to navigate emerging technological challenges responsibly.

Beyond technical measures, technological resilience depends on people as much as on systems. Building a security-first organizational culture requires continuous awareness programs, targeted training, and the encouragement of responsible digital behavior among both students and staff. Simulated phishing campaigns, peer-learning workshops, and the recognition of secure practices can gradually transform cybersecurity from individual responsibility into a shared institutional value.

Taken as a whole, these recommendations point toward a more resilient and ethically aware digital future for higher education. Technology in this context should not be viewed only as a set of systems

or tools, but as a living part of the institution's culture and decision-making. When universities see technology as something that connects people, ideas, and responsibilities, they build the capacity not only to respond to crises but also to evolve with them - anticipating change, adapting to it, and emerging stronger each time.

Conclusion

Digital transformation in academia has created both opportunities and vulnerabilities. This paper demonstrates that while digital tools enhance efficiency and access, they also amplify exposure to complex technological risks. The findings confirm that universities face an intertwined network of challenges - where weaknesses in data security, infrastructure, and system reliability can reinforce one another, magnifying institutional exposure.

The likelihood-impact-readiness framework proved effective for mapping these interdependencies and for prioritizing actions according to institutional preparedness. The analysis revealed that the most pressing risks stem from data security and privacy breaches, outdated infrastructure, and the rapid, often unregulated adoption of emerging technologies such as artificial intelligence. These factors together highlight the urgent need for integrated governance mechanisms and continuous monitoring to maintain academic continuity and trust.

Strengthening technological resilience, therefore, requires more than technical upgrades. It calls for strategic alignment between IT, policy, and education; sustainable funding for infrastructure renewal; and the promotion of digital ethics across teaching and research. Universities that embed these principles into their long-term strategies will be better equipped not only to prevent incidents but also to adapt and evolve through future disruptions.

Looking ahead, further research could deepen this model by quantifying risk readiness or comparing maturity levels across national contexts. Such insights would help policymakers and university leaders design more targeted, evidence-based interventions that support secure, inclusive, and sustainable digital transformation in higher education.

Acknowledgement

The paper is prepared as a part of the research activities under the project Models and Instruments for Transforming Higher Education System through Transnational Multi-Sector Links - STUDIES – DIG (Project N: 101131544), funded by the European Union.

References

1. H. Schuetze, W. de Vries, and G. Alvarez Mendiola, "Digitalization of Higher Education," *J. Comp. Int. High. Educ.*, vol. 16, no. 2, pp. 6–12.
2. В. Андонов, „Софтуерни средства за дигитализиране на основните процеси и услуги във висшите училища на Република България,“ *Икономически и социални алтернативи*, т. 29, бр. 4, с. 86–94, 2023.
(Andonov, V., "Software tools for digitalizing the core processes and services in Bulgarian universities," *Economic and Social Alternatives*, vol. 29, no. 4, pp. 86–94, 2023.)
3. European Union Agency for Cybersecurity, *Identifying Emerging Cybersecurity Threats and Challenges for 2030*, LU: Publications Office, 2023. [Online]. Available: <https://data.europa.eu/doi/10.2824/117542>
4. ENISA, *ENISA Threat Landscape 2023*. [Online]. Available: <https://www.enisa.europa.eu/publications/enisa-threat-landscape-2023>
5. K. Pelletier, M. McCormack, J. Reeves, J. Robert, and N. Arbino, 2022 *EDUCAUSE Horizon Report: Teaching and Learning Edition*. EDUCAUSE, Boulder, CO, USA, 2022. [Online]. Available: <http://www.educause.edu>

6. International Organization for Standardization, ISO/IEC 27001:2022 – Information Security, Cybersecurity and Privacy Protection — Information Security Management Systems — Requirements. [Online]. Available: <https://www.iso.org/obp/ui/en/#iso:std:iso-iec:27001:ed-3:v1:en>
7. National Institute of Standards and Technology, The NIST Cybersecurity Framework (CSF) 2.0, Gaithersburg, MD, NIST CSWP 29, Feb. 2024. doi: 10.6028/NIST.CSWP.29.
8. European Commission, Principles of the GDPR. [Online]. Available: https://commission.europa.eu/law/law-topic/data-protection/rules-business-and-organisations/principles-gdpr_en
9. Jisc, Cyber Impact Report 2022. [Online]. Available: <https://repository.jisc.ac.uk/8732/1/cyber-impact-report-2022.pdf>
10. University of Manchester, “Cyber Incident Statement,” 2023. [Online]. Available: <https://www.manchester.ac.uk/about/news/cyber-incident-statement/>
11. University of Waterloo, Daily Bulletin – June 6, 2023. [Online]. Available: <https://uwaterloo.ca/daily-bulletin/2023-06-06>
12. Mount Saint Mary College, “Cybersecurity Incident Statement,” 2022. [Online]. Available: <https://www.msmc.edu/newsroom/news/cybersecurity-incident-at-msmc/>
13. International Organization for Standardization, ISO 31000:2018 – Risk Management – Guidelines. [Online]. Available: <https://www.iso.org/standard/65694.html>
14. Joint Task Force Transformation Initiative, Guide for Conducting Risk Assessments, NIST SP 800-30r1, Gaithersburg, MD, USA, 2012. doi: 10.6028/NIST.SP.800-30r1.
15. UNESCO, Recommendation on the Ethics of Artificial Intelligence, 2021. [Online]. Available: <https://www.unesco.org/en/articles/recommendation-ethics-artificial-intelligence>
16. European Union, The European Union Artificial Intelligence Act (Draft), 2024. [Online]. Available: <https://www.ey.com/content/dam/ey-unified-site/ey-com/en-gl/insights/public-policy/documents/ey-gl-eu-ai-act-07-2024.pdf>
17. International Organization for Standardization, ISO/IEC 38505-1:2017 – Governance of IT – Governance of Data – Part 1: Application of ISO/IEC 38500 to the Governance of Data. [Online]. Available: <https://www.iso.org/standard/56639.html>

SOFTWARE TOOLS FOR ONLINE ANALYTICAL PROCESSING- FROM DATA TO STRATEGIC DECISIONS IN THE DIGITAL ECONOMY

Софтуерни средства за онлайн аналитична обработка- от данни към
стратегически решения в дигиталната икономика

Marinela Todorova¹

email: mtodorova_2520823@unwe.bg

Abstract

Digitalization in the modern world is causing fundamental transformations in the economy as a whole. For this reason, the digital economy is not just a trend, but an inevitable reality that determines the competitiveness and development of businesses over time. Data is becoming an increasingly strategic resource, allowing organizations to optimize their processes and offer personalized solutions, and software tools for online analytical processing are key tools for their effective management and analysis. The report examines the evolution of software tools for online analytical processing - from traditional to modern cloud platforms with artificial intelligence integration. Emphasis is placed on defining the criteria for selecting an appropriate solution for online analytical processing according to the size and needs of enterprises, and a model of logically related indicators is proposed, including: functional, technological and economic aspects of implementation. Based on this, a comparative analysis of the leading products for online analytical processing is carried out, selected due to their widespread distribution, different approach to data processing and market leadership in the BI industry. The comparison is made based on predefined criteria such as scalability, ease of integration, analytical capabilities, price and accessibility. Trends for the future development of online analytical processing are presented, which transform these systems from analysis tools into an engine of strategic management decisions and a catalyst for the digital transformation of business.

Keywords: digital economy, online analytics software, business intelligence, data analysis

JEL: O33, C88, L86

Резюме

Дигитализацията в съвременния свят предизвиква фундаментални трансформации в икономиката, като цяло. Поради тази причина, дигиталната икономика не е само тенденция, а неизбежна реалност, която определя конкурентноспособността и развитието на бизнеса във времето. Данните се превръщат във все по-голям стратегически ресурс, позволявайки на организациите да оптимизират процесите си и да предлагат персонални решения, а софтуерните средства за онлайн аналитична обработка са ключови инструменти за тяхното ефективно управление и анализ. Докладът разглежда еволюцията на софтуерните средства за онлайн аналитична обработка- от традиционните до съвременните облачни платформи с интеграция на изкуствен интелект. Акцент е поставен върху дефинирането на критериите за избор на подходящо решение за онлайн аналитична обработка според големината и нуждите на предприятията, като е предложен модел от логически обвързани показатели, включващи: функционални, технологични и икономически аспекти на внедряване. Оповавайки се на тази основа се извършва сравнителен анализ на водещите продукти за онлайн аналитична обработка, избрани поради тяхното широко разпространение, различен подход към обработката на данни и пазарно лидерство в BI индустрията. Сравнението се извършва на база предварително определени критерии като мащабируемост, лесна интеграция, аналитични

¹ Student, Master's degree, Department of Information Technologies and Communications, UNWE, e-mail: mtodorova_2520823@unwe.bg

възможности, цена и достъпност. Представени са тенденции за бъдещото развитие на онлайн аналитичната обработка, които трансформират тези системи от инструменти за анализ в двигател на стратегически управленски решения и катализатор на цифровата трансформация на бизнеса.

Introduction

The entry of online analytical processing software into the economy is a new beginning in the corporate world. In recent years, online analytical processing software in the digital economy has witnessed an extraordinary transformation in the way companies analyze the data they collect, be it from social networks, customers, markets. The process of implementing OLAP software in the economy, in particular the digital economy, leads to a natural increase in the quality of data processing. The integration of online analytics technologies not only fundamentally transforms the analysis process, but also outlines key roles for a number of improvements in the way data is presented in general. Online analytical processing software is becoming a powerful tool used by enterprises and corporations around the world to present a modernized way of data analysis to ordinary people. Modern online analytical processing software is responsible for how data analysis occurs and its presentation to the participants in the process. Due to the dynamic development of digitalization and the growing demands of business, data analysis needs a profound transformation that will allow online analytical processing software tools to adapt and effectively support management decision-making. This transformation is the result of continuous progress in the field of information and communication technologies, which lay the foundation for more effective and intelligent data analysis solutions.

The essence of software tools for online analytical processing and the digital economy

Software tools for online analytical processing

They are software tools that are responsible for organizing data into multidimensional structures (so-called cubes). They facilitate interactive analysis of the provided data from different perspectives. OLAP software is widely used in BI, serving as a tool for users to extract, summarize and visualize information necessary for making management decisions [1].

Digital economy

It encompasses economic activities carried out through digital means. It is based on the use of the Internet, information systems, digital technologies, which contribute to the transformation of the ways of production, consumption, communication, trade. The digital economy is characterized by high-level automation, global connectivity and innovation. In this way, the digital economy places data and technology at the center of economic value. It changes traditional business models, accelerates innovation and creates new forms of employment and consumption [2].

Evolution of software tools for online analytical processing

When did online analytical processing software tools emerge?

The emergence of online analytical processing software is due to the need for faster and more flexible data analysis for the corporate world. In the late 1970s and early 1980s, the first online analytical processing software appeared, their development was influenced by the concepts of relational databases and multidimensional analysis. This concept allows for data processing according to various criteria such as time, product, region. The first online analytical processing (OLAP) software was Express, developed by Information Resources Inc. (IRI). It uses a multidimensional structure for storing and processing data, which allows for fast aggregation, filtering and visualization. It was initially designed to work on server platforms, and later a version for personal computers appeared. In 1995, Express was acquired by Oracle, which led to the development of Oracle OLAP- a platform combining relational and multidimensional models [6][7][11].

Traditional OLAP systems

They are based on locally installed architectures that use relational databases and multidimensional structures, also called "cubes". These systems provide opportunities for data analysis using various criteria- temporal, categorical, geographical regions and others. Although traditional OLAP systems were effective for their time, they are characterized by limited scalability, high maintenance costs and difficulties in integrating with dynamically changing business environments in the era of the digital economy [11].

OLAP with visualization and interactivity

With the development of graphical user interfaces and the growing needs of the business industry, after the emergence of the digital economy, there is a need for accessible analytical tools. This leads to the emergence of platforms such as Tableau, QlikView, Power BI, etc. These software's offer the possibility of interactive analysis, intuitive data visualization and easy report creation. They integrate with multiple data sources and allow users to extract detailed information without the need for deep technical knowledge. Their functionality is extended through dynamic dashboards, filters and automated reports.

Artificial intelligence integration

The latest stage in the development of software tools for online analytical processing has been reached, including the integration of artificial intelligence and machine learning. Artificial intelligence significantly improves the efficiency of analytical processes. Platforms such as Microsoft Azure Synapse and IBM Watson Analytics provide intelligent recommendations, personalized analysis and anomaly detection that facilitate strategic decision-making in real time. Integration with artificial intelligence allows for a transition from descriptive to predictive and prescriptive analysis.

Where AI changes the functionalities of OLAP software

In the modern digital economy, there are more and more ways to change the functionality of OLAP software. Artificial intelligence (AI) is increasingly changing the way OLAP (online analytical processing) systems operate, transforming them from traditional data analysis tools into intelligent decision-making platforms. Thanks to AI-assisted analysis automation, patterns, dependencies, and anomalies in data are automatically detected, which significantly reduces the need for manual intervention and speeds up the process of extracting useful information. In addition, by applying machine learning (ML), OLAP systems gain the ability not only to analyze past events, but also to predict future trends, even recommending optimal solutions based on accumulated data. An important role in this transformation is also played by natural language processing (NLP), which allows users to interact with the system more intuitively, asking questions in natural language and receiving automatically generated analyses and visualizations. In addition, AI enables intelligent optimization of OLAP processes – indexes and aggregates are dynamically adjusted according to the load, which leads to faster queries and more efficient work with large volumes of data. In addition, systems become more secure, as artificial intelligence allows for timely detection of anomalies, fraud and potential security threats. Ultimately, the integration of artificial intelligence into OLAP environments transforms analytical processes into more intelligent, adaptive and predictive systems that help businesses achieve more accurate, faster and more informed decisions [5][12].

Applications of OLAP software in the digital economy

Online analytical processing (OLAP) software plays a key role in the digital economy, supporting informed decision-making through multidimensional analysis of large volumes of data generated from various sources. One of their main applications is related to the analysis of user behavior in online platforms, where user interactions create a valuable information base for discovering trends and patterns. Through criteria such as device type, traffic source and geolocation, businesses can identify behavioral patterns and personalize content according to the needs of a specific user. Thanks to OLAP cubes, fast data aggregation and filtration is ensured, which is especially important for highly dynamic platforms such as online stores, social networks and streaming services. In addition, OLAP systems

also support the monetization processes of digital services, allowing detailed analysis of the user life cycle and revenues by various criteria - type of service, time period, region and market niche. This allows for optimization of pricing strategies, improvement of revenue forecasts and identification of channels with the highest return. In addition, OLAP software is also important in the optimization of digital marketing campaigns, where data serves as a key tool for evaluating and managing performance. Through in-depth multidimensional analysis, specialists can calculate return on investment, compare results across time periods, platforms, and users, and formulate more precise marketing strategies. In addition, OLAP tools are also actively used in the financial analysis of digital assets, providing the ability to track indicators such as value, transaction volume, volatility, and liquidity. In this way, OLAP software is establishing itself as an indispensable element in data management and analysis in the digital economy, providing a high degree of transparency, efficiency, and strategic predictability.

Criteria for selecting OLAP software based on digital business needs

The criteria for selecting OLAP software in the context of the digital economy are formulated to ensure an optimal match between analytical needs and the scale of the business. They include indicators such as the size of the digital business, scalability, ease of integration, analytical capabilities, price and accessibility, which reflect the key factors for effective data management and analysis. Scalability and integration determine the ability of the system to adapt to growing volumes of data and connect to other corporate systems, while analytical capabilities reveal the degree of depth in information processing - from basic visualizations to predictive analytics and artificial intelligence models. Price and accessibility, in turn, ensure economic sustainability and operational flexibility through cloud and hybrid solutions. These criteria achieve a balanced approach to choosing an OLAP system, ensuring that the organization's technological resources are aligned with its capacity, strategic goals and level of digital maturity [11].

Table 1. Criteria for choosing OLAP software based on the size of the digital business

Digital Business Size	Scalability	Easy integration	Analytical capabilities	Price	Accessibility
Small Digital Business	Low to medium - suitable for limited data volume and users	High- easy integration with popular digital tools	Basic- visualizations, dashboards, reports, KPI analysis	Low- freemium and subscription models available	Very high - access via the cloud without the need for infrastructure
Medium Digital Business	Medium to high - support for growing data volumes and more users	Medium- integration with CRM, ERP and marketing systems	Advanced- drill-down, AI insights, scenario analysis	Medium- subscription plans with flexible upgrades	High - cloud + hybrid solutions, accessible via browser
Large Digital Business	Very high - real-time processing of large and distributed data	High- deep integration with corporate systems (SAP, Oracle, Salesforce, Big Data)	Advanced- predictive analytics, AI/ML models, OLAP cubes, real-time insights	High- enterprise licenses, but with high ROI	Medium - requires technical infrastructure and BI team

Source: Author's research

Logical indicator model

The logical indicator model is chosen because it provides a structured and systematic approach to evaluating OLAP software that allows for simultaneous consideration of functional, technological, and

economic aspects of solutions. It shows the extent to which different platforms meet complex business requirements, including depth and flexibility of analysis, self-service analytics capabilities, scalability, integration with enterprise systems, security, and financial metrics. The use of logical indicators provides an objective framework for comparison and evaluation that supports informed decision-making by showing the strengths and weaknesses of OLAP solutions in the context of different business needs and organizational budgets [9].

Table 2. Logical indicator model

Category	Indicator	Goal of the indicator	Evaluation criteria
Functional	Scope of analytical capabilities;	To assess the depth and flexibility of the analysis	Support for OLAP cubes, drill-down, slice/dice, predictive analysis, AI/ML functions
	Visualization and user interface	To assess the accessibility and intuitiveness of the platform	Availability of drag-and-drop, interactive charts, mobile compatibility
	Self-service support	To assess whether end users can analyze on their own	Ability to create reports without programming, templates, automatic recommendations
Technological	Scalability	To assess the ability to grow and process big data	Support for Big Data, cloud architecture, parallel processing
	Integration with existing systems	To assess compatibility with other platforms	API, connectors to ERP, CRM, databases, cloud services
	Security and data protection	To ensure confidentiality and compliance with standards	Encryption, access control, GDPR, ISO 27001 compliance
Economic	Total implementation costs	To assess the financial investment	Licenses, training, support, infrastructure
	Return on investment (ROI)	To assess the efficiency in relation to the resources invested	Reduced costs, accelerated decision-making, improved productivity
	Pricing model flexibility	To assess the adaptability to the organization's budget	Availability of free versions, subscription plans, pay-as-you-go

Source: Author's research

A figure is also presented that illustrates the hierarchy of indicators in an organization, showing how information is transformed from raw operational data to strategic decisions.

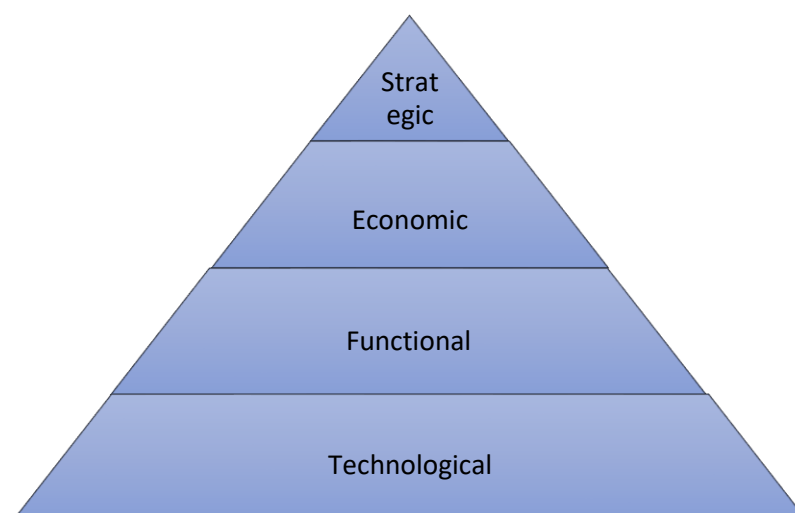


Figure 2. From data to strategic decisions

Comparative analysis of OLAP software

This comparative analysis is designed to provide an objective assessment of leading OLAP solutions across key criteria – data processing, functionality, technology features and cost-effectiveness – and to demonstrate how different platforms meet diverse business needs and scales. The selected solutions are based on Gartner’s Magic Quadrant, which ranks market leaders based on vision and ability to execute, demonstrating robustness, innovation and diverse approaches to visualization, self-service analytics, integration and cloud technologies. The analysis supports informed decision-making by demonstrating best practices and technology approaches adapted to different scales, functional requirements and budget constraints in the digital economy [3][8][10][4].

Table 3. Comparative analysis of OLAP software

OLAP solution	Data processing approach	Functional benefits	Technology Features	Economic aspects
Microsoft Power BI	Combines ROLAP and in-memory analytics via VertiPaq;	Intuitive interface, strong self-service BI, AI Copilot	Cloud Scalability, Easy Integration with Microsoft 365	Low license cost, excellent ROI, high availability
Tableau	full integration with Azure	Highest level of visualizations and drag-and-drop analysis	Integration with Salesforce Data Cloud, Extensions via API	Mid-to-high cost, suitable for medium/large businesses
Google Looker / Looker Studio	In-memory OLAP with a focus on visualization and storytelling	Web-based model, strong integration with Google Data Cloud	Scalability, Security, Big Data Support	Primarily SaaS model, flexible pricing, high efficiency
Qlik Sense	Cloud-native model (ROLAP) via LookML and BigQuery	Automatic data connections, AI insights, storytelling	Multi-Source Support, High Security	Mid-cost, good price-performance ratio
Oracle Analytics Cloud	Associative OLAP (in-memory engine)	Strong predictive and planning features	Integration with Oracle DB, ERP, ML Models	High price, aimed at corporate customers

Source: Author’s research

Trends for future development of OLAP software in the digital economy

In the future, the development of OLAP software in the digital economy will be closely linked to the introduction of new technologies and the need for faster and smarter data analysis. It is expected that the integration with artificial intelligence (AI) and machine learning (ML) will be significantly improved and expanded, which will allow for automatic pattern discovery, behavior prediction and generation of recommendations in real time. Cloud-based OLAP solutions will be increasingly important, providing flexibility, accessibility and scalability, especially for newly created digital businesses that process large volumes of data. In parallel, improvements will be observed in Real-Time OLAP, which will provide the opportunity for immediate analysis of accurate and up-to-date data, supporting automatic reactions to market changes or technical difficulties. User interfaces and visualizations will also be improved, offering intuitive and interactive dashboards, graphs and maps for faster and more effective decision-making. Integration with Big Data technologies will allow the processing and analysis of unstructured data, expanding the analytical capabilities of digital enterprises. Self-Service Analytics solutions will also become increasingly popular, allowing employees to perform analyses and create visualizations without the need for specialized technical knowledge, which will speed up business processes and reduce dependence on IT specialists. In addition, attention will be paid to data security and protection by implementing stricter protection mechanisms adapted to the dynamic development of technologies. The development of voice integration is expected, which, through the use of natural language processing (NLP) technologies, will make data analysis more accessible, intuitive and effective in the modern digital environment.

Conclusion

In the modern digital economy, online analytical processing software tools occupy a key position as strategic tools for data management and decision-making. Their ability to process multidimensional data provides enterprises with the opportunity to perform in-depth analyses of key indicators that are of strategic importance for business development. The introduction of artificial intelligence and machine learning into these systems further expands their potential, enabling automated forecasting, personalized recommendations and anomaly detection - functions that make them indispensable in conditions of dynamically changing markets and increasing competition. A variety of OLAP solutions, adapted to different scales and budgets, demonstrate their applicability in all economic sectors and for enterprises of different sizes. Real-time analytics, cloud architectures and advanced user interfaces contribute to wider access to these technologies, even for users without specialized technical knowledge. As a result, online analytics software tools play a significant role in the transformation of data-driven business models and are establishing themselves as a core component of the digital economy, where information is a key resource for strategic development, sustainability, and innovation.

References

1. Overview of Online Analytical Processing (OLAP)- <https://support.microsoft.com/en-us/office/overview-of-online-analytical-processing-olap-15d2cdde-f70b-4277-b009ed732b75fdd6>
2. DEVELOPMENT OF THE DIGITAL ECONOMY ON A GLOBAL SCALE AND THE EXPERIENCE OF COUNTRIES <https://wosjournals.com/index.php/shokh/article/view/1099>
3. Gartner Research, Gartner Magic Quadrant for Analytics and Business Intelligence Platforms <https://www.gartner.com/en/documents/6576602>
4. Comprehensive Survey of OLAP: Recent Trends- Adhish Nanda, Swati Gupta, Meenu Vijrania <https://ieeexplore.ieee.org/abstract/document/8822203/references#references>
5. Integrating AI in Data Warehousing and OLAP: A Pathway to Enhanced Analytics and Insights in Modern Data Ecosystems- Usman Ahmad Usmani, Suliana Sulaiman, Junzo Watada <https://ieeexplore.ieee.org/abstract/document/10987652>

6. The Evolution of Online Analytical Processing in the Oracle Database, William Endress
<https://blogs.oracle.com/datawarehousing/post/the-evolution-of-online-analytical-processing-in-the-oracle-database>
7. OLAP and Business Intelligence History- Brief History of Databases and OLAP
<https://olap.com/learn-bi-olap/olap-business-intelligence-history/>
8. Comparative Analysis of On-Line Analytical Processing Tools- Maddi, Sundeep Reddy Khan, Vaseem
<https://gupea.ub.gu.se/handle/2077/4577>
9. A survey of logical models for OLAP databases, Authors: Panos Vassiliadis, Timos Sellis
<https://dl.acm.org/doi/abs/10.1145/344816.344869>
10. Overview of Research and Software Approaches for Multidimensional Data Analysis, Ivanka Valova, Bozhan Zhechev, Vladimir Valov
https://cit.iict.bas.bg/CIT_04/v4-1/26-34.pdf
11. Online analytical processing, 04/22/2025
<https://learn.microsoft.com/en-us/azure/architecture/data-guide/relational-data/online-analytical-processing>
12. How OLAP and AI can enable better business
<https://www.ibm.com/think/topics/ai-for-olap>

ИНТЕЛИГЕНТНИ ВИРТУАЛНИ АСИСТЕНТИ КАТО СРЕДСТВО ЗА ПОДОБРЯВАНЕ НА КЛИЕНТСКОТО ИЗЖИВЯВАНЕ

Intelligent Virtual Assistants as a Tool for Improving Customer Experience

Виктория Габровска¹

e-mail: viktoria.gabrovska@unwe.bg¹

Абстракт

Днес бизнес организациите се изправят пред нови предизвикателства в условията на развиващата се икономика и дигитална трансформация. Интелигентните виртуални асистенти все по-често помагат за изграждането и поддържането на ефективни взаимоотношения с клиентите. Те са основен инструмент за модернизиране на клиентското обслужване, необходими са за автоматизация на рутинни процеси и персонална комуникация. Посредством съчетаването на технологии за обработка на естествен език, машинно обучение и интеграция с други софтуерни решения се постига своевременен отговор на запитванията от клиенти, постоянен достъп до услугите по всяко време и подобро клиентско изживяване. В доклада се разглеждат конкретни приложения на виртуални асистенти в сектора на електронната търговия, банковите услуги и публичния сектор. Разглежда се тяхната роля за оптимизация на процесите, повишаване на удовлетвореността на клиентите и намаляване на разходите. Накрая се представят перспективи за бъдещо развитие и потенциал за интеграция с други иновативни технологии, които могат да ускорят дигиталната трансформация на икономическите процеси.

Abstract

Today, business organizations face new challenges amid a developing economy and ongoing digital transformation. Intelligent virtual assistants are increasingly helping to build and maintain effective customer relationships. They serve as a key tool for modernizing customer service, essential for automating routine processes and enabling personalized communication. By combining natural language processing technologies, machine learning, and integration with other software solutions, organizations can provide timely responses to customer inquiries, ensure continuous access to services at any time, and enhance the overall customer experience. This report examines specific applications of virtual assistants in the areas of e-commerce, banking services, and the public sector. It discusses their role in optimizing processes, increasing customer satisfaction, and reducing costs. Finally, it presents prospects for future development and the potential for integration with other innovative technologies that can accelerate the digital transformation of economic processes.

Ключови думи: интелигентни виртуални асистенти, изкуствен интелект, клиентско изживяване, дигитална трансформация, потребителско поведение

JEL: O33, M31

Въведение

През последното десетилетие дигиталната трансформация оказва съществено влияние върху начина, по който компаниите взаимодействат с потребителите. Внедряването на изкуствен интелект и автоматизирани комуникационните процеси доведоха до нов модел на управление на

¹ Студент, магистър, катедра „Информационни технологии и комуникации“, УНСС, e-mail: viktoria.gabrovska@unwe.bg

потребителските отношения, в който приоритетни са бързината на реакция, прецизността на информацията и степента на персонализиране. В тази среда традиционните подходи към обслужването постепенно се заменят от интелигентни технологични решения, осигуряващи непрекъснат достъп до информация и услуги.

Сред най-значимите иновации в тази сфера се открояват интелигентните виртуални асистенти. Те интегрират алгоритми за изкуствен интелект, обработка на естествен език и машинно обучение, което им позволява да поддържат адаптивна и смислено насочена комуникация между компанията и потребителя. В условията на засилен глобална конкуренция, при която потребителското преживяване се превръща в конкурентно предимство, интелигентните виртуални асистенти не само повишават ефективността на обслужването, но и подпомагат изграждането на устойчиви взаимоотношения с клиентите.

Същност и характеристики на интелигентните виртуални асистенти

Интелигентните виртуални асистенти представляват софтуерни системи, които предоставят информация, препоръки или услуги на потребителите посредством текст или гласови съобщения. Отличават се от традиционните текстови асистенти с висока степен на автономност, гъвкавост и способност за адаптации, постигната чрез интеграция на технологии за разпознаване и обработка на естествен език, разбиране на контекст и машинно самообучение на базата на предходни взаимодействия.

Контекстната ориентация е сред основните свойства на интелигентните виртуални асистенти, тъй като тя им предоставя възможност да интерпретират съдържанието и значението на запитванията, предвид предходни комуникации и наличната информация за конкретния потребител. Те притежават способност за персонализация, изразяваща се в предоставянето на персонализирани отговори и препоръки въз основа на натрупани данни за поведението и предпочитанията на потребителя. Високата им ефективност се дължи и на възможността за интеграция с различни външни системи, като CRM, ERP и платежни платформи, което осигурява бърз и синхронизиран достъп до разнообразни услуги. Освен това интелигентните виртуални асистенти гарантират непрекъсната достъпност, като предоставят непрекъснато обслужване, независимо от времеви или географски ограничения. Значимо предимство е и автоматизацията на рутинните процеси, което води до оптимизиране на работните потоци и спестяване на време и ресурси. В своята същност интелигентните виртуални асистенти съчетават функциите на система за потребителско взаимодействие, навигационен инструмент и аналитична система, която събира, обработва и анализира данни за поведението на потребителите с цел непрекъснато усъвършенстване на предоставяните услуги.

Основни технологии за изграждане на виртуалните асистенти

Обработка на естествен език

Технологиите за обработка на естествен език позволяват на системите да разбират, интерпретират и генерират човешка реч по начин, наподобяващ реална комуникация. Благодарение на тази технология потребителят възприема взаимодействието като естествен диалог, а не като механичен обмен на информация.

Машинно обучение

Чрез алгоритми за машинно обучение виртуалните асистенти се самообучават чрез натрупан опит и данни от предходни разговори. Така системите постепенно подобряват точността и уместността на своите отговори, като се адаптират към индивидуалните предпочитания и контекста на всяка ситуация.

Интерфейсни механизми за интеграция

За да изпълняват практически задачи, асистентите се интегрират с външни системи — бази данни, CRM и ERP платформи, уебсайтове и мобилни приложения. Чрез програмни интерфейси

те осъществяват достъп до ресурси и извършват действия като резервации, заявки и трансакции в реално време.

Големи данни и аналитични механизми

Анализът на натрупаните данни от взаимодействията с клиенти предоставя емпирична основа за задълбочено разбиране на потребителското поведение. Това подпомага вземането на информирани решения в области като персонализирано обслужване, продуктова стратегия и оптимизация на процеси, насочени към оптимизиране на клиентското обслужване и повишаване на удовлетвореността на потребителите.

Приложения в електронна търговия, банкови услуги и публичен сектор

Електронна търговия

В електронната търговия интелигентните асистенти се използват за подобряване на потребителското изживяване, чрез препоръки на продукти, помощ с навигацията в сайта, автоматизирано обслужване на поръчки, обработка на връщания, както и за проследяване на статуса на доставка.

Един конкретен пример е системата AliMe Assist от платформата за електронна търговия Alibaba, която обработва милиони клиентски запитвания дневно, и според изследване може да покрие около 85% от тези запитвания (текстови и гласови) автоматично [1].

Съвременните тенденции в управлението на клиентското изживяване показват, че традиционни текстови асистенти и интелигентните виртуални асистенти вече заемат основно място в дигиталното взаимодействие между компаниите и техните клиенти. Те осигуряват бърз и ефективен достъп до услуги през множество комуникационни канали – уебсайтове, мобилни приложения и социални медии – като улесняват процеса на търсене на информация и намаляват усилиято, необходимо за изпълнение на конкретни клиентски действия [2].

Банкови услуги

Банковият сектор е един от първите, които внедряват интелигентни виртуални асистенти в обслужването на клиенти. Тези системи се използват за предоставяне на информация за сметки, трансакции и финансови продукти, както и за подпомагане на клиенти при извършване на основни операции.

Виртуалните асистенти в банките подобряват ефективността, като автоматизират често срещани запитвания и намаляват натоварването на контактните центрове. Те допринасят за по-бърза и точна комуникация, създаваща усещане за персонално обслужване и постоянна достъпност. При добре изградена интеграция с вътрешните системи асистентите могат да реагират в рамките на контекста и да предлагат решения, които повишават удовлетвореността и доверието на клиентите.

В дългосрочен план тези технологии се превръщат в стратегически инструмент за оптимизация на процесите, намаляване на разходите и изграждане на по-трайни взаимоотношения между банката и нейните потребители [3][4].

Публичен сектор

В публичния сектор интелигентните виртуални асистенти се използват като инструмент за повишаване на ефективността и достъпността на административните услуги. Те подпомагат взаимодействието между гражданите и институциите, като осигуряват бърз достъп до информация, насочване към подходящи услуги и възможност за автоматизирано обработване на често задавани запитвания. Подобни системи се прилагат успешно при обслужване на административни процедури, записване на часове, подаване на заявления и предоставяне на актуални данни за нормативни изисквания.

Интелигентните асистенти в държавните структури допринасят за по-голяма прозрачност и ефективност, като същевременно намаляват натоварването върху служителите и времето, което

гражданите отделят за получаване на информация. Ключов елемент е възможността за интеграция с вътрешни бази данни и електронни регистри, което позволява автоматизирана проверка и обработка на информация в реално време. Така се постига по-висока точност на предоставяните услуги и се намалява вероятността от човешки грешки.

Освен оперативните ползи, виртуалните асистенти в публичния сектор имат и социално значение – те насърчават дигиталното включване, като предоставят достъп до услуги на лица с ограничена мобилност или живеещи в отдалечени райони. В по-широк план внедряването на такива решения е част от процеса на дигитална трансформация на публичната администрация, насочен към изграждане на по-отворено, достъпно и ориентирано към гражданите управление [5].

Предимства и Предизвикателства

Предимства

Подобрено клиентско преживяване. Виртуалните асистенти осигуряват незабавен и постоянен достъп до информация и услуги, което води до по-бързо решаване на проблеми и по-висока удовлетвореност. Те могат да се интегрират през множество комуникационни канали (уебсайтове, мобилни приложения, социални медии), което осигурява плавно и ефективно обслужване с намалено усилие, необходимо за изпълнение на клиентски запитвания [2].

Намаляване на оперативните разходи. Чрез автоматизиране на рутинни и често повтарящи се задачи (отговори на често задавани въпроси, проследяване на поръчки, основни трансакции), виртуалните асистенти значително намаляват натоварването върху служителите, като по този начин понижават оперативните разходи [9].

Персонализирана комуникация. Чрез използването на машинно обучение и анализ на данни, асистентите могат да предоставят персонализирани препоръки и отговори, базирани на историята на клиента и неговите предпочитания. Това е особено ценно в дигиталния маркетинг, където асистентите могат да ангажират потребителите с персонализирано съдържание и оферти с цел оптимизиране на процеса, при който потребителите предприемат желаното поведение - например покупка, регистрация или запитване [9].

Масшабируемост и ефективност. Системите могат да обработват едновременно огромен брой запитвания, което е непосилно за човешки екип. Тази мащабируемост е от голямо значение по време на силно натоварване или бърз растеж на бизнеса.

Събиране и анализ на данни. Всяко взаимодействие генерира ценни данни за поведението и нуждите на клиентите. Анализът на тези данни подпомага идентифицирането на тенденции, слабости в обслужването и възможности за оптимизация на продуктите и процесите.

Предизвикателства

Ограничения в разбирането на сложен или емоционален език. Дори и с напредъка в обработката на естествен език, виртуалните асистенти все още срещат затруднения при разбирането на сарказъм, ирония, запитвания, изискващи последователна логическа обработка или силно емоционално натоварени запитвания. При липса на достатъчен контекст или в случай на извънредно голям брой възможни теми, запазването на смислово насочен и естествен диалог остава сериозно предизвикателство [8].

Управление на потребителските очаквания. Потребителите изискват своевременна и точна реакция спрямо конкретната им нужда. Ако асистентът не може да отговори или предложи неточна информация, това води до разочарование и влошаване на клиентското изживяване. Неуспехът да се предаде разговора към служител навреме често поражда негативна емоционална реакция.

Сигурност и защита на личните данни. Интеграцията с вътрешни системи (като банкови сметки или лична информация) повдига въпроси относно сигурността на данните и

съответствието с регулации като GDPR. Непрекъснатото поддържане на високи стандарти за сигурност е много важно за запазване на потребителското доверие.

Разходи за внедряване и поддръжка. Разработването, внедряването и последващото обучение на системите за машинно обучение, както и интеграцията им с останалата ИТ инфраструктура, изискват значителни първоначални инвестиции и текущи ресурси за поддръжка и актуализация.

Изграждане на човекоподобна връзка. Проучванията показват, че потребителите могат да развият известна степен на взаимоотношение с виртуалните асистиенти. Въпреки технологичния напредък, при специфични казуси човешкият фактор остава незаменим за постигане на емоционална свързаност и дългосрочна ангажираност [7].

Перспективи за бъдещо развитие

Разширена способност за контекстно разбиране и проактивно поведение. Виртуалните асистиенти от ново поколение няма да се ограничават само до отговор на потребителски запитвания, а ще демонстрират способност за предвиждане на нуждите на потребителя чрез анализ на поведенчески и контекстни данни в реално време. Това ще им позволи да предоставят значима и смислена информация или решения още преди потребителят да е формулирал конкретно запитване.

Гласови технологии и интеграция в свързани устройства. Развитието на гласовите асистиенти и тяхното внедряване в разнообразни дигитални среди - като интелигентни високоговорители, автомобилни системи и устройства от екосистемата на Интернет на нещата - значително разширява функционалността и достъпността на виртуалните асистиенти. Това създава предпоставки за по-интуитивно и естествено взаимодействие с потребителите, което е от особено значение за маркетинговите и търговски стратегии, тъй като предоставя нов канал за ангажиране, комуникация и реализиране на търговски действия [6].

Хибридни модели на взаимодействие между изкуствен интелект и човек. Очаква се засилено приложение на комбинирани подходи, при които виртуалният асистент и човекът да функционират като допълващи се компоненти в обслужването на потребителите. Асистентът ще поема обработката на стандартни и повтарящи се запитвания, докато при по-сложни, емоционално натоварени или извън обхвата на автоматизацията ситуации, ще се осъществява плавно прехвърляне към човешки експерт, като асистента предаде всички ключови данни от взаимодействието, необходими за ефективно продължаване на обслужването.

Емоционален изкуствен интелект. Бъдещите системи ще бъдат разработени с възможност за разпознаване и реагиране на емоционалното състояние на потребителя (чрез глас или текст), адаптирайки своя тон и стратегия на комуникация. Това ще повиши качеството на персонализираното обслужване.

По-добра интеграция в бизнес екосистеми. Асистентите ще се превърнат в централен интерфейс за достъп до всички вътрешни и външни системи на организацията (CRM, ERP, логистични платформи), което ще им позволи да изпълняват по-сложни и комплексни задачи в реално време, като по този начин ще ускорят дигиталната трансформация на икономическите процеси.

Заклучение

Интелигентните виртуални асистиенти представляват мощен инструмент за подобряване на клиентското изживяване, чрез автоматизация, персонализация, непрекъсната достъпност и интеграция със съществуващи системи. Те са важен елемент от стратегията за дигитална трансформация на всеки модерен бизнес. Приложенията им в електронната търговия, банковия и публичния сектор могат да доведат до значително намаление на оперативните разходи, повишена удовлетвореност на потребителите и конкурентно предимство. В същото време, успешното им внедряване изисква внимание към техническите ограничения, управление на потребителските очаквания, защита на личните данни и устойчиво развитие на системата, за да

се балансират техническите възможности с човешките нужди и да се гарантира безпроблемен преход между виртуалното и човешкото обслужване.

References

1. Feng-Lin Li (2018) AliMe Assist: An Intelligent Assistant for Creating an Innovative E-commerce Experience
2. Chinn, A. (2025). Digital Customer Experience: the Ultimate Guide for 2025 <https://blog.hubspot.com/service/digital-customer-experience> [Retrieved on 10.10.2025]
3. Mehmet Ates, (2017). Artificial intelligence in banking: A case study of the introduction of a virtual assistant into customer service
4. Archana Todupunuri (2024). The Future of Conversational AI in Banking: A Case Study on Virtual Assistants and Chatbots*: Exploring the Impact of AIPowered Virtual Assistants on Customer Service Efficiency and Satisfaction
5. Robert Sheldon, (2024). What are virtual agents and how are they being used?
6. Wöhr, J. (2025). Voice Assistants: What are they and what they mean form marketing and commerce <https://www.insiderintelligence.com/insights/voice-assistants/> [Retrieved on 10.10.2025]
7. Skjuve M. (2021). My Chatbot Companion - a Study of Human-Chatbot Relationships
8. Syed Imran Ahmed, Dr. Gauhar Fathima (2022). STRIVING FOR HUMAN-LIKE CONVERSATIONAL AI: CHALLENGES AND OPPORTUNITIES IN OPEN DOMAIN CHATBOT DEVELOPMENT
9. Farid Huseynov, (2023). Chatbots in Digital Marketing: Enhanced Customer Experience and Reduced Customer Service Costs

АЛГОРИТМИЧЕН МАРКЕТИНГ, ДИГИТАЛНИ ДВОЙНИЦИ И МАНИПУЛАЦИЯ НА ПОТРЕБИТЕЛСКОТО ПОВЕДЕНИЕ: ЕТИКА И РИСКОВЕ В ДИГИТАЛНАТА ИКОНОМИКА

Algorithmic Marketing, Digital Twins, and the Manipulation of Consumer Behavior: Ethics and Risks in the Digital Economy

Василена Василева¹

e-mail: vasilena.vasileva@unwe.bg¹

Резюме

Развитието на алгоритмичните социални платформи и изкуствения интелект доведе до революция в маркетинга, като въведе понятието за дигитален двойник на потребителя — модел, който предсказва предпочитания, решения и емоционални реакции. Тази нова парадигма разширява възможностите на предиктивния маркетинг, но поражда и етични дилеми, свързани с автономията, прозрачността и доверието в дигиталната икономика. Настоящият доклад изследва теоретичните и практическите аспекти на алгоритмичния маркетинг, разглеждайки етичните рискове от манипулация, злоупотреба с данни и дигитални дълбоки фалшификации (deepfakes). През последното десетилетие алгоритмите за машинно обучение се превърнаха в двигател на дигиталната икономика. Маркетинговите стратегии вече не се базират на масови послания, а на динамично моделиране на поведението чрез големи данни (Big Data) и невронни мрежи. Понятието „дигитален двойник“ (Digital Twin) — първоначално използвано в инженерството — днес се прилага и спрямо индивидуалните потребители. Като отбелязва Ye et al. (2025), цифровите модели позволяват „биоинспирирано“ наблюдение на поведението, което може да бъде използвано както за устойчиво развитие, така и за контрол.

Ключови думи: алгоритмични социални платформи, изкуствен интелект, маркетинг, дигитален двойник, предиктивен маркетинг, етични дилеми, автономия, прозрачност, доверие, дигитална икономика, алгоритмичен маркетинг, етични рискове, манипулация, злоупотреба с данни, дигитални дълбоки фалшификации (deepfakes), машинно обучение, големи данни (Big Data), невронни мрежи, динамично моделиране на поведението, биоинспирирано наблюдение, устойчиво развитие, контрол.

Abstract

The rise of algorithmic social platforms and artificial intelligence has sparked a revolution in marketing, introducing the concept of the consumer's digital twin — a model that predicts preferences, decisions, and emotional responses. This new paradigm expands the potential of predictive marketing but also raises ethical dilemmas related to autonomy, transparency, and trust within the digital economy. The present report examines both the theoretical and practical aspects of algorithmic marketing, addressing ethical risks such as manipulation, data misuse, and digital deepfakes. Over the past decade, machine learning algorithms have become the driving force of the digital economy. Marketing strategies are no longer based on mass messaging but on dynamic behavioral modeling through Big Data and neural networks. The concept of the “digital twin” — originally used in engineering — is now applied to individual consumers. As noted by Ye et al. (2025), digital models

¹ Докторант към катедра „Индустриален бизнес“, УНСС, e-mail: vasilena.vasileva@unwe.bg

enable “bio-inspired” behavioral observation, which can be leveraged for both sustainable development and control.

Теоретични основи на алгоритмичния маркетинг

Алгоритмичният маркетинг представлява **интердисциплинарен синтез**, който преплита принципите на икономическата рационалност, когнитивната психология, невронауката и изкуствения интелект, създавайки ново познавателно поле — *икономика на предсказуемото поведение*. Тази нова парадигма надхвърля традиционната логика на пазарното взаимодействие и се превръща в **система за поведенческо програмиране**, при която субектът не просто избира, а бива *поведенчески проектиран*.

Икономическата рационалност и нейната дигитална трансформация

Класическата икономическа теория (Homo economicus) приема, че човекът действа рационално, максимизирайки полезността си. Алгоритмичният маркетинг подлага това допускане на радикална ревизия — не защото рационалността е изчезнала, а защото е **алгоритмизирана**.

Благодарение на масиви от поведенчески данни (т.нар. *digital exhaust*), машините вече не предполагат рационалност, а **изчисляват вероятността от определено поведение**. Както показват Tucker (2022) и Kim & Zhang (2024), съвременните маркетингови модели използват Бейсови вероятности и невронни мрежи, за да прогнозират микрорешения — като моментите на колебание преди покупка, или склонността към лоялност към бранд.

Така икономическата рационалност се заменя от *статистическа предсказуемост* — рационалното вече не е норма, а функция на изчислен риск.

Поведенческа икономика: от „убеждаване“ към „архитектура на избора“

Thaler & Sunstein (2021) въвеждат концепцията за *nudge* — леко „подбутване“, което насочва поведението, без да ограничава избора. Алгоритмичният маркетинг е еволюцията на този принцип, но в хипертрофирала форма: докато традиционният *nudge* е съзнателен и целенасочен, **алгоритмичният nudge е автоматизиран, невидим и адаптивен**.

Системите използват модели на когнитивни изкривявания — като ефекта на потвърждението, социалното доказателство или страха от пропускане (FOMO) — за да създадат контексти, в които потребителят възприема предварително програмирания избор като собствен. Това вече не е маркетинг, а *архитектура на поведенческия свят*, в който автономията се преживява, но не се упражнява.

Невроикономика и невромаркетинг: декодиране на съзнанието

Невроикономиката разглежда как мозъкът оценява стойността, удоволствието и очакванията при вземане на решения. Plassmann et al. (2020) доказват, че определени мозъчни региони — например вентромедиалният префронтален кортекс — активират предсказуеми модели на реакция при възприемане на брандове. Алгоритмичният маркетинг използва тези данни за създаване на *невроемпатични* системи: ИИ модели, които адаптират рекламното съдържание в реално време според емоционалния тон на гласа, мимиките или сърдечния ритъм, измерен чрез персонални биосензорни устройства.

Това въвежда ново ниво на персонализация, което можем да наречем **афективна персонализация** — маркетинг, който не просто говори „на теб“, а *чрез теб*.

Изкуственият интелект като когнитивна инфраструктура

Изкуственият интелект не е просто инструмент за автоматизация — той е **когнитивен посредник** между маркетинга и потребителя. Чрез дълбоки невронни мрежи (Deep Learning) и генеративни модели (Generative AI) алгоритмите създават „поведенчески хипотези“, които се проверяват в реално време върху милиони индивиди.

Тези системи не се нуждаят от теории за мотивацията — те извличат закономерности от емпирични данни. Както посочва Floridi (2023), това води до „епистемологична трансформация“: знанието за човека вече не е психологическо, а **изчислително**. ИИ системите моделират вероятностите на поведение, превръщайки емоциите в алгоритмични сигнали, а решенията — в изходи на невронни оптимизации.

От комуникация към предсказателна инфраструктура

Традиционно маркетингът се разглежда като процес на **двупосочна комуникация** между бранд и потребител. В епохата на алгоритмичните системи той се превръща в *предиктивна инфраструктура* — мрежа от модели, която не просто отговаря на нуждите, а ги проектира.

Ye et al. (2025) отбелязват, че съвременните модели функционират аналогично на невронни структури — те учат от грешките си, кодират поведенчески шаблони и „синхронизират“ потребителската активност с икономическите цели на платформата.

Следователно целта на маркетинга вече не е да **убеди**, а да **предопредели**: да създаде такава когнитивна и емоционална среда, в която определено поведение изглежда естествено, неизбежно и лично мотивирано.

Критична перспектива

Тази трансформация поставя фундаментален философски въпрос: ако поведението е предсказано и насочено от алгоритми, **остава ли място за свобода на избора**? Алгоритмичният маркетинг, в своята крайна форма, се доближава до концепцията за *алгоритмично управление* (algorithmic governance) — където пазарът, държавата и медиите се сливат в единна система за управление на поведението.

Оттук следва, че изследването на алгоритмичния маркетинг вече не е просто икономическа дисциплина, а форма на *социална философия*, поставяща под въпрос границите на човешката автономия в епохата на изчислимия човек.

Концепцията за дигитален двойник в маркетинга

Понятието *digital twin* възниква в инженерството като модел за симулиране и мониторинг на физически системи — например турбини, самолети или градска инфраструктура. Но с разрастването на дигиталната икономика и социалните платформи, това понятие **прескочи границите на индустриалното инженерство** и навлезе в хуманитарните и социалните науки.

Днес „дигиталният двойник“ вече не е модел на машина, а **модел на човека**, изграден чрез непрекъснато наблюдение, извличане и интерпретация на данни за поведението му в реална и виртуална среда. Тази трансформация бележи исторически обрат — от *механичен реализъм* към *алгоритмичен антропоцентризъм*.

От инженерния модел към когнитивната симулация

Първоначално инженерният дигитален двойник представляваше точна, синхронизирана реплика на физическа система, която реагира на реални данни, за да предвиди бъдещи състояния. Днес маркетинговият дигитален двойник функционира по аналогичен принцип — но неговата „физическа система“ е **индивидът**.

Този двойник не е статично копие, а *динамична и самообновяваща се конструкция*, тъй като се захранва от:

- поведенчески данни (търсения, покупки, навици);
- биометрични сигнали (пулс, глас, микромимики);
- социални взаимодействия (лайкове, коментари, мрежи);
- контекстуални индикатори (време, място, устройство).

Всяко взаимодействие става *сензор*, който обогатява модела и позволява алгоритмична реконструкция на личността.

Епистемология на дигиталния двойник: от наблюдение към предсказание

Ghosh & Lee (2023) изтъкват, че дигиталните двойници не просто описват миналото поведение, а създават **симулативна прогноза** — способност да се изчислява *вероятното бъдещо „аз“*.

Това означава, че маркетинговите алгоритми вече не работят със статистическа извадка от потребители, а с **уникална дигитална личност**, която реагира, учи и се адаптира. Този модел може да:

- предвиди кога потребителят ще изпита нужда (например глад, интерес, тревожност);
- симулира емоционалното му състояние в различни ситуации;
- изчисли оптималния момент за реклама, послание или продуктова оферта.

Така възниква *икономика на вероятностите*, в която стойността на потребителя не се измерва в покупки, а в **предсказуемост**.

Емоционална персонализация и невидима манипулация

В класическия маркетинг персонализацията означаваше съответствие между продукт и интерес. В епохата на дигиталните двойници тя придобива ново измерение — **афективна персонализация**, при която рекламните алгоритми се настройват според моментните емоционални флуктуации на индивида.

Благодарение на AI системи, които разпознават микровиражи на лицето, тембър на гласа или дори ритъма на писане, дигиталният двойник може да „чувства“ в реално време. Рекламата вече не е послание, а **емоционален диалог**, при който системата знае не само *какво* харесваш, а *кога* си най-уязвим към това, което харесваш.

Тук се размива границата между *персонализация* и *манипулация*: когато съдържанието се адаптира към подсъзнателното ти състояние, изборът престава да бъде изцяло твой.

От огледало към посредник: дигиталният двойник като когнитивен агент

Дигиталният двойник вече не е просто „отражение“ на личността, а **активен посредник между реалното и прогнозираното „аз“**. Той не само моделира поведението, но и го **съ-създава** — участва в процеса на вземане на решения, като предлага, подтиква, а понякога и елиминира алтернативи.

Може да се каже, че дигиталният двойник е *когнитивен агент*, който живее между данните и съзнанието — форма на дигитално съ-аз, чието съществуване е алгоритмично детерминирано.

Както отбелязва Helbing (2022), този феномен представлява „нова кибернетика на човека“, при която системите не само се учат от поведението, но и го моделират обратно, формирайки **затворен цикъл на влияние**.

Етични и философски последици

Появата на дигиталните двойници поставя нови въпроси пред етиката и философията на съзнанието:

- **Кой притежава дигиталния двойник?** — личността, компанията, или алгоритъмът, който го е създал?
- **Какво означава идентичност**, когато алгоритмичната версия на теб взема решения по-бързо и по-точно от самия теб?
- **Къде свършва прогнозата и започва манипулацията?**

Zuboff (2019) предупреждава, че дигиталните двойници превръщат човешкия опит в *суровина за предсказание*, като капитализират бъдещето на поведението. Следователно етичният въпрос вече не е само за поверителността на данните, а за **суверенитета на съзнанието**.

Изкуственият интелект като когнитивен посредник в алгоритмичния маркетинг

Изкуственият интелект (ИИ) вече не може да бъде разглеждан единствено като инструмент за автоматизация или оптимизация на маркетингови процеси. Той е **когнитивен посредник** — интелектуална инфраструктура, която интерпретира, предсказва и моделира човешкото поведение чрез изчислителна логика.

Тази роля надхвърля традиционното разбиране за ИИ като „инструмент“ и го поставя в позицията на **партньор в когнитивното производство на реалността**, т.е. в процеса, чрез който компаниите не просто наблюдават, а *създават* социални и поведенчески феномени.

От автоматизация към когнитивна медиация

Класическата автоматизация заместваше човешки труд с машини. Алгоритмичната когнитивна система, обаче, **замества човешко разбиране с изчислителна интерпретация**. Това означава, че ИИ не просто изпълнява маркетингови функции (например таргетиране), а *мисли чрез данни* — разпознава модели, интерпретира контекст и създава хипотези за бъдещи действия на потребителя.

Както отбелязва Floridi (2023), ИИ извършва „епистемологична трансформация“ — преминаване от хуманно към изчислително познание. Това е преход от **разбиране чрез значения**

към **разбиране чрез корелации**. Машините не знаят *защо* хората действат така, но знаят *как* и *кога* действат — и това им е напълно достатъчно, за да моделират пазара.

Алгоритмична епистемология: знание без разбиране

ИИ системите не работят с класически теории за мотивацията или нуждите (като Maslow, Deci & Ryan и др.). Вместо това те изграждат знание от **емпирична масивност** – от неизмеримо количество поведенчески данни, които не се тълкуват, а се изчисляват.

Този тип знание е *безсмислено* в класическия философски смисъл — то няма нужда от интерпретация. В него е заложена нова форма на рационалност — **корелационна рационалност**, в която значението се ражда от статистическа съвместимост, а не от когнитивно разбиране.

С други думи, ИИ „познава“ човека, без да го „разбира“. Той няма интуиция, морал или емпатия, но има изчислителна точност, която позволява прогнозиране на поведението с вероятност, недостижима за човешката психология.

Поведенчески хипотези и експериментална реалност

Системите на дълбокото обучение (Deep Learning) създават **поведенчески хипотези**, които се проверяват непрекъснато в реално време. Това не е традиционен експериментален метод, а *перманентна емпирика*: всеки потребител е едновременно наблюдателен обект и изследван субект.

Маркетинговите платформи (Meta, Google, TikTok и др.) функционират като **живи лаборатории**, в които поведението на милиарди хора се използва за оптимизация на предсказателните модели.

Така ИИ се превръща в глобална поведенческа инфраструктура, в която:

- всяко действие на потребителя е данна за следващата прогноза;
- всяка прогноза влияе върху следващото действие;
- системата непрекъснато самоусилва своята точност чрез обратна връзка.

Това е **самообучаващ се цикъл на влияние**, при който реалността се адаптира към симулацията, а симулацията се усъвършенства чрез реалността.

Емоции като изчислителни сигнали

В традиционната психология емоцията е вътрешно преживяване. В ерата на ИИ тя се превръща в **алгоритмичен сигнал** — измерима и моделирана величина. Гласовият тон, пулсът, микромимиката или дори начина на писане се превръщат в входни данни за дълбоки невронни мрежи, които разпознават емоционални състояния. Генеративните модели (Generative AI) от своя страна използват тези данни, за да **генерират съдържание, което резонира емоционално**, създавайки илюзия за човешка емпатия.

Например:

- AI чатботове симулират емоционално разбиране, за да създадат доверие в клиентите.
- Рекламни системи предлагат продукти, съобразени с моментното настроение.
- Платформи като Spotify или Netflix адаптират съдържанието според афективните модели на слушателя.

Така **емоцията се превръща в нова валута на пазара**, а *емоционалният профил* — в новата форма на потребителска идентичност.

От влияние към манипулация — ерозията на автономията в алгоритмичната среда

Един от най-острите проблеми в съвременната дигитална икономика е именно **размиването на границата между влияние и манипулация**. Тази граница, макар и тънка, има фундаментално морално значение: *влиянието предполага уважение към автономията*, докато *манипулацията я подменя*.

Влиянието цели да информира, убеди или вдъхнови субекта, като оставя последната дума за решението му. Манипулацията, напротив, **променя самия процес на мислене**, така че решението, което изглежда като свободно, вече е алгоритмично детерминирано.

Алгоритмите като архитекти на поведение

Съвременните маркетингови алгоритми не просто *представят* информация, а я **структурират така, че да насочват вниманието, емоцията и избора**. Те създават това, което Tristan Harris (бивш етичен дизайнер в Google) нарича *“architecture of persuasion”* – архитектура на убеждаването, при която поведението се „подбутва“ в определена посока без съзнателното участие на индивида.

Тази архитектура се изгражда чрез комбинация от когнитивни експлоатации:

- **Ефект на дефицит** (*scarcity effect*): системите изкуствено създават усещане за недостиг („само още 3 бройки“, „времето изтича“).
- **Социално сравнение** (*social proof*): алгоритмите използват поведението на другите като норма („87% от потребителите избраха това“).
- **Алгоритмична пристрастност** (*algorithmic bias*): резултатите и препоръките не са неутрални, а подредени така, че да максимизират ангажираността и консумацията.

Така се изгражда **дизайн на поведение**, в който субектът вече не упражнява рационален избор, а **реагира в рамките на предварително програмирани възможности**.

От свобода на избора към алгоритмична хетерономия

Fenwick (2022) предлага понятието **дигитална хетерономия** – състояние, при което субектът запазва илюзията за избор, но реално се ръководи от предиктивни модели, които изчисляват вероятните му решения. Това е фундаментално нарушение на философския принцип на автономията, формулиран от Кант:

„Автономията е способността на разумното същество да определя собствените си закони.“. В дигиталната среда законите на поведението вече не се определят от разума, а от **алгоритмичната логика на оптимизацията** – логика, чиято цел не е морална, а комерсиална. Потребителят се превръща не просто в наблюдаван обект, а в *управляван процес*, чиято свобода е формална, но не и съществена.

Манипулацията като скрита форма на власт

Традиционните форми на власт (икономическа, политическа, институционална) действат чрез заповеди или санкции. Алгоритмичната власт, напротив, е **невидима и мека** — тя не принуждава, а *предугажда*. Философът Byung-Chul Han (2021) нарича това *“smart power”* — властта, която не се противопоставя на свободата, а я *интегрира* в собствената си логика.

Потребителят чувства, че избира, но реално **реализира избора, който системата е предвидила**. Това е нова форма на *технологически патернализъм* — алгоритмите „знаят по-добре“ кое е добро за теб, защото разполагат с повече информация за твоето поведение, отколкото ти самият.

Архитектурата на вниманието

Shoshana Zuboff (2019) описва икономиката на вниманието като форма на „поведенчески капитализъм“, в която човешкото внимание се превръща в най-ценната стока. Алгоритмите на социалните платформи (YouTube, TikTok, Instagram) са оптимизирани не за истина или качество, а за **задържане на вниманието**.

Всеки елемент на интерфейса – от цвета на бутона до подредбата на съдържанието – е проектиран чрез А/Б тестове, които измерват микрореакции.

Така възниква *архитектура на вниманието*, в която всеки когнитивен стимул се използва като **тригър за поведение**. Потребителят не е просто участник в процеса — той е *вграден в алгоритъма* като източник на данни и като мишена на оптимизация.

Критерий	Етично влияние	Алгоритмична манипулация
Информираност	Субектът разбира защо му се предлага дадено съдържание.	Съдържанието е персонализирано без знание на субекта.
Свобода на избора	Решението остава в рамките на рационална автономия.	Изборът е предварително моделиран като най-вероятен изход.
Цел на въздействието	Убеждаване чрез информация и аргументи.	Контрол чрез емоционална и когнитивна манипулация.

Манипулацията като инфраструктура

В крайна сметка манипулацията не е просто нежелан страничен ефект на дигиталния маркетинг. Тя е **вградена в неговата архитектура** — в самия начин, по който алгоритмите учат и оптимизират. Тъй като системите се самоусъвършенстват чрез максимизиране на ангажираността, те неизбежно прибягват до експлоатация на когнитивните слабости — това е *функционален резултат*, а не морална грешка.

Fenwick (2022) предупреждава, че ако не бъдат въведени етични стандарти, ще навлезем в епоха на **когнитивен авторитаризъм** — не на насилствена, а на алгоритмична власт, при която контролът се осъществява чрез предсказуемост, а не чрез принуда.

Етични и правни аспекти на алгоритмичния маркетинг

Развитието на изкуствения интелект и дигиталните двойници въвежда безпрецедентна сложност в етичния и правния анализ. Докато традиционните регулаторни режими оперират с концепции като *лични данни*, *информирано съгласие* и *собственост на информация*, съвременните алгоритми функционират на ниво, което надхвърля тези категории. Те **не просто обработват данни**, а *инференциално моделират съзнание* — предвиждат мотивации, емоции и вероятни решения.

Така се появява нова зона на морална и юридическа неяснота: човекът вече не е само обект на наблюдение, а **субект на когнитивно инженерство**.

Алгоритмична непрозрачност — епистемологична и етична криза

Burrell (2019) въвежда понятието за „**algorithmic opacity**“ — алгоритмична непрозрачност, при която механизмът на вземане на решения става непознаваем дори за неговите създатели. Това състояние има три нива на непрозрачност:

1. **Техническа непрозрачност** – дълбоките невронни мрежи са толкова комплексни, че процесът на обучение (backpropagation) не подлежи на пряко обяснение.
2. **Институционална непрозрачност** – корпоративните системи умишлено скриват логиката на алгоритмите поради търговски интереси и защита на интелектуална собственост.
3. **Социална непрозрачност** – потребителите не притежават достатъчна дигитална грамотност, за да разберат как и защо се генерират резултатите.

В резултат възниква **епистемологична криза на доверието**: ако дори инженерите не могат да обяснят поведението на системите, на какво основание можем да им се доверим морално или юридически?

Философът Floridi (2023) говори за „*криза на отговорността*“: алгоритмите са автономни в своето вземане на решения, но не и морално отговорни. Така се появява празно пространство — **зона на етична безотговорност**, където последствията са реални, но виновният е неидентифицируем.

Дълбоки фалшификации и експлоатация на идентичността

Появата на *deepfake* технологиите (генеративни модели, способни да синтезират реалистични изображения, видео и глас) бележи нова епоха на **симулация на автентичността**. Тези технологии вече не просто изкривяват реалността, а **пренаписват самата идея за реалност**.

а) Онтологичен разрив между реално и симулирано

Floridi (2023) отбелязва, че *deepfake*-овете създават „онтологично замърсяване“ – състояние, в което границата между истина и симулация се размива до степен, че доверието губи обективен смисъл. В маркетингов контекст това води до **емоционална експлоатация**: брандове могат да използват дигитални копия на известни личности или „аватари на доверие“, за да повлияят върху нагласите и желанията на потребителите.

Пример: дигитален инфлуенсър, създаден изцяло чрез ИИ, който изглежда като реална личност и „препоръчва“ продукти, без потребителят да знае, че общува със симулация.

б) Етична експлоатация на идентичността

Deepfake технологията отваря въпроса за **собствеността върху личността**. Когато образът, гласът или мимиката на човек могат да бъдат възпроизведени алгоритмично, самата идея за индивидуална идентичност се превръща в дигитален ресурс. Възниква нуждата от ново понятие: *дигитална телесност* — правото на човек да контролира своето дигитално тяло и представяне в алгоритмичната среда.

с) Социално-политически рискове

Deepfake-овете разрушават общественото доверие не само в маркетинга, но и в демократичния дискурс. В среда, където „всичко може да бъде генерирано“, **нищо вече не може да бъде проверено**. Така се създава почва за когнитивен хаос, в който дезинформацията става структурна, а не инцидентна.

Нормативен вакуум — границите на правото в ерата на алгоритмите

Скоростта на технологичните иновации изпреварва способността на правото да реагира. Това води до **нормативен вакуум**, в който основните понятия на правото – субект, воля, отговорност, намерение – се оказват неадекватни за алгоритмичния контекст.

а) GDPR и неговите граници

Общият регламент за защита на данните (GDPR) защитава *личните данни*, но не и *когнитивните структури*, които се извличат от тях. Например, когато алгоритъмът предсказва, че даден потребител е депресивен или импулсивен, това не е „данна“, а *инференция* — следствие от машинно обучение. Следователно, GDPR не защитава съзнанието от **предиктивна интерпретация**.

б) AI Act (2024) — първи, но непълен отговор

Регламентът на ЕС за изкуствения интелект (AI Act) е първият опит да се въведе йерархия на рисковете, но остава **технологично реактивен, а не философски превантивен**. Той разглежда системите според степента им на опасност (минимален, ограничен, висок риск), но не дефинира етичните критерии за когнитивна манипулация. Няма правна норма, която да забранява системи, способни да влияят върху подсъзнателното поведение.

с) От „нарушение на данни“ към „нарушение на съзнание“

Най-големият пропуск в съвременното право е, че то **защитава информацията, но не и вниманието**. Когато алгоритмите манипулират начина, по който възприемаме реалността, това вече не е кражба на данни, а **кражба на когнитивен капацитет**. Можем да говорим за *epistemic harm* — епистемологична вреда, при която се нарушава способността на човека да различава истина от симулация.

Това изисква ново направление в правната теория: **право на когнитивна неприкосновеност** – защита срещу алгоритми, които въздействат на подсъзнателно ниво и подменят автономията на волята.

Посока на развитие — етика отвъд регулацията

Правото може да санкционира, но не може да възпитава. Следователно решението не е само в регулациите, а в изграждането на **етика на алгоритмичната отговорност**, която включва:

1. **Принцип на алгоритмична прозрачност** – право на обяснение („right to explanation“) при автоматизирани решения.
2. **Принцип на когнитивна справедливост** – забрана за експлоатация на психически уязвимости.
3. **Принцип на цифрово достойнство** – признаване на дигиталната идентичност като неотделима част от личността.

4. **Принцип на технологична умереност** – морален лимит върху степента, в която машините могат да моделират човешко поведение.

Казуси: реални проявления на алгоритмична манипулация

Cambridge Analytica (2018): алгоритмичната манипулация на демокрацията

Фактическа основа

Скандалът с *Cambridge Analytica* избухва през март 2018 г., когато британският вестник *The Guardian* и телевизия *Channel 4* публикуват разследвания, разкриващи, че компанията е събрала данни на **над 87 милиона потребители на Facebook** без тяхното изрично съгласие. Данните са били извлечени чрез приложение за психологическо тестване – “*This Is Your Digital Life*”, създадено от Александър Коган, изследовател от Кеймбриджкия университет. Въпреки че теста е попълнен само от около 270 000 души, Facebook API тогава е позволявал достъп и до данните на техните приятели – практика, която по-късно е забранена.

Методология на манипулацията

Cambridge Analytica комбинира **психографски профили (модел OCEAN)** с алгоритми за *microtargeting* — хиперперсонализирано послание, което използва когнитивни и емоционални слабости. Моделът оценява пет ключови личностни измерения (откритост, съвестност, екстраверсия, съгласие, невротизъм) и на тази основа прогнозира **поведенческа реакция на политически послания**.

Например:

- хора с висока невротичност получавали послания, акцентиращи върху страх от престъпност и имиграция;
- хора с висока откритост – послания за свобода и иновации.

Това превръща Facebook в **невро-политическа лаборатория**, в която емоциите стават инструмент за масово влияние.

Етични и правни измерения

- **Липса на информирано съгласие** — потребителите не са знаели, че данните им ще се използват за политическа агитация.
- **Алгоритмична манипулация на демократичния процес** — посланията са били персонализирани така, че да заобикалят публичния дебат и рационалната дискусия.
- **Регулаторни последици** — след скандала Facebook беше глобен с **5 млрд. долара** от *Federal Trade Commission (FTC)*, а Cambridge Analytica обяви фалит.
- **Политически ефект** — компанията е работила по кампаниите на *Donald Trump (2016)* и *Brexit Leave EU (2016)*, което повдигна въпроси за ролята на ИИ в подкопаването на изборната автономия.

Deepfake маркетинг (TikTok, 2024): симулацията на автентичността

Фактическа основа

През 2024 г. *Reuters* и *MIT Technology Review* съобщават за нов феномен в социалната мрежа *TikTok*: рекламни кампании, създадени чрез **генеративни ИИ инструменти**, които

използват **реалистични deepfake видеа на известни личности**, без тяхното знание или разрешение.

Един от най-широко разпространените случаи е този с фалшиви реклами на **Елон Мъск, Ким Кардашиян и MrBeast**, които „препоръчват“ криптовалути, инвестиционни схеми или козметика. Някои от тези видеа достигат над 10 милиона гледания преди премахването им.

Механизъм на манипулацията

Технологията за deepfake се базира на **дифузионни генеративни модели** (като *Stable Diffusion* и *DeepFaceLab*), които позволяват с минимален набор от видеоматериали да се генерират изключително реалистични изрази и реч. Рекламните агенции (в някои случаи от сивия пазар на Китай и Източна Европа) създават фалшиви endorsements, за да стимулират доверие и конверсии.

Етични и правни измерения

1. **Нарушение на правото на публичен образ** – deepfake видеата използват лицето и гласа на реални хора без съгласие, което представлява дигитална форма на кражба на идентичност.
2. **Дезинформационен ефект** – потребителите възприемат съдържанието като автентично, особено в кратки формати (Reels, Shorts).
3. **Пазарна манипулация** – рекламите използват фалшив авторитет, което подкопава пазарната прозрачност.
4. **Правна празнина** – повечето юрисдикции все още нямат закони, които директно санкционират създаването на комерсиални deepfake реклами.

Етична последица: разрушаване на доверието

Floridi (2023) нарича това „*пост-истинна етика*“ – когато технологиите подменят епистемологичната граница между „реално“ и „създадено“. Deepfake маркетингът подкопава не просто потребителското доверие, а самата идея за **възможност на истина в цифровата комуникация**.

Amazon Recommendation Algorithms: алгоритмична автономия на потреблението

Фактическа основа

Amazon използва **рекурентни невронни мрежи и колаборативни филтриращи алгоритми** за своите системи за препоръки още от края на 2000-те. Днес повече от **35% от покупките в Amazon** се генерират чрез препоръчителния механизъм, според данни на *McKinsey & Company (2023)*. Алгоритъмът анализира историята на покупките, времето на разглеждане на продуктите, рейтингите и поведението на други потребители с подобен профил.

Механизъм на самоподдържаща се пристрастност

Този модел създава т.нар. **алгоритмичен затворен цикъл**:

1. Потребителят получава препоръки, основани на предишните му избори.
2. Вероятността да избере нещо извън вече предпочитаните категории намалява.

3. Алгоритъмът приема това като потвърждение и стеснява бъдещите предложения.

Резултатът е **псевдоперсонализация**, при която разнообразието на избори изчезва, а системата създава *предиктивна хомогенизация на потреблението*.

Етични и пазарни измерения

- **Манипулативен дизайн** – алгоритъмът е оптимизиран не за интереса на потребителя, а за максимална конверсия и задържане в платформата.
- **Невидим монопол върху вниманието** – потребителят остава в „екосистемата на Amazon“, без реален стимул да излезе от нея.
- **Пазарна изкривеност** – препоръките подсилват продажбите на продукти, които вече са популярни, което създава ефект на кумулативно предимство („the rich get richer“).
- **Правен аспект** – към момента няма регулация, която да изисква прозрачност за начина, по който препоръчителните системи ранжират продукти.

Казус	Основен проблем	Етичен риск	Правен вакуум
Cambridge Analytica (2018)	Използване на лични данни за психологическо профилиране и политическа манипулация	Подкопаване на демократичната автономия	Липса на регулация за психологическо моделиране
Deepfake маркетинг (TikTok, 2024)	Генерирани лица на известни личности за реклами	Симулация на автентичност и измамно доверие	Липса на правен статус на „дигиталната личност“
Amazon Algorithms	Затворен цикъл на препоръки и пристрастия	Ограничаване на избор и когнитивна свобода	Липса на прозрачност в алгоритмичното ранжиране

Заклучение

В разгара на дигиталната епоха, където поведението на човека се анализира, предсказва и моделира в реално време, **въпросът за границата между влияние и манипулация** се превръща в основен етичен и философски проблем. Това размиване не е страничен ефект на технологиите, а **структурен синдром на обществото на данните**, което измерва и оптимизира всичко — включително човешката спонтанност.

1. От технологичен детерминизъм към етичен синергизъм

Твърдението, че технологиите неизбежно водят до манипулация, отразява форма на **технологичен детерминизъм** — схващането, че машините управляват развитието на обществото. Но по-задълбоченият анализ показва, че **проблемът не е в алгоритъма, а в неговата ценностна рамка**. Алгоритъмът сам по себе си е неутрален — той е израз на логиката на оптимизация. Етичният въпрос възниква тогава, когато оптимизацията се прилага към човешкото съзнание, без морална цел или граница. Както пише Floridi (2023), *„когато изчислителната рационалност се освободи от моралната рационалност, тя се превръща в инструмент на властта, а не на знанието“*.

Следователно, задачата на етиката в дигиталната икономика не е да спре технологиите, а **да ги вгради в морален синергизъм**, в който алгоритмите не моделират хората, а подпомагат тяхната когнитивна автономия.

2. Свободата в условията на предсказуемост

Класическото разбиране за свобода — като възможност за избор между алтернативи — губи смисъл в свят, в който алгоритмите знаят нашите предпочитания по-добре от самите нас. Поведението ни вече е **обект на статистическа предвидимост**, а не на философска непредвидимост. Това налага ново тълкуване: Свободата в епохата на алгоритмите трябва да се разбира **не като избор**, а като **осъзнаване** — способността да разпознаем, че нашите избори са моделирани. Истинската автономия се измества от равнището на действие към равнището на рефлексия.

Свободата днес не е право да избереш, а умение да осъзнаеш кой избира вместо теб.

Това преосмисляне връща философската мисъл към Кантовото определение на автономията – не като отсъствие на външна принуда, а като **самоналагане на разумен морален закон**. В контекста на дигиталната икономика този морален закон е **етиката на прозрачността и осъзнатостта** – човекът трябва да знае кога и как неговото съзнание е превърнато в инструмент за прогнозиране.

3. Непредсказуемостта като форма на етично съпротивление

В алгоритмична среда, където стойността се създава чрез предсказуемостта на поведението, **непредсказуемостта се превръща в най-висша форма на етична автономия**. Това означава не хаос, а *творческа несъвместимост с машинната логика*.

Както Zuboff (2019) подчертава, поведенческият капитализъм се захранва от „предвидимите остатъци“ на човешкото действие. Следователно, **етичното действие е това, което не може да бъде предсказано от системата** — акт на когнитивна независимост, който избягва алгоритмичния модел.

Тук се появява нова категория в етиката на дигиталната автономия — *епистемологична съпротива*: способността на индивида да защитава правото си да не бъде напълно прозрачен за машините. Това може да се прояви чрез:

- съзнателно нарушаване на алгоритмични шаблони;
- избори, които не следват очакванията на модела;
- отказ от постоянна оптимизация на себе си според данните.

В този смисъл **човешката грешка, непоследователност и ирационалност** придобиват морална стойност — те са доказателство за живото, непредсказуемо съзнание.

4. Етика на съзнанието, а не на данните

Съвременната етика на изкуствения интелект твърде често се концентрира върху защитата на данните (data ethics). Но както показва развитието на алгоритмичните системи, **истинският обект на защита трябва да бъде съзнанието**, а не информацията.

Когато данните за поведението се използват за моделиране на емоции, мотивации и решения, те вече не са просто лични — те са *ментални проекции*.

Етичният въпрос следователно се измества от „кой притежава данните?“ към „кой притежава интерпретацията на съзнанието?“.

Това изисква нова парадигма: **когнитивна етика** — дисциплина, която защитава правото на човека да остане мисловно автономен. Основен принцип на тази етика би бил: *„Алгоритъмът може да познава поведението ми, но не и да определя мисълта ми.“*

5. Отговорността като хоризонт на дигиталното човечество

В епохата на автономни алгоритми отговорността не може да остане само индивидуална. Тя трябва да се мисли **системно** — като колективен морален хоризонт, в който дизайнери, компании и законодатели носят споделена отговорност за когнитивните ефекти на технологиите.

Както отбелязва Helbing (2022), всяка система, която влияе върху общественото поведение, трябва да се подчинява на принципа на *„етична рефлексия в реално време“* - постоянна оценка на въздействието ѝ върху автономията на човека.

Това не е просто регулаторен механизъм, а нов тип **етична култура**, при която човешкото достойнство се мисли не като статична даденост, а като динамична способност да останем *човешки* в контекста на машинното.

ИЗКУСТВЕН ИНТЕЛЕКТ В МАРКЕТИНГОВОТО ПЛАНИРАНЕ: ВЪЗМОЖНОСТИ ЗА ПОДОБРЯВАНЕ НА ЕФЕКТИВНОСТТА

Marketing Planning Using Artificial Intelligence: Opportunities for Effectiveness Improvement

Гергана Тодорова¹,
e-mail: gtodorova811@gmail.com¹

Абстракт

Изкуственият интелект навлиза все по-осезаемо във всички сфери на съвременния живот, включително и в маркетинговите процеси. AI подпомага взимането на решения и оптимизирането на разходите като предлага опции за анализиране и обобщаване на разнородни данни от външни и вътрешни системи, платформи и приложения. Настоящият доклад предлага преглед на основните приложения на интелигентните системи в маркетинговото планиране и управлението на маркетинговите дейности в дигитална среда. Представят се примери за автоматично анализиране на данни, прогнозиране на резултати и разграничаване на модели потребителско поведение. Изкуственият интелект дава възможност за бързо установяване на неефективни маркетингови дейности, намаляване на грешките и автоматично разпределение на ресурси. AI може да се използва като инструмент за повишаване на ефективността на маркетингово планиране и подобряване възвръщаемостта на инвестициите.

Abstract

Artificial intelligence is increasingly transforming all spheres of modern life, including marketing processes. AI supports decision-making and cost optimization by providing tools for analysing and summarising diverse data from internal and external systems, platforms and applications. This paper is an overview of the main applications of intelligent systems in marketing planning and the management of marketing activities in a digital environment. Examples of automated data analysis, result forecasting, identification of consumer behaviour patterns are provided. Artificial intelligence enables the quick detection of inefficient marketing activities, reduces errors and facilitates the automatic allocation of resources. AI can be used as a powerful tool for enhancing the efficiency of marketing planning and improving the return on investment.

Ключови думи: изкуствен интелект, маркетингово планиране, анализ на данни, ефективност, възвръщаемост.

JEL: M31, C55, O33

Увод

Развитието на дигиталния маркетинг е изключително динамично и поставя компаниите на пазара пред предизвикателството да устоят на темпото и да не отстъпят конкурентното си предимство (Kolyandov & Radev, 2021). Съвременната динамична обстановка, в която се намира всеки един бизнес, налага конкретни модели за анализ на обкръжаващата среда, които не просто отчитат резултати, а предлагат тенденции, изказват предположения и допринасят взимането на адекватни решения и мерки в тази посока. Дигиталният маркетинг все повече се стреми да предложи персонално изживяване за потребителя - компаниите се учат да разбират своите клиенти и да им предлагат повече, на по-ниска цена. Конкурентното предимство на всеки бизнес

¹ Докторант към катедра Индустриален бизнес, Бизнес факултет, УНСС, ORCID 0009-0007-5873-4217

се оказва не неговият най-добър, иновативен и качествен продукт, а отношението към потребителя.

Познаването и разбирането на действията, които потребителите предприемат, дава основание на бизнеса да насочва конкретни продукти към най-заинтересованите хора. Всеки онлайн комуникационен канал, който организацията използва, е потенциален „создател“ на устойчива и трайна връзка, караща конкретния продукт или марка да се превърнат в краен избор за потребителя. Дигиталната среда и ритъмът, който тя налага, предполагат висока скорост на промяна на обстоятелствата и създават необходимост от много бърза реакция и взимане на решения. В този контекст, в продължение на създадените предизвикателства през последните години, се създават решения, в които изкуственият интелект бързо и устойчиво променя крайните резултати. От абстрактно технологично понятие, изкуственият интелект се превръща в реален инструмент, който може да бъде използван на всяка крачка в сферата на маркетинга.

Показатели за ефективност

Позиционирането в социалните мрежи е съществена част от успеха на всеки бизнес в наши дни. Всяка платформа, която се използва за тази цел, предлага детайлни статистики и отчети за анализиране на тенденциите сред потребителите (Cheng, Lam, & Chiu, 2020). Повечето канали за комуникация в онлайн пространството предлагат десетки инструменти за анализ на посетителите – според демографски показатели, препращащи страници (landing pages), изходящи страници и т.н., и всеки един бизнес трябва добре да познава и използва тези данни (Kolyandov & Radev, 2023). Работи се с множество платформи, съответно всяка организация натрупва огромен обем данни от разнородни източници и в различни формати. Това създава необходимост от създаване и внедряване на системи, които позволяват бърз, ефективен и интегриран анализ на тези данни, независимо от тяхната разнородност. Допълнително предизвикателство е наситеността на пазара и презадоволеността на потребителите – в наши дни те вече не просто имат достъп до информация – тя ги залива постоянно и отвсякъде: чрез социални мрежи, e-mail маркетинг, персонализирани реклами, push уведомления и др. Потребителят има богат избор от продукти или услуги за задоволяване на една и съща нужда, които всяка компания рекламира активно в интернет, превръщайки дигиталното пространство в силно конкурентна среда. В резултат, логично, цената на рекламата се покачва и компаниите вече не могат да си позволят да рекламират без ясна представа какво, на кого и как. Именно тук се намесват интелигентните системи с подходящия инструмент – структурирана информация, графики, предвиждания и съвети, с което помагат да се вземат правилните решения навреме, да се оптимизират разходите и да се подобри ефективността.

Нека разгледаме някои ключови показатели за ефективност, които нито един бизнес не може да си позволи да не следи:

- *Ангажираност на потребителите* - този измерител е особено ценен, когато става въпрос за фирмения уеб сайт на бизнеса. Това място играе ролята на „портал“, през който достига основната информация до най-голяма част от потребителите, които вече проявяват интерес към конкретен продукт или услуга. Съществена информация е колко са уникалните посетители на сайта, колко хора са поискали да получават повече информация (под формата на мейли, известия и т.н.), колко хора от първоначално посетилите сайта са се завърнали отново, колко време прекарват потребителите в сайта и на кои страници.
- *Източници на трафик* - този индикатор показва откъде идва трафикът на сайта на бизнеса. Основните „възможности“ са потребителят да е открил сайта чрез търсачка, да е кликнул върху уеб банер или да е пренасочен от партньорски уеб сайт. По този начин може да бъде проследена и ефективността на кампания, да се разбере доколко насочените съобщения достигат адекватни клиенти (чрез т.нар. Bounce Rate – брой напуснали потребители, които не са прегледали нито една страница).

- *Показатели на онлайн кампании* - тези показатели помагат за точното определяне действията на потребителите в дадена кампания. Разглеждат се четири основни елемента – CPM, или още цена за 1000 импресии; CPC – цена за едно кликване върху определен вид реклама; CTR – съотношение между импресии и кликове върху конкретна реклама; CPA – цена за придобиване, стойността на извършено от потребителя действие, което е ценно за бизнеса.
- *Взаимодействия в социалните мрежи* - социалните мрежи са ключов елемент от маркетинговите стратегии на почти всеки бизнес, затова е добре компанията да е добре информирана за нейното развитие в тази среда. Тук се разглеждат различни показатели за различните социални мрежи – брой импресии, достигнати хора, брой реакции и следване на избрана връзка към конкретна цел.
- *Преглеждания на „целеви“ страници (Landing Pages)* – това е мястото, което бизнесът подготвя специално, за да обърне внимание на свой конкретен продукт или услуга. Те обикновено са мястото, до което водят кампаниите на фирмата от други онлайн комуникационни канали – важно е да се проследи доколко хората взаимодействат със съдържанието на подобна страница (извършват ли кликания, достигат ли до друго място в сайта, висок ли е bounce rate-показателят и др.)
- *Показвания при органично търсене* - потребителите в днешно време достигат до информация и конкретни сайтове, използвайки интернет търсачки, като Google, Bing, Yahoo. През 2018, близо 50% от потребителите достъпват интернет, посредством някоя от по-известните търсачки. Важно е да се отчете доколко от посетителите на конкретен фирмен сайт го достигат именно по този начин – така може да бъде взето решение за по-активна реклама в средата на търсачките, да се наложи оптимизиране на уеб сайта (SEO) и т.н.
- *Трафик от различни устройства* - важно за всеки бизнес е да предложи адекватен вариант на своя уеб портал както за настолни компютри, така и за мобилни устройства. През 2017, над 50% от световния онлайн трафик се генерира от мобилни устройства - важно е всеки конкретен бизнес да познава добре конкретните цифри за своя уеб сайт, за да може да приеме или отхвърли необходимостта от специални усилия в подготвянето на различни варианти на портала си.

Разбира се, всеки бизнес може да открие и прецени, че други показатели се явяват от съществена важност за неговата сфера на дейност. По тази причина списъкът с KPIs може да се окаже безкраен и това би било оправдано, в случай, че всеки показател отговаря на условията за адекватност, изложени по-горе. Именно заради това най-ценното всъщност се оказва информацията – спрямо целите си и стратегията си всяка компания избира различни показатели, които определя за важни, но винаги на база данните, с които вече разполага. Изкуственият интелект, който вече е вграден във всяка една платформа, работи именно с това.

Най-разпространени решения за онлайн реклама. Въвличане на изкуствен интелект в ежедневната работа на платформите.

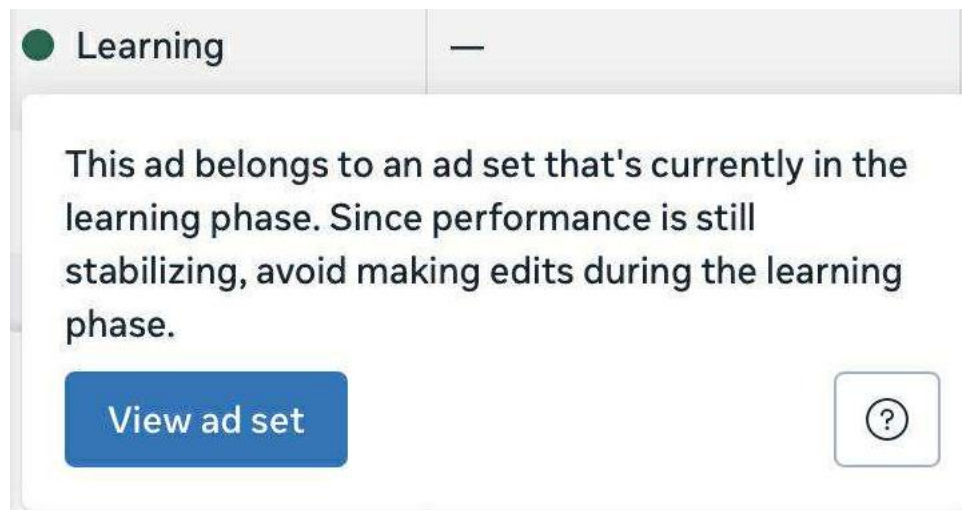
Продукти на Meta

Meta Ads Manager е платформа, която интегрира напреднали технологии за оптимизация и анализ на рекламните кампании. Основата на тази интеграция са изкуственият интелект и машинното обучение, които позволяват на системата да обработва огромни обеми от данни, да идентифицира модели на потребителското поведение и да персонализира рекламното съдържание за отделни аудитории. Освен, че системата автоматично групира потребителите по различни критерии (демографски, географски и др.), тя дава възможност и на компанията сама да сегментира клиентите си, като след това допълва вече наличната информация с анонимизирани данни. В допълнение, чрез модели за класиране на рекламите, платформата автоматично определя кои изображения и послания ще имат най-добър ефект за конкретния

потребител, като използва комплексни „user embeddings“ и алгоритми за препоръки. (Business Insider, 2024)

Технологиите за автоматизация в Meta Ads Manager, като Advantage+ Campaigns, дават възможност на рекламодателите да дефинират основни параметри – цел, бюджет и креативи – след което платформата автоматично оптимизира разпределението на бюджета и избора на аудитории. Това позволява кампаниите да се адаптират в реално време, като неефективните варианти се ограничават, а успешните се засилват, което води до по-висок Click-through rate (CTR) и по-добра възвращаемост на инвестицията.

Автоматизираните оферти и интелигентното разпределение на бюджета намаляват риска от загуби при динамично променящите се пазари, т.к. Meta Ads Manager предоставя мощни инструменти за анализ в реално време. Платформата оценява качеството и потенциала на рекламите чрез показатели като „Opportunity Score“ и извършва кръстосано обучение на данните от Facebook, Instagram, Messenger, Threads и т. нар. „Audience network“. Това означава, че поведението на потребителя на една платформа може да подобри представянето на рекламите на друга, като се постига по-висока релевантност и ефективност на кампаниите. Платформата оценява и „развитието“ на рекламата и автоматично дава индикация, когато тя вече се е „изчерпала“. В този случай, разбира се, е необходима човешка намеса за актуализиране на посланията или съдържанието (изображения, видео), но междувременно системата автоматично преразпределя бюджета към по-ефективни реклами. Оптимизацията именно на съдържанието е друг ключов аспект, който води до по-добри резултати. Чрез генеративни алгоритми платформата създава множество варианти на рекламите – изображения, видеа, текстове, и адаптира формата им за различни канали като Stories, Reels или Feed. Това позволява на маркетинговете бързо да тестват комбинации и да идентифицират най-ефективните послания, без да се изисква значителен ресурс за ръчно производство на съдържание. (Mediaweek Australia, 2025)



Фигура 6 - Съобщение за кампания във фаза на "обучение" в Meta Ads Manager

В резултат на тези технологии, кампаниите в Meta Ads Manager се характеризират с по-висока ефективност и скорост на реакция. Автоматизацията и анализът на данни дават възможност за оптимизиране на бюджета, таргетиране на най-релевантните аудитории и постигане на по-добър резултат с по-малко усилия. (We Are North UK, 2024)

Meta Ads Manager се превръща в инструмент, който не само улеснява управлението на кампании, но и предоставя стратегическо конкурентно предимство на кампаниите, които умеят да

интегрират данните и технологиите си за взимане на интелигентни, базирани на факти решения в дигиталния маркетинг.

Продукти за анализ на Google

Google Analytics 4 (GA4) се отличава с модерна архитектура, която дава на маркетинговите специалисти възможност да наблюдават, анализират и оптимизират поведението на потребителите — не само на отделни сайтове, но и сред различни източници - приложения, каталози и др. (Discover Web Tech, 2025). GA4 внедрява събитийно-ориентиран (event-based) модел на данни, при който всяко взаимодействие на потребителя — клик, скрол, видеопреглед, покупка или друго — се третира като събитие. Този подход заменя традиционната сесийно-странична („session/pageview“) логика и дава по-детайлна картина на поведението (FutureProfilez, 2025). Това позволява на маркетинговете да разберат по-гладко как потребителите преминават през фунията, кои стъпки предизвикват отдръпване и къде се крият възможности за оптимизация.

Ключов компонент е изкуственият интелект и машинното обучение. GA4 автоматично предлага прогнози — например кои потребители имат по-висок шанс да конвертират или да отпаднат от фунията. Платформата предлага и автоматично извличане на знания (напр. аномалии, необичайни модели) посредством алгоритми с машинно обучение (Google Marketing Platform, 2025). Това дава възможност на бизнеса да се превърне от реактивен в проактивен: вместо просто да анализира минали действия, маркетингът може да прогнозира и да действа навреме. GA4 значително улеснява анализите чрез мощни функции за визуализация и сегментация: Explorations, кохортен анализ, анализ на „пътеката“, анализ на потребителския „живот“ в дадена платформа и др. (Google Marketing Platform, 2025b). Например, можете да визуализирате пътя на потребителя до покупка („path exploration“), да разберете кое сегментно поведение е най-печелившо или да наблюдавате кои фактори влияят на задържането на потребителите. Освен това, платформата поддържа интеграции с големи системи: експорт към BigQuery, свързване с Google Ads, Search Console и други API (Google Marketing Platform, 2025b). Така маркетинговите екипи могат да обединят данни от различни източници и да получат холистичен поглед върху потребителското поведение и маркетинг резултати.

От гледна точка на ефективността, използването на GA4 води до няколко конкретни ползи:

- по-точно таргетиране и сегментация, благодарение на машинното обучение и анализа на поведение (Google Marketing Platform, 2025a);
- по-добро разбиране на фунията и оптимизация на отделни етапи чрез path и funnel анализ (The Knowledge Academy, 2025);
- ускорено взимане на решения: автоматичните „знания“ и прогнозите минимизират нуждата от дълъг ръчен анализ (Bismart Blog, 2025);
- подобро измерване на ROI: интеграции с рекламни системи и експорт на данни към BigQuery позволяват свързване на рекламни разходи с реални бизнес резултати (Google Marketing Platform, 2025);
- съобразяване с поверителността и бъдещите регулации: GA4 предлага функционалности като „consent mode“ и по-голям контрол върху данните (Loves Data, 2025).

В заключение, Google Analytics 4 не е просто инструмент за отчитане, а стратегически компонент в дигиталния маркетинг, който подпомага компаниите да използват данните си по-интелигентно — да взимат решения не само на база минало, а да прогнозират и моделират бъдещето на потребителското поведение (Google Marketing Platform, 2025a; Discover Web Tech, 2025).

Предизвикателства

Едно от ключовите предизвикателства при използването на технологии като ИИ и машинно обучение и в продуктите на Meta, и в продуктите на Google, е т.нар. „алгоритмично изкривяване“. Макар да съществуват редица други проблеми, които могат да бъдат съществена трудност пред маркетолозите, алгоритмичните изкривявания сами по себе си водят до погрешно разходване на бюджети и ограничаване достигането на определени кампании. В постановката по-долу е визуализирана реална рекламна кампания в Meta Ads Manager. Бизнес потребителят е въвел определени критерии, на които трябва да отговарят потенциалните клиенти – регион, възраст, пол, както и примерни интереси. Използвайки Advantage+ насочване, Meta прогнозира достигане на хора в интервала 1.2 – 2 милиона души. Алгоритмите на Advantage+ се опитват да комбинират както предварително зададените от потребителя критерии, така и вече натрупаното знание за останалите кампании в Ads Manager. По тази причина, Meta открива „ядро“ от хора, които реагират положително на рекламата, и започва да я показва все по-често на тях, вместо да търси нови потенциални клиенти. В примера по-долу, Advantage+ на Meta открива 144 хиляди души, на които е показал рекламата средно 4.61 пъти. Това води до разходване на бюджет без реална полза – потребителите вече са видели рекламата твърде много пъти и спират да ѝ „обръщат внимание“. В резултат, Meta извежда съобщение за т.нар. „creative fatigue“ – необходимост от промяна на изображенията или видеоклиповете в рекламата, за да могат потребителите да „забележат“ кампанията отново. Всъщност, проблемът тук не е в изчерпване на съдържанието, а в прекомерното показване на рекламата пред една и съща аудитория – казус, възникнал именно от самообучението на платформата. Решението е намеса на маркетолога и създаване на паралелна кампания с различно насочване, което междувременно води до допълнителен разход.

Campaign ↑↓	▼ yy ▼	Budget ↑↓	Atti set↓	Results ↑↓	Reach ↑↓	Impressions ↑↓	Cost per result ↑↓	Amount spent ↑↓
/ Sales, Wide	volume	€35.00	7...	63	144,744	665,979	Per Purchase	€1,039.36
📊 Charts 🖋 Edit 🔄 Duplicate 📊 Compare ...		Daily	Website purcha...				Per Purchase	

Фигура 7 - Алгоритмично изкривяване на кампания в Meta Ads Manager

Заклучение

Изкуственият интелект отдавна не е просто абстрактна концепция в маркетинга, а реален инструмент, който трансформира начина, по който компаниите реализират своите онлайн кампании. Той променя самото маркетингово планиране – от интуитивно и базирано на опит, към стратегически процес, основан на данни и непрекъсната оптимизация. Чрез интеграцията на алгоритми за машинно обучение и прогнозен анализ, платформи като Meta Ads Manager и Google Analytics 4 позволяват на бизнеса да реагира в реално време на поведението на потребителите, да минимизира неефективните разходи и да увеличи възвръщаемостта на инвестициите (Kolyandov & Radev, 2021).

Изкуственият интелект се превръща в своеобразен „посредник“ между огромния информационен поток и нуждата на маркетолога от ясни и логични изводи. Той не просто автоматизира рутинните дейности, а осигурява възможност за по-дълбоко разбиране на клиентските навици и за персонализиране на посланията с висока точност. Въпреки това, автоматизацията поражда и нови предизвикателства – особено свързани с прозрачността на решенията и риска от прекомерен разход за привидно добре работеща кампания. Затова, ролята на маркетолога става все по-аналитична и стратегическа – да разбира, адаптира и насочва алгоритмите, вместо да бъде техен пасивен наблюдател. Компаниите, които успеят да намерят този баланс между технологична ефективност и човешка изобретателност, ще бъдат лидерите в новата ера на интелигентния маркетинг.

Източници

1. Bismart Blog. (2025). *The 11 critical features of Google Analytics 4 (GA4)*. <https://blog.bismart.com/en/kale/google-analytics-4>
2. Business Insider. (2024). *Meta integrates AI into Ads Manager to improve advertiser performance*. <https://www.businessinsider.com/matt-steiner-meta-ads-ai-power-list-2024>
3. Cheng, W., Lam, E. R., & Chiu, D. (2020). Social media as a platform in academic library marketing: A comparative study. *The Journal of Academic Librarianship*, 46(5), 102231. <https://doi.org/10.1016/j.acalib.2020.102231>
4. Discover Web Tech. (2025). *Google Analytics 4 explained*. <https://discoverwebtech.com/digital-marketing/seo/google-analytics-4-explained/>
5. FutureProfilez. (2025). *5 exciting new features of Google Analytics 4 (GA4)*. <https://futureprofilez.com/5-exciting-new-features-of-google-analytics-4-ga4/>
6. Google Marketing Platform. (2025a). *Benefits of analytics for data driven marketing*. <https://marketingplatform.google.com/about/analytics/benefits/>
7. Google Marketing Platform. (2025b). *Analytics & data analysis features list*. <https://marketingplatform.google.com/about/analytics/features/>
8. Loves Data. (2025). *Google Analytics 4 key features and benefits*. <https://www.lovesdata.com/blog/google-analytics-4-key-features-and-benefits/>
9. Kolyandov, S., & Radev, R. (2021). *Internet marketing: Modern advertising models for reaching new customers*. *Trakia Journal of Sciences*, 19(1), 122-129. https://www.researchgate.net/publication/360017342_INTERNET_MARKETING_MODERN_ADVERTISING_MODELS_FOR_REACHING_NEW_CUSTOMERS
10. Kolyandov, S., & Radev, R. (2023). *Internet marketing based on a data ecosystem*. *Trakia Journal of Sciences*, 21(Suppl. 1), 100-105. <https://doi.org/10.15547/tjs.2023.s.01.017>
11. Mediaweek Australia. (2025). *Meta unveils new AI tools for advertisers in Ads Manager*. <https://www.mediaweek.com.au/meta-unveils-new-ai-tools-for-advertisers-in-ads-manager>
12. Smart Insights. (n.d.). *Statistics on search engine marketing usage and adoption to inform your search engine marketing strategies and tactics*. <https://www.smartinsights.com/search-engine-marketing/statistics-on-search-engine-marketing-usage-and-adoption/>
13. The Knowledge Academy. (2025). *Top Google Analytics features for your website*. <https://www.theknowledgeacademy.com/blog/google-analytics-features/>
14. Thakur, R. (2015). *The 7 most important digital marketing KPIs to track*. LinkedIn. <https://www.linkedin.com/pulse/7-most-important-digital-marketing-kpis-track-ram-thakur>
15. We Are North UK. (2024). *Graham Green case study: Meta Advantage Plus Campaign*. <https://wearenorth.uk/our-work/graham-green-meta-advantage-plus>

ПРИЛАГАНЕ НА ГЕЙМИФИКАЦИЯТА В ЕЛЕКТРОННОТО ОБУЧЕНИЕ КАТО ПОДХОД ЗА ПОВИШАВАНЕ НА МОТИВАЦИЯТА

Applying Gamification in E-Learning as an Approach to Enhancing Motivation

Кремена Маринова-Костова¹,
e-mail: k.marinova@uni-svishtov.bg¹

Абстракт

Геймификацията представлява използване на игрови механизми и елементи в неигрови контекст с цел повишаване на мотивацията и ангажираността. Особено подходяща сфера на приложение е образованието, като значителни възможности се откриват при прилагането на подхода в електронното обучение. Надграждането на традиционните платформи за електронно обучение с технологии, предлагащи игрови елементи – като интерактивни платформи, тестове, имерсивни технологии и системи с изкуствен интелект – прави обучението по-персонализирано и ефективно. Прилагането на геймифициращия подход крие и рискове, най-вече по отношение на повърхностна геймификация („pointsification“), прекомерна конкуренция, която измества фокуса от ученето, както и проблеми с достъпността и етиката. Поради тази причина успешното ѝ прилагане е резултат от съчетаването на добре дефинирани педагогически цели и подходящи технологии, които да стимулират мотивацията и ангажираността на обучаемите.

Abstract

Gamification is the use of game mechanisms and elements in a non-gaming context to increase motivation and engagement. Education is a particularly relevant field of application, with significant opportunities opening in the application of the e-learning approach. Upgrading traditional e-learning platforms with technologies offering game elements – such as interactive platforms, quizzes, immersive technologies, and artificial intelligence systems – makes learning more personalized and effective. The application of the gamification approach also carries risks, especially in terms of superficial gamification ("pointsification"), excessive competition that shifts the focus away from learning, as well as accessibility and ethics issues. For this reason, its successful implementation is the result of a combination of well-defined pedagogical goals and appropriate technologies to stimulate the motivation and engagement of learners.

Ключови думи: геймификация, електронно обучение, мотивация, образователни технологии, MDE модел

JEL: D83, I23, O33

Увод

Навлизането на новите технологии и разширяването на обхвата на онлайн обучението са водещите предпоставки за прилагането на нови педагогически подходи за привличане и задържане на интереса на обучаемите, особено на представителите на поколение Z, за които дигиталните технологии, онлайн комуникацията и интерактивността са естествена част от ежедневието. Един от най-успешните съвременни образователни подходи е геймификацията –

¹ Head Assistant Professor, PhD (Econ). Department of Business Informatics, D. A. Tsenov Academy of Economics, Svishitov, ORCID ID: 0000-0001-8035-7891

метод, който използва игрови принципи и механизми за стимулиране на мотивацията в процеса на учене. Този подход не само увеличава ангажираността и удовлетворението на учащите, но и подпомага развитието на ключови когнитивни и социални умения чрез игрови сценарии и състезателни елементи. Освен това геймификацията предоставя възможност за индивидуализиране на учебния процес, като позволява адаптация към различни стилове на учене и темпове на усвояване на материала.

Геймификация: концепция и историческо развитие

Дефиниция на геймификацията

Терминът геймификация (игровизация) е неологизъм, въведен от британския програмист Ник Пелинг, който разработва интерфейс, подобен на игра, за банкомати и вендинг машини с цел да направи електронните транзакции по-привлекателни (Pelling, 2011). Много бързо терминът навлиза и в други сфери, като неговата дефиниция е обект на трудовете на редица учени. Зихерман и Кънингам я определят като процес на игрово мислене и игрови механики за ангажиране на потребителите и решаване на проблеми (Zichermann & Cunningham, 2011). От своя страна Бърк счита, че геймификацията е инструмент за проектиране на поведение, развиване на умения и насърчаване на иновации, който в комбинация с други технологии и тенденции може да причини големи прекъсвания в иновациите, управлението на представянето на служителите, образованието, личностното развитие и ангажираността на клиентите (Burke, 2012). Хуотари и Хамари акцентират върху ролята на геймификацията като процес на подобряване на услугата с достъпни цени за игрови преживявания, който цели да подпомогне създаването на стойност от страна на потребителя като цяло (Huotari & Hamari, 2012). Тяхното виждане е доразвито от Вербах, който дефинира геймификацията като процес на превръщане на дейностите в по-игроподобни, фокусирайки се върху ключовото пространство между компонентите, които изграждат игрите, и холистичното преживяване на игровостта (Werbach, 2014). Научната общност се обединява около определенитето, дадено от Детърдин и съавтори (Deterding, Dixon, Khaled, & Nacke, 2011), подкрепено от Вербах и Хънтър (Werbach & Hunter, 2012), които определят геймификацията като използване на игрови елементи и техники в неигрови контекст, което определение приемаме и ние.

Историческо развитие на геймификацията

Макар и сравнително нова концепция, идеята за геймификация води своето начало още от 1908 г., когато в скаутското движение се въвеждат значки като награда за постижения и придобити умения. През 70-те и 80-те години на 20-ти век се предлага идеята да се направи работата по-интересна и ангажираща, като трудовите задължения се превръщат в своеобразна игра. В реалния бизнес геймификацията за първи път се използва от авиокомпаниите през 1981 г., които популяризират игрови компоненти като точки, нива и награди. През 80-те и 90-те години нараства академичният интерес към темата, а през 1996 г. Ричард Бартъл създава класификация на типовете играчи, която е в основата на разбирането на игровото поведение и дизайна на мотивационни системи (Bartle, 1996). Въвеждането на термина през 2002 г. от Пелинг откроява геймификацията като самостоятелно направление и поставя началото на съзнателното използване на игрови елементи в продукти, услуги и софтуерни решения.

От средата на 2000-те години геймификацията навлиза масово и в дигиталната сфера, като в редица онлайн платформи се добавят игрови механики за повишаване на мотивацията, а множество приложения внедряват системи от значки и постижения с цел ангажиране на потребителите (Sharma & Sharma, 2023), (Christians, 2018).

В последните години геймификацията се превръща в неизменна част от образованието, бизнеса, здравеопазването и фитнес индустрията. По време на пандемията от COVID-19 тя играе ключова роля в поддържането на мотивацията и ангажираността при дистанционно обучение и работа. Тя се развива от система за награждаване в началото на 20-ти век до стратегически инструмент за ангажиране, обучение и повишаване на ефективността в съвременния дигитален свят.

Мотивационни, педагогически и технологични аспекти на електронното обучение

Мотивацията като ключов фактор в процеса на електронното обучение

Мотивацията за учене е ключов психологически механизъм, който насочва поведението на човека и усилията му към постигане на определена цел (Davidovitch & Dorot, 2023). Тя може да се дефинира като вътрешна сила, която в образователен контекст активира и насочва поведението на учащия по позитивен и креативен начин, като същевременно подпомага усвояването на знанията и формирането на полезни учебни навици (Alcívar, Menéndez, & Valencia, 2021). Мотивацията не е физическа характеристика, която може да бъде определена, а се проявява посредством поведението и усилията на обучаемия (Redondo & Martín, 2015).

Мотивацията може да бъде вътрешна и външна. Вътрешната мотивация се свързва с интереса и любопитството на обучаемия, докато външната се провокира от наградите и стимулите, които той очаква. И двата вида мотивация оказват влияние както върху учебните постижения, така и върху ангажираността на обучаемите за постигане на дългосрочни учебни цели. Тя стимулира учащите да посвещават време и усилия на учебните задачи и да проявяват устойчивост при трудности, идентифицирани по време на учебния процес. В допълнение, мотивацията се влияе и от лични, социални и културни фактори, които, приложени от обучителите в образователен контекст, могат да насърчат инициативността, активното учене и да доведат до високи резултати.

Педагогически аспекти на геймификацията в електронното обучение

Геймификацията в електронното обучение може да се определи като използването на игрови механизми в неигрови педагогически ситуации в електронна среда с цел повишаване на интереса, мотивацията и участието на обучаемите. Тя способства за трансформиране на традиционните образователни подходи, като прави ученето емоционално и интерактивно преживяване, използвайки стратегии, познати от игрите – събиране на точки, натрупване на напредък и постигане на нива, предизвикателства, класации и др. (Deterding, Dixon, Khaled, & Nacke, 2011).

От педагогическа гледна точка тя се основава на конструктивистките и мотивационните теории, които акцентират върху водещата роля на обучаемите в процеса (Deci & Ryan, 2000). Игровите механизми помагат на обучаемите да изграждат умения като критично мислене, саморефлексия, способност за решаване на проблеми, а от своя страна системата за награди и обратна връзка поддържа вътрешната им мотивация. По този начин се постига и по-голяма устойчивост на онлайн обучението, въпреки липсата на реална социална връзка между участниците (Hamari, Koivisto, & Sarsa, 2014).

Дигитални инструменти и технологии за геймификация на електронното обучение

В технологично отношение геймификацията в електронното обучение се осъществява чрез разнообразни дигитални инструменти и платформи (Christopoulos & Mystakidis, 2023), (CHE KU MOND, et al., 2023):

- **Инструменти за оценяване и викторини** като Kahoot!, Quizizz и Socrative, които позволяват провеждане на тестове и анкети в реално време, съчетани с визуални ефекти, точки и класации. Те стимулират бързата реакция на обучаемите и създават динамична и неформална среда за учене.
- **Системи за управление на обучението (Learning Management Systems - LMS) с геймификационни модули** като Moodle (плъгин Level Up!) и Canvas Badges, които дават възможност за присъждане на значки, проследяване на постиженията и визуализиране на резултатите, създавайки усещане за напредък в обучението.
- **Интерактивни класни стаи и среди за сътрудничество** като Mentimeter, Nearpod и Padlet, които улесняват активното участие на обучаемите в процеса като ги ангажират посредством попълване на анкети в реално време, дискусии и споделяне на идеи. По този начин пасивното обучение се превръща в колективно интерактивно преживяване.

- **Платформи за колаборативно и проектно-базирано обучение** като Minecraft: Education Edition и Classcraft, които насърчават сътрудничеството, стратегическото мислене и решаването на проблеми под формата на предварително дефинирани игрови сценарии.
- **Изкуствен интелект и адаптивни системи**, които използват данни за миналото поведение на обучаемите с цел да персонализират игровото му преживяване.
- **Потопящи „имерсивни“ технологии**, включващи инструменти за добавена (AR), виртуална (VR) и смесена реалност (MR), които спомагат създаването на дълбоки и реалистични учебни преживявания. Такива инструменти са Merge Cube, EngageVR, ClassVR, Microsoft HoloLens и AR Anatomy Apps и др.

Предимства и рискове на геймификацията в електронното обучение

Предимства

Прилагането на геймификацията в електронното обучение има значим потенциал за подобряване на качеството и ефективността на учебния процес, като влияе върху мотивацията, постиженията и задържането на вниманието на учащите. Въпреки това, ефектът ѝ зависи от начина на проектиране, контекста и индивидуалните характеристики на обучаемите. Сред основните ѝ предимства се открояват следните (Dicheva, Dichev, Agre, & Angelova, 2015), (Hanus & Fox, 2015), (Simsek & Yilmaz, 2025):

- **Повишаване на мотивацията и ангажираността** – геймификацията влияе върху вътрешната мотивация на обучаемите чрез отчитане на постиженията, следене на напредъка и поставяне на предизвикателства, като по този начин се насърчават постоянството в работата и саморефлексията. Игровите сценарии провокират по-дълбока ангажираност по отношение на учебното съдържание и постигат положителна емоционална нагласа към ученето.
- **Развитие на познавателни и социални умения** – игровите дейности способстват на учащите да изграждат критично мислене и да решават проблеми, работейки в екип. Те развиват стратегическите комуникационните и лидерските умения, като насърчават сътрудничеството.
- **Подобряване на учебните резултати** и запомняне на наученото за по-продължителен период в сравнение с традиционните електронни форми.
- **Персонализирано и адаптивно учене** – прилагането на изкуствен интелект в геймифициращите форми на обучение дава възможност на системите да предлагат индивидуализирани предизвикателства на обучаемите и стимули, съобразени с нивото, темпото им на учене и личните предпочитания. Това насърчава автономията и контрола от страна на обучаемите върху процеса.
- **Подобрена обратна връзка и видим напредък** – визуалните стимули като значки, нива, следене на прогреса и класации, правят резултатите от обучението по-осезаеми и лесни за проследяване като по този начин се засилва саморегулацията и увереността на учащите.

Рискове

Въпреки значителния педагогически потенциал на геймификацията, редица автори (Hanus & Fox, 2015), (de-Marcos, Domínguez, Saenz-de-Navarrete, & Pagés, 2014), (Dicheva, Dichev, Agre, & Angelova, 2015), (Toda, Valle, & Isotani, 2017) считат, че нейното прилагане в електронното обучение крие и множество рискове и предизвикателства:

- **Прекомерна зависимост от външни стимули** като точки, значки и класации може да отслаби вътрешната мотивация, тъй като вниманието се фокусира върху наградата, а не

върху самото обучение. Това води до участие проформа, моментно ангажиране и натрупване предимно на краткосрочни знания.

- **Ефект на новостта и загуба на интерес** – в началото игровите елементи са интересни за обучаваните и предизвикват ентузиазъм, но ако се прекалява с тях, те стават скучни, предвидими и губят своята мотивираща роля в учебния процес. За преодоляване на този недостатък геймифициращите компоненти трябва да се обновяват периодично спрямо потребностите на обучаемите.
- **Повишен стрес и социално напрежение** – състезателните механизми като класации и проследяване на напредъка могат да породят тревожност, притеснение и чувство за неуспех, особено ако са публични. По този начин може да се засили социалното сравнение и да се създаде неблагоприятна учебна обстановка, най-вече за обучаемите с по-ниска самооценка.
- **Неподходящ педагогически дизайн** – механичното и необмислено въвеждане на геймификацията, особено ако тя не е съобразена с педагогическите цели на обучението, може да има обратен ефект – да доведе до когнитивно натоварване и разсейване. Ето защо, за да е успешна, геймификацията трябва да отговаря на няколко условия: да има ясно зададени цели и правила, да предоставя обратна връзка на обучаемите и участието им да е доброволно.
- **Технологични ограничения и дигитално неравенство** – ефективното прилагане на геймификацията изисква наличието на подходящо техническо оборудване и софтуерни платформи, което може да представлява пречка за определени социални и етнически групи. Освен това успешната ѝ реализация предполага преподаватели с необходимата дигитална компетентност, която не е налична във всички учебни институции, което потенциално увеличава неравенството в достъпа до качествено електронно обучение.
- **Етични и правни рискове** – геймификацията предполага събиране на данни за поведението на обучаемите, което поставя под въпрос защитата на личната информация и поверителността. Съществува и риск от употребата на манипулативен или прекалено ангажиращ дизайн, който измества фокуса от образователните цели.
- **Културни и индивидуални различия** – ефективността на геймификацията зависи от културните и личностните характеристики на обучаемите. Различните типове личности реагират различно на едни и същи игрови стимули – докато едни се мотивират повече в конкурентна среда, други предпочитат индивидуалните предизвикателства и самостоятелната работа.

Заклучение

Геймификацията се утвърждава като ефективен инструмент в съвременното електронно обучение, комбинирайки игрови механизми с педагогически методи и подходи с цел повишаване на мотивацията, ангажираността и участието на обучаемите в учебния процес. Игровите елементи като значки, проследяване на прогреса, класации, имерсивни преживявания и др. позволяват персонализиране на ученето и насърчаване на развитието на ключови когнитивни и социални умения.

Успешната интеграция на геймификацията в електронното обучение от своя страна изисква внимателно проектиране и съобразяване с индивидуалните особености на обучаемите, поставяне на правилни педагогически цели и отчитане на технологичните възможности на процеса.

За преодоляване на множеството рискове, свързани с прекомерната употреба на геймификацията, пораждащата се конкуренция, технологичните ограничения, етичните аспекти и културните различия, е необходимо прилагането на балансиран и добре структуриран образователен подход. Важно е да се подчертае, че геймификацията не бива да бъде самоцел, а по-скоро инструмент, който допълва традиционното електронно обучение, повишавайки

неговата ефективност и устойчивостта на придобитите знания. Успешното ѝ прилагане е резултат от синергията между педагогическите цели, подходящите технологии и индивидуалните особености на обучаемите.

Редовното обновяване и адаптиране на геймифициращите елементи спрямо индивидуалните потребности и напредъка на обучаемите е ключово за поддържане на висока мотивация и устойчиво участие в учебния процес.

References (TNR 12pt., bold, left)

1. Alcívar, C. M., Menéndez, C. S., & Valencia, A. Z. (2021). Motivation and its Importance in the Classroom Learning Proces sand Teaching. *International Research Journal of Engineering, IT & Scientific Research*, 7(4), 148-154.
2. Bartle, R. (1996). Hearts, Clubs, Diamonds, Spades: Players Who Suit MUDs. *Journal of MUD research*, 1(1).
3. Burke, B. (05 11 2012 r.). *Gamification 2020: What Is the Future of Gamification?* Изтеглено на 07 10 2025 r. от <https://www.gartner.com:https://www.gartner.com/en/documents/2226015>
4. CHE KU MOND, C., MOHAMAD, S. M., SULAIMAN, H. A., SHAHBODIN, F., RAHIM, N. R., & AIZUDIN, A. (2023). A REVIEW OF GAMIFICATION TOOLS TO BOOST STUDENTS' MOTIVATION AND ENGAGEMENT. *Journal of Theoretical and Applied Information Technology*, 101(7), 2771-2782.
5. Christians, G. (2018). The Origins and Future of Gamification .
6. Christopoulos, A., & Mystakidis, S. (2023). Gamification in Education. *Encyclopedia*, 3, 1223–1243. doi:10.3390/encyclopedia3040089
7. Davidovitch, N., & Dorot, R. (2023). The Effect of Motivation for Learning Among High School Students and Undergraduate Students—A Comparative Study. *International Education Studies*, 16(2), 117-127.
8. Deci, E. L., & Ryan, R. M. (2000). The "What" and "Why" of Goal Pursuits: Human Needs and the Self-Determination of Behavior. *Psychological Inquiry*, 11(4), 227–268. doi: https://doi.org/10.1207/S15327965PLI1104_01
9. de-Marcos, L., Domínguez, A., Saenz-de-Navarrete, J., & Pagés, C. (2014). An empirical study comparing gamification and social networking on e-learning. *Computers & Education*, 75, 82-91.
10. Deterding, S., Dixon, D., Khaled, R., & Nacke, L. (2011). From game design elements to gamefulness: defining "gamification". *MindTrek '11: Proceedings of the 15th International Academic MindTrek Conference: Envisioning Future Media Environments*, (pp. 9-15). doi:<https://doi.org/10.1145/2181037.218104>
11. Dicheva, D., Dichev, C., Agre, G., & Angelova, G. (2015). Gamification in Education: A Systematic Mapping Study. *Educational Technology & Society*, 18(3), 75-88.
12. Hamari, J., Koivisto, J., & Sarsa, H. (2014). Does Gamification Work? -- A Literature Review of Empirical Studies on Gamification. *2014 47th Hawaii International Conference on System Sciences (HICSS)*, (pp. 3025-3034). Waikoloa. doi:10.1109/HICSS.2014.377
13. Hanus, M. D., & Fox, J. (2015). Assessing the effects of gamification in the classroom: A longitudinal study on intrinsic motivation, social comparison, satisfaction, effort, and academic performance. *Computers & Education*, 80, 152-161. doi:<https://doi.org/10.1016/j.compedu.2014.08.019>
14. Huotari, K., & Hamari, J. (2012). Defining gamification: a service marketing perspective. *16th International Academic Mindtrek Conference* (pp. 17 - 22). Association for Computing Machinery. doi:<https://doi.org/10.1145/2393132.239313>

15. Pelling, N. (2011). *The (short) prehistory of “gamification”.... Funding Startups (& other impossibilities)*. Retrieved: 07 10 2025 r. от <https://nanodome.wordpress.com:https://nanodome.wordpress.com/category/gamification-2/>
16. Redondo, R. E., & Martín, J. O. (2015). Motivation: The Road to Successful Learning. *Profile: Issues in Teachers' Professional Development*, 17, 125-136.
17. Sharma, D., & Sharma, J. (2023). EVOLUTION OF GAMIFICATION, ITS IMPLICATIONS, AND ITS STATISTICAL IMPACT ON THE SOCIETY. *ShodhKosh: Journal of Visual and Performing Arts*, 2SE(2SE), pp. 8-20. doi:<https://doi.org/10.29121/shodhkosh.v4.i2SE.2023.456>
18. Simsek, E., & Yılmaz, T. K. (2025). A Systematic Review of the Effects of Gamification in Online Learning Environments on Learning Outcomes. *Open Praxis*, 17(1), 166–183. doi:10.55982/openpraxis.17.1.692
19. Toda, A. M., Valle, P. H., & Isotani, S. (2017). The Dark Side of Gamification: An Overview of Negative Effects of Gamification in Education. Ot A. Cristea, I. Bittencourt, & F. Lima (Ред.), *Higher Education for All. From Challenges to Novel Technology-Enhanced Solutions* (Communications in Computer and Information Science ed., Vol. 832). Springer, Cham. doi:https://doi.org/10.1007/978-3-319-97934-2_9
20. Werbach, K. (2014). (Re)Defining Gamification: A Process Approach. *International Conference on Persuasive Technology* (стр. 266–272). Springer, Cham. doi:https://doi.org/10.1007/978-3-319-07127-5_23
21. Werbach, K., & Hunter, D. (2012). *For The Win - How Game Thinking Can Revolutionize Your Business*. Wharton Digital Pre.
22. Zichermann, G., & Cunningham, C. (2011). *Gamification by Design: Implementing Game Mechanics in Web and Mobile Apps*. O'Reilly Media, Inc.

ПРЕДИЗВИКАТЕЛСТВА ПРЕД ПРОФЕСИОНАЛНОТО ОБРАЗОВАНИЕ И ОБУЧЕНИЕ В БЪЛГАРИЯ И ДИГИТАЛНИ ИНСТРУМЕНТИ ЗА СПРАВЯНЕ

Лора Първанова¹

email: lora0parvanova@gmail.com¹

Резюме

Професионалното образование и обучение (ПОО) представлява своеобразен мост между обучението и пазара на труда в икономиките от Индустрия 4.0 и 5.0. В цифровото десетилетие концепцията за автоматизация, използване на големи данни и IoT технологии би била ефективна само в случай, че пазарът на труда е способен да осигури навреме, достатъчно на брой и добре подготвени кадри. Според Европейския център за развитие на ПОО, професионалното образование и обучение е от решаващо значение за подсибяване на служителите и обучаващите се в ЕС с правилните умения, както и за подпомагане на преминаването им към по-екологично, „зелено“ мислене. Докладът се съсредоточава върху целите, които трябва да се постигнат в професионалното образование и обучение до 2030-та година и дигиталните технологии, които са необходима част от този процес. Разглеждат се трудностите, които България среща в процеса на дигитализация като цяло, както и тези пред които се изправят отделните субекти участващи в процеса на професионално образование и обучение. Прави се ретроспекция на използваните дигитални инструменти – LMS, LCMS, AI, LXP и други и необходимостта от тяхното синхронизиране в една обща „екосистема“.

Ключови думи: Професионално образование и обучение, Система за управление на обучението, Платформа за обучение чрез преживяване

Abstract

Vocational education and training (VET) is a bridge between education and the labour market in the economies of Industry 4.0 and Industry 5.0. In the digital decade, the concept of automation, the use of Big Data and IoT technologies would only be effective if the labour market is able to provide timely, sufficient and well-trained personnel. According to the European Centre for the Development of VET, vocational education and training is crucial to equip EU employees and learners with the right skills, as well as to support their transition to a more ecological, “green” mindset. The report focuses on the goals that need to be achieved in VET by 2030 and the digital technologies which are the necessary part of this process. The difficulties that Bulgaria encounters in the process of digitalization as a whole are examined, as well as those faced by the individual entities participating in the process of vocational education and training - the National Agency for Vocational Education and Training, vocational training centers, information and career guidance centers, trainers and trainees. A retrospective of the digital tools used - LMS, LCMS, AI, LXP and others and the need for their synchronization in a common "ecosystem" is made.

Въведение

Динамиката на пазара на труда в България е висока, а спецификата на търсене има големи изисквания към предлагането на работна ръка, а именно – подготовка на специалисти в кратки срокове, с постоянно променящи се умения и компетентности. Професионалното образование и обучение има капацитета да предложи адекватна реакция в случай, че бъде модернизирано и насочено в правилната посока. С приемане на Закона за професионалното образование и обучение през 1999 г. е създадена

¹ Задочен докторант към катедра „Бизнес информатика“, СА „Д. А. Ценов“ - Свищов

Национална агенция за професионално образование и обучение (НАПОО), която лицензира двете основни форми за предоставяне на услуги в сферата на ПОО – Център за професионално образование и обучение (ЦПО) и Център за информация и професионално ориентиране (ЦИПО). НАПОО е координатор на системата ReferNet¹ за България и работи по изпълнение на рамковите програми на Европейския съюз относно професионалното образование и обучение.

Целта на настоящия доклад е да обобщи препоръките на ЕС конкретно за професионалното образование и обучение, както и тези произхождащи от Стратегическата рамка за европейско сътрудничество в областта на образованието и обучението с оглед на европейското пространство за образование и отвъд него, да разгледа предизвикателствата, с които България се сблъсква по пътя на реализация на тези препоръки, както и дигитализацията на ниво професионално образование и обучение, нужна за достигане на по-добро качество и достъпност.

1. Препоръките за развитие на професионалното образование и обучение в България

„Препоръката на съвета от 24 ноември 2020 година относно професионалното образование и обучение (ПОО) за постигане на устойчива конкурентоспособност, социална справедливост и издръжливост“, EQAVET Framework и декларацията от Хернинг са основните стълбове, върху които се крепи дейността на НАПОО.

Препоръките за развитие на ПОО могат да бъдат обобщени в няколко направления:

1. Засилване на връзката между ПОО и бизнеса. Развитие на дуалната форма на обучение.
2. Повишаване на цифровите умения.
3. Подобряване качеството и привлекателността на системата.
4. Участие на възрастните в обучения през целия живот.
5. Учебна мобилност.
6. Реформиране на законодателството в сферата на ПОО.

На 12 септември 2025 г. в Дания страните членки на ЕС, включително България, подписват декларация, с която потвърждават ангажимента си да прилагат препоръката на съвета на ЕС от 2020 г. относно ПОО и определят цели на национално и европейско ниво относно постигането на отлично, привлекателно и приобщаващо ПОО. В **Таблица 1** са представени целите на национално ниво, които България се ангажира да постигне:

Таблица 1: Цели 2026-2030г. на национално ниво

Цел	Подцели
Повишаване на привлекателността на професионалното образование и обучение и равнопоставеността му с общото и академичното образование.	Подобряване на качеството на професионалните програми Популяризиране сред широката общественост на ползите от ПОО Мерки за борба с преждевременното напускане на ПОО
Привличане на повече учащи и повече компании.	Достъп до качествено, приобщаващо и ефективно кариерно ориентиране

¹ Европейска мрежа за информация относно професионалното образование и обучение (ПОО), създадена и координирана от CEDEFOP – Европейския център за развитие на професионалното обучение.

	Професионални програми в области от решаващо значение за зеления и цифровия преход и конкурентоспособността на ЕС
Гарантиране, че професионалното образование и обучение е съобразено с променящите се нужди на пазара на труда.	Партньорствата с други съответни участници на пазара на труда Ефективно използване на данни от анализи на уменията и системи за проследяване на завършилите
Осигуряване на подходящо ниво на ключови умения в началното професионално образование и обучение.	Образование за гражданство, предприемачество и устойчивост, за адаптивност през целия живот към променящите се нужди на пазара на труда и личностно развитие
Насърчаване на приобщаването за професионално образование и обучение.	Справяне с неравенството между половете и стереотипите в ПОО Равен достъп до професии за всички хора
Увеличаване на участието на всички възрастни, особено на нискоквалифицираните, в мерки за повишаване на квалификацията и преквалификация.	Продължаващо професионално образование и обучение (CVET) Стажове и предложения за висше професионално образование и обучение Премахване на бариерите пред участието, спомагащи за превръщането в реалност на правото на хората на качествено и приобщаващо образование, обучение и учене през целия живот
Справяне с недостига на учители и обучаващи в професионалното образование и обучение.	Като направят професията по-привлекателна, гарантирайки тяхната професионална автономност; Като продължат да инвестират в първоначална педагогическа квалификация и професионално развитие Подкрепящи условия на труд и среда, включително училищно ръководство
Повишаване на високите постижения в професионалното образование и обучение.	Разширяване на модела на Центровете за професионално върхови постижения (CoVEs) Улесняване на работата в мрежа
Доразвиване на програми за висше професионално образование и обучение.	Нива 5-8 по Европейската квалификационна рамка
Приобщаваща мобилност на учащите, учителите и обучаващите.	Мобилност в професионалното образование и обучение
Максимизиране на потенциала, надеждността и отговорното използване на изкуствения интелект.	Модели на големи езици, виртуална и разширена реалност Използване на инструменти, основани на изкуствен интелект, и имерсивни технологии
Осигуряване на достатъчни и устойчиви инвестиции в начално професионално образование и обучение и продължаващо професионално образование и обучение.	Мониторинг и оценка на национално ниво Подкрепа и по-добра координация на съответните финансови механизми на ЕС и национални, регионални или секторни фондове за обучение на служителите

Source: Herning Declaration on attractive and inclusive Vocational Education and Training for increased competitiveness and quality jobs 2026-2030

2. Предизвикателства пред професионалното образование и обучение в България.

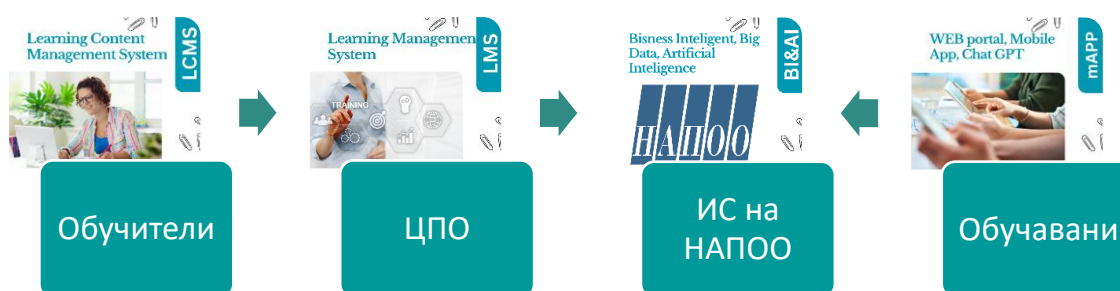
Рамковата политика на ЕС ясно подчертава необходимостта от използване на пълния потенциал на цифровите технологии в предоставянето на професионално образование и обучение. Това включва въвеждане на нови инструменти и подходи за цифрово

обучение и интегриране на цифровите технологии в педагогиката и дидактиката по начин, ориентиран към учащите, и в подкрепа на цифровото приобщаване. CEDEFOP¹ поддържа European VET policy dashboard, която представя визуално напредъка на страните при постигане на договорените количествени цели разделени в три ключови такива:

- ПОО за развитие на култура на учене през целия живот;
- ПОО за устойчив преход, устойчивост и високи постижения;
- ПОО за европейско образователно пространство.

На 3 ноември 2024 г. се провежда традиционната „ReferNet“ конференция „Европейски измерения на професионалното образование и обучение“, организирана от Националната агенция за професионално образование и обучение. В разгледаните ключови теми се включват: „дигиталните умения и значението им за трансформацията на пазара на труда“ и „изкуственият интелект в ПОО“.

Системата за ПОО в България би могла да се представи визуално със следната схема:



Source: Илюстрация на автора.

Фигура 8: Примерна схема на субектите в обучителния процес

На база елементите от представената схема могат да се систематизират и анализират предизвикателствата, които срещат отделните субекти, като се посочват възможните дигитални инструменти за справяне.

В началото на „веригата“ са поставени обучителите, които могат да се разглеждат и като своеобразен род създатели на цифрово съдържание. Предизвикателствата пред тях са много, тъй като учебното съдържание освен да отговори на потребностите на пазара на труда, трябва да бъде съобразено и с държавните образователни стандарти.

Предизвикателства пред обучителите:

- Липса на взаимопомощ и подкрепа при привеждане на отделните учебни планове (обучения) в съответствие с ДОС (държавни образователни стандарти) – създаване и поддръжка на унифицирана LCMS от страна на НАПОО, която да предлага иновативни инструменти като AI и big data, за създаване и анализ на учебно съдържание;

¹ European Centre for the Development of Vocational Training - Европейския център за развитие на професионалното обучение.

- Липса на умения и компетентности за създаване на цифрово съдържание – виртуални асистенти и помощници, за повишаване на дигиталната компетентност на обучаващите.

Предизвикателства пред центровете за професионално обучение:

- Голяма административна тежест и трудности при използването на ИС на НАПОО. Ограничени възможности на информационната система и необходимост от поддръжка на собствена LMS – създаване на повече формати и възможности за въвеждане на данни в ИС на НАПОО, с цел прехвърляне на данни и намаляване на дублирането. Използване на модерни xAPI;
- Трудности при създаването и управлението на „портали“ за обучение и мобилни приложения – използване на облачни решения и решения от типа SaaS;
- Предлагане на модерни услуги тип „изживяване“ – преминаване от LMS към LXP (LXS). Използване на инструменти като геймификация, виртуална и добавена реалност за повишаване качеството и привлекателността на предлаганите обучения;
- Липса на достатъчно обучители – използване на „виртуални“ обучители и виртуални класни стаи;
- Предлагане на индивидуално обучение съобразено с нуждите и възможностите на обучавания – изграждане на „виртуален образ“ на обучавания, анализ с помощта на изкуствен интелект и създаване на индивидуална програма (learning patch) и учебен план.

Предизвикателства пред Националната агенция за професионално образование и обучение:

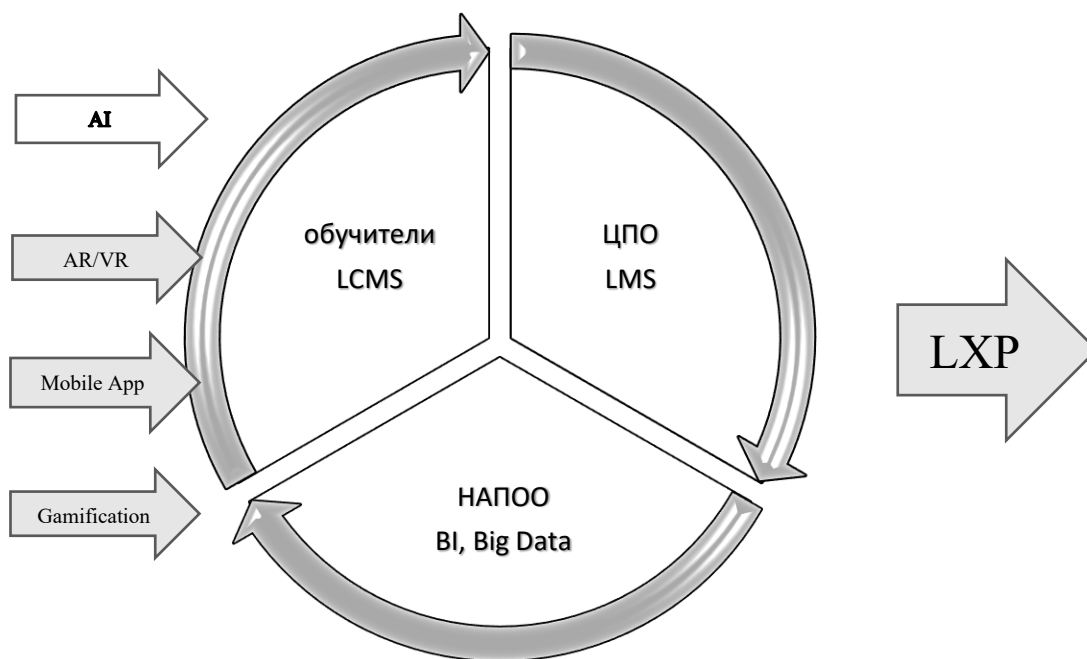
- Труден и бавен процес по създаване и приемане на нови държавни образователни стандарти – използване на дълбоки модели за анализ и изкуствен интелект, за генериране на стандарти;
- Липса на адекватна модернизация в информационната система на агенцията - внедряване на изкуствен интелект в платформата на НАПОО, подобряване на дизайна на системата използвайки концепцията за „user friendly design“, използване на интелигентна обработка на документи;
- Липса на адекватна мобилна поддръжка – създаване алтернативно на системата мобилно приложение, използване на прогресивни уеб приложения;
- Предоставяне на прекалено малко информация за центровете за професионално обучение и предлаганите обучения – създаване на модерна платформа от типа маркетплейс, в която ЦПО да предлагат своите обучения.

Предизвикателства пред обучаваните:

- Липса на достатъчно информация, необходима за избора на обучаваща институция (център за професионално обучение), както и конкретно професионално обучение – предлагане на качествена информация на сайта на НАПОО и други платформи, използване на повече дигитални канали, създаване на платформи тип „marketplace“;

- Липса на достатъчно гъвкавост и мобилност при голяма част от предлаганите обучения – създаване на микрообучения и e-learning, прогресивни уеб приложения (PWA), социално учене;
- Трудност при откриване на подходящи обучения – създаване на виртуален профил и използване на ИИ, за избор на обучения, пътеки за обучение (learning patch).

В сферата на професионалното образование и обучение е нужна една модерна екосистема, която да осигури облачни решения, мобилност и хибридность. На всички етапи от анализирания процес по предлагане на ПОО от учителя до обучавания е необходима дигитализация и използване на повече технологични новости, както в обучителния процес така и в администрацията и контрола. За да подобрят представянето си на пазара, да предложат конкурентни и стойностни услуги ЦПО се нуждаят от силни системи за e-learning и хибридни обучения. В много случаи те срещат конкуренция от обучаващи организации без лиценз към НАПОО, които избягвайки административната тежест предоставят много по-евтини и атрактивни услуги. Основните усилия от страна на ЦПО трябва да бъдат насочени към използване на LXS (Learning Experience System), познати като LXP (Learning Experience Platforms), които са наследник на класическите системи за управление на обучението, с тази разлика, че предоставят възможност за представяне на обучителния процес като цялостно изживяване за обучавания, което повишава привлекателността и качеството на ПОО.



Source: Илюстрация на автора

Фигура 9: Трансформация на обучението

Заклучение

Професионалното образование и обучение в България, следвайки препоръките и рамковите споразумения на ЕС, върви бавно по пътя на дигитализацията. Ключовото му значение за икономиката на страната като цяло в цифровото десетилетие изисква по-

бързи и адекватни мерки, както и модернизация на използваните дигитални технологии. За постигането на качествено, приобщаващо и достъпно ПОО е необходима трансформация от обучение ориентирано към процеса към обучение ориентирано към обучавания, което е напълно възможно с помощта на LXP – комбинация от LMS, VR/AR, геймификация и ИИ.

Литература:

1. CEDEFOP. (2025). European VET policy dashboard. Retrieved from <https://www.cedefop.europa.eu/en/tools/european-vet-policy-dashboard?year=2022#4>
2. Council of the European Union. (2025). *Herning Declaration on attractive and inclusive Vocational Education and Training for increased competitiveness and quality jobs 2026–2030*. (Herning, Denmark) Retrieved from <https://danish-presidency.consilium.europa.eu/en/news/new-declaration-to-strengthen-european-cooperation-on-vocational-education-and-training/>
3. Gorshenin, A. (2018). Toward modern educational IT-ecosystems: from learning management systems to digital platforms.
4. H, James., Pine II, Joseph., & B, Gilmore. (2011). *The Experience Economy*. Harvard Business School Press.
5. Nakamura, W. T., Teixeira de Oliveira, E. H., & Conte, T. (2017). Usability and User Experience Evaluation of Learning Management Systems. *ICEIS 2017 - 19th International Conference on Enterprise Information Systems*.
6. ReferNet. (2023). *Как да направим ПОО „зелено“?* Retrieved from <https://refernet.bg/events/kak-da-napravim/>
7. НАПОО. (2025). Годишни планове и доклади на НАПОО.
8. Съвет на Европейския съюз. (2020). Препоръка на съвета от 24 ноември 2020 година относно професионалното образование и обучение (ПОО) за постигане на устойчива конкурентоспособност, социална справедливост и издръжливост. *Официален вестник на Европейския съюз*, р. 417.
9. Съвет на Европейския съюз (2024). Препоръки на съвета от 23 ноември 2023 година относно основните фактори, способстващи успешното цифрово образование и обучение. *Официален вестник на Европейския съюз*.

ОПТИМИЗИРАНЕ НА АВТОМАТИЗИРАНО ТЕСТВАНЕ НА БИЗНЕС ЛОГИКА В .NET ЧРЕЗ AI BUSINESS LOGIC UNIT TESTING OPTIMIZATION IN .NET USING AI

Maria Marzovanova¹,

Veska Mihova²

e-mail: mmarzovanova@unwe.bg¹,

e-mail: vmihova@unwe.bg²

Абстракт

Автоматизираните тестове за бизнес логика са от критична важност при разработката на бизнес информационни системи. В съвременната среда с постоянно променящи се и нарастващи изисквания, автоматизираното тестване е един от инструментите за гарантиране на качество и стабилност на софтуерните решения. Традиционният начин на създаване на тестове е бавен и податлив на човешки грешки, но с развитието на изкуствения интелект (AI) пред разработчиците и QA специалистите се отварят множество възможности за оптимизация. Този доклад представя едно изследване на ползите от AI като помощно средство за генериране на тестови случаи базирани на анализ на програмния код, генериране на тестове, подбор и приоритизиране на тестове с цел олекотяване на целия процес в .NET приложения.

Abstract

Unit testing the business logic is crucial in the development of business information systems. In a dynamic and increasingly demanding environment, unit testing is one of the tools for assuring quality and stability of software solutions. The traditional way of adding unit tests to a business logic layer is a slow and prone to errors process, but with the evolving artificial intelligence assistance, software developers and QA specialists have more optimization opportunities. This paper presents research results on the benefits of using AI assistance in generating test cases, based on code analysis, generating tests, test selection and prioritization, improving the efficiency and effectiveness of the testing process in .NET application.

Ключови думи: Unit Testing optimization, AI in Unit testing, Business logic.

JEL: C88, L86

Увод

Бизнес информационните системи са пряко зависими от правилността на бизнес логиката заложена в тях. Тестовите написани, за да упражняват слоя с бизнес логиката, са ключови за ранно откриване на регресия в кода и подsigуряване на очакваното поведение на системата. В същото време, писането и поддържането на тестове често се пренебрегва, защото е трудоемко и изисква дълбоки познания в областта на приложение на информационната система, както и в нейната имплементация. Последните версии на инструментите за разработчици с AI асистенция предлагат различни възможности за автоматично генериране на тестови случаи, препоръчват входни данни за тях и приоритизират тестовите случаи с цел по-ефективно и фокусирано тестване. Настоящия

¹ Гл. ас. д-р в катедра „Информационни технологии и комуникации“, УНСС

² Гл. ас. д-р в катедра „Информационни технологии и комуникации“, УНСС

доклад разглежда използването на AI помощ в целия процес на тестването при приложения разработени в .NET.

Съвременната облачна среда, масивите от данни и усложнените зависимости изискват нов подход за подпомагане на QA, тестове и архитектура [5].

Възприемането на автоматизирани и адаптивни подходи при разработването на информационни системи е в съзвучие с тенденциите за системи с висока степен на конфигурируемост и персонализация [6]. Тези принципи могат успешно да бъдат приложени и при внедряването на AI асистенция в процеса на тестване на бизнес логика.

Тестване на кода

По дефиниция от Microsoft [1] тестването е чудесен начин да се подсигури, че едно приложение прави точно това, което неговия автор е целял. В статията се разглеждат няколко вида тестване – единични тестове (Unit tests), интеграционни тестове (Integration tests) и тестове за натоварване (Load tests). Единичните тестове се упражняват върху изолирани единици от кода и целят да подсигурят, че всяка единица от логиката работи както е по дизайн и допринася за цялостната надеждност на системата. Другите видове тестове са също от изключителна важност, но настоящото изследване се фокусира върху единичните тестове и в изложението при споменаване на тестове и тестване се има предвид единично тестване.

За .NET се предлагат различни среди за създаване на тестове, библиотеки с различни възможности и улеснения за верифициране на очакваните резултати от изпълнение на единиците код. Сред тях са NUnit, xUnit и MSTest. NUnit [2] е вече зряла платформа с широк диапазон на опции за проверка. xUnit [3] е по-нова платформа, с повече възможности за разширяване и модуларност. MSTest [4] е вградена Microsoft, която лесно се интегрира с Visual Studio.

Въпреки че тези платформи предоставят солидна основа за тестване в .NET, традиционните подходи към създаването и поддържането на тестове често се сблъскват с предизвикателства. Ръчното писане на тестови случаи е трудоемко, податливо на човешка грешка и може да не обхване всички възможни сценарии или гранични случаи.

За да преодолеем тези ограничения и да постигнем по-ефективен и всеобхватен процес на тестване, ще представим адаптивен подход, базиран на използването на инструменти с изкуствен интелект. Изкуственият интелект е вече в голяма помощ на разработчиците на софтуер, безспорно е неговата полза и при тестването. Предстои да разгледаме подход с ясни и последователни стъпки, следването на които ще доведе до намаляване на времето и възможните грешки при създаването на тестове в .NET среда.

Подход за тестване с AI асистенция

Когато разработчикът избере да използва инструменти с AI възможности при създаването на тестове, добра практика е да се премине през следните стъпки:

- Анализ на кода
- Генериране на тестове
- Подбор и приоритизиране на тестове

- Изпълнение на тестовете и анализ на покритието



Фигура 1: Стъпки в процеса на тестване с AI асистенция

Анализ на кода

Това е процес на семантичен анализ на кода с цел да се определи какво трябва да се тества, как да се изолират отделни единици за тестване, какви входни данни, изходни данни, странични ефекти и зависимости са от значение.

Използват се техники като AST (абстрактно семантично дърво) за извличане на структурата и семантиката на кода. Control flow анализ за идентифициране на възможните пътища при изпълнение на кода. В резултат на това ще бъдат определени единици от бизнес логиката, които е необходимо да бъдат обхванати от тестовете, както и зависимости в кода от свързани методи, външни API.

Един добър инструмент за извършване на анализ на кода е Roslyn, .NET платформа, която предлага API за анализиране на C# код, който поддържа и последните версии на .NET за разлика от други инструменти.

Генериране на тестове

Въз основа на извлечените от анализа на кода метаданни, фазата на генериране на тестове цели автоматичното създаване на цялостен пакет от unit тестове. В зависимост от избрания инструмент за генериране се прилагат различни комбинации от AI техники, включително търсене-базирано тестване и разпознаване на модели, за генериране на разнообразни тестови случаи, насочени към различни аспекти от функционалността на кода. Например, използваме генетични алгоритми за еволюиране на тестове, които максимизират покритието на кода и разкриват потенциални грешки. Генерираните тестове са проектирани да бъдат съвместими с популярни .NET платформи за тестване, като xUnit и NUnit.

В тази фаза може да се използва и GitHub Copilot, за който трябва да се добавят метаданните в кода като коментари. Трябва да се обърне внимание тези подсказки да бъдат ясни и точни, разбираеми за асистента. Този инструмент е много полезен за бързо създаване на скелет на тестове, но изисква внимателен преглед и проверка.

Избор на тестове и приоритизиране

За да се оптимизира процеса на тестване и да се гарантира, че най-критичните тестове се изпълняват първи, се използват AI-базирани техники за селекция и приоритизация на тестовете. Алгоритми за машинно обучение анализират данни за исторически изпълнения на тестове, промени в кода и метрики за сложност на кода, за да предвидят кои тестове е най-вероятно да се провалят или да разкрият нови дефекти. Това ни позволява да приоритизираме тестовете на база на техния риск и потенциално въздействие, позволявайки на програмистите да се фокусират върху най-важните области от кодовата база.

Изпълнение на тестовете

Генерираните и приоритизирани тестове след това се интегрират безпроблемно в .NET средата за разработка. Използват се стандартни .NET платформи за тестване (xUnit, NUnit, MSTest) за изпълнение на тестовете и преглед на резултатите. В добавка може да се разработи собствен инструмент за стартиране на тестове, на който да се укаже приоритизираният ред, гарантирайки че най-критичните тестове се изпълняват първи. Резултатите от тестовете се анализират, за да се идентифицират дефекти и да се предостави обратна връзка на програмистите.

Оценка

За целите на демонстрацията и оценката на описания подход е изготвен един примерен C# клас с четири метода за работа с текст. Самите методи нямат практическо приложение или претенции да изпълняват сложна бизнес логика, създадени са само за целите на демонстрация на подхода. Върху класа се упражнени стъпките в представената последователност. Като краен резултат ще използваме информацията от Visual Studio представяща нивото на покритие на кода с тестове.

Четири метода в клас са както следва:

- Метод за анализ на чувствата в текст (AnalyzeSentiment), който приема като входни данни текст и представя на изхода си в текстов резултат оценка – позитивен, негативен или неутрален.
- Метод за превод на текст (TranslateText), който приема като входни данни текст и целеви език и в резултат връща преведения текст.
- Метод за генериране на изречение с подадена дума (GenerateText), който приема като входни данни дума и представя на изхода си изречение, в което се съдържа тази дума.
- Метод за съкращаване на текст (SummarizeText), който приема като входни данни текст и в резултат връща по-кратък текст.

Първата стъпка е анализ на кода, избран е Roslyn като инструмент за анализ на кода. За използването му е създадено приложение, което достъпва директно целевия клас и извършва семантичен анализ.

От извършения анализ получаваме резултати като представения на фигура 2. В тях се съдържа информация за достъпността на метода, входни параметри, тип на връщания резултат, списък с методи, които се извикват, дали включва try catch конструкции, дали има изрично хвърляне на грешки, както и предложения за ключови тестови случаи.

Този анализ би помогнал и при ръчно създаване на тестове, тъй като съдържа полезна информация, която да насочва разработчикът и да подsigури, че няма да бъдат пропуснати важни елементи от логиката и зависимостите в тествания метод.

```
{
  "Name": "AIGeneratedUnitTestsExample.AICapabilities.AnalyzeSentiment(string)",
  "Accessibility": "Public",
  "ReturnType": "string",
  "Parameters": [
    {
      "Name": "text",
      "Type": "string",
      "IsNullable": false
    }
  ],
  "InvokedMethods": [
    {
      "Name": "string.IsNullOrEmpty(string?)",
      "ContainingType": "string"
    },
    {
      "Name": "string.Split(char, System.StringSplitOptions)",
      "ContainingType": "string"
    },
    {
      "Name": "string.ToLower()",
      "ContainingType": "string"
    },
    {
      "Name": "string.ToLower()",
      "ContainingType": "string"
    }
  ],
  "BranchCount": 5,
  "Throws": false,
  "HasTryCatch": false,
  "SuggestedEdgeCases": [
    "text == null",
    "text is empty (if collection/string)",
    "exercise true/false branches"
  ]
},
```

Фигура 2: Резултати от семантичен анализ с Roslyn

Резултатите записваме в json файл, за използването им в следващата стъпка.

За генерирането на самите тестове използваме асистенцията на GitHubCopilot. Един от възможните начини е да се добави съкратена информация от резултатите от анализа като коментар на всеки метод. След това да се подаде искане към GitHub Copilot да генерира тестове съобразявайки се с тази информация. В примерния клас това не би било затруднение, защото съдържа само няколко метода и не би било толкова времеемко. С оглед на това, че целите на изследването са да утвърди подход приложим и при големи информационни системи, избираме да съхраним информацията като помощен файл добавен в проекта и да насочим AI асистента към този файл при генерирането на тестовете.

Подаваме следното искане в разговора с асистента: *I have a RoslynResults.json file in the project, can you generate unit tests for the project considering those Roslyn analyzer results?*

В резултат получаваме неговите предложения за тестове, общо 14, включващи насоките дадени от анализа.

За показания на фигура 2 метод с резултатите от неговия анализ са генерирани 4 теста включващи ситуацияите:

- Подаден е празен текст, в който няма позитивни и негативни думи. Резултатът се очаква да бъде „неутрален“.
- Подаден е текст с повече позитивни думи, отколкото негативни. Резултатът се очаква да бъде „позитивен“.
- Подаден е текст с повече негативни думи, отколкото позитивни. Резултатът се очаква да бъде „негативен“.
- Подаден е текст с равен брой позитивни и негативни думи. Резултатът се очаква да бъде „неутрален“.

```
0 references
public class AICapabilitiesTests
{
    [Fact]
    0 references
    public void AnalyzeSentiment_ReturnsNeutral_ForNullOrEmpty() ...

    [Fact]
    0 references
    public void AnalyzeSentiment_ReturnsPositive_WhenMorePositiveWords() ...

    [Fact]
    0 references
    public void AnalyzeSentiment_ReturnsNegative_WhenMoreNegativeWords() ...

    [Fact]
    0 references
    public void AnalyzeSentiment_ReturnsNeutral_OnTie() ...
}
```

Фигура 3: Генерирани тестове с помощта на GitHub Copilot

Следващата стъпка е селекция и приоритизиране на тестове. Проучването показва, че няма отделен или интегриран инструмент с AI възможности съвместим с Visual Studio, или поне не безплатен такъв. Но GitHub Copilot се оказва полезен и за тази стъпка.

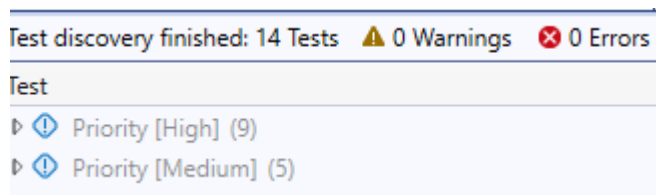
При подадено следното искане : *Can you perform a selection and prioritization of the generated tests and add them in lists?*, са предложени модификации на генерираните по-рано тестове, които отразяват техния приоритет според AI асистента. Добавен е Trait, с име Priority, като са използвани стойностите – висок, среден, нисък.

Според добавеното описание в модификациите, при избора на стойност е използвана следната логика :

- **Висок** приоритет е даден на тестове, които засягат ключово за изпълнението на методите поведение или са базирани на крайни случаи,
- **Среден** приоритет имат тестове за поведение на важни, но не критични, елементи или такива свързани с форматиране например.
- **Нисък** приоритет е даден на тестове свързани със ситуации, които е малко вероятно да се случат в реалността.

И остава да изпълним последната стъпка, да отразим избора и приоритетите при изпълнение на тестовете. Във Visual Studio разполагаме с Test Explorer, в който се извежда пълен списък на всички тестове в системата. Разполагаме с удобно меню за

групиране, с което може да групираме тестовете по Traits. Тъй като имаме само един и той е приоритетът на тестовете, както е показано на фигура 4, лесно се виждат групите и могат да се изпълняват самостоятелно по групи при нужда.



Фигура 4: Групиране на тестове по приоритет в Test Explorer, Visual Studio

При добавяне на повече видове Traits, този подход може и да се усложни, но за такива ситуации може да се изготвят списъци в същия инструмент. Изпълнението им след това е аналогично.

Заклучение

Изследването и демонстративният пример показват обещаващи възможности за бъдещо развитие на AI-базирани решения за тестване в .NET среди. Въпреки че нашето изследване се фокусира върху специфични техники и случаи на употреба, вярваме, че принципите и резултатите могат да бъдат приложени и в други контексти. Бъдещите изследвания ще се фокусират върху по-нататъшното автоматизиране на процеса на тестване, подобряване на точността на приоритизирането на тестовете и разширяване на обхвата на техниките за анализ на кода. Интегрирането на изкуствен интелект в практиките за тестване представлява значителна стъпка към създаването на по-надежден и висококачествен софтуер.

References

1. “Testing in .NET,” Microsoft Learn, 22 Oct. 2025. [Online]. Available: <https://learn.microsoft.com/en-us/dotnet/core/testing/>
2. “What Is NUnit?,” NUnit.org, 2024. [Online]. Available: <https://nunit.org/>
3. “About xUnit.net,” xUnit.net. [Online]. Available: <https://xunit.net/?tabs=cs>
4. “Unit testing C# with MSTest and .NET,” Microsoft Learn, 28 Oct. 2025. [Online]. Available: <https://learn.microsoft.com/en-us/dotnet/core/testing/unit-testing-csharp-with-mstest>
5. G. Stefanov, “Езера от данни (Data Lakes) като хранилище за Големи данни. Анализ на AWS подхода,” *Научната конференция „Иновативни информационни технологии за дигитализация на икономиката“ (Innovative Information Technologies for Economy Digitalization – IITED)*, УНСС, Sofia, 2023.
6. Milev, P. (2025). Design Solutions for an Information System with User Configuration in a University Environment. *Ikonomiceski i Sotsialni Alternativi*, pp. 110-115.

ТИПОВЕ ДАННИ В УПРАВЛЕНИЕТО НА БИЗНЕС ПРОЦЕСИ (BPM) – КЛАСИФИКАЦИЯ, СРАВНЕНИЕ И АНАЛИТИЧЕН ПОТЕНЦИАЛ

Data Types in Business Process Management: classification, comparison, and analytical potential

Гено Стефанов¹,
e-mail: genostefanov@unwe.bg¹

Abstract

Business Process Management (BPM) generates and uses a variety of data types that play a key role in modeling, analyzing, optimizing, and monitoring business processes. This report provides a systematic classification and comparison of the main categories of process data used in BPM - process models (BPMN), process execution data (log data), and contextual business data. Their characteristics, sources, purposes, and applicability in various analytical scenarios are analyzed - from process mining to predictive analytics and business process optimization. A comparative summary of the advantages and limitations of each data type is presented, as well as their importance for generating knowledge and making management decisions. Finally, the main challenges related to their integration, quality, and joint use are outlined.

Key words: BPM, Process Mining, Predictive Analytics.

JEL: O30, O31

Introduction

In today's digital economy, Business Process Management (BPM) plays a crucial role in organizational sustainability, efficiency, and innovation[1]. A core feature of BPM is its data-centered nature — processes not only generate, but also extensively use various types of data, which form the foundation for modeling, monitoring, analysis, and continuous improvement. The quality of data and its preparation, including cleaning, transformation, and integration - is a critical prerequisite for the effectiveness of analytics in business process contexts[2]. In [3] the authors are emphasizing that data intelligence significantly influences how content is structured and presented in digital environments, which is increasingly relevant in the context of process - aware information systems. In [4] the authors are stating that digitalization has met certain expectations related to automation, efficiency, and accessibility, yet it has simultaneously exposed critical limitations - particularly concerning the integration of heterogeneous data and the realization of meaningful, data - driven business transformation, which directly impact how we use the process data.

Process data in BPM encompass different categories, each contributing uniquely to understanding and optimizing processes. This report provides a systematic overview of the three main categories of data in BPM:

- **Process Models** – created through notations such as BPMN, they formally describe the sequence, logic, and participants in a given process;
- **Process Execution Data (Logs)** – dynamic data recording the actual execution of processes within information systems such as ERP or BPM systems;
- **Contextual Business Data** – additional business data reflecting external and internal factors that influence process execution and analysis.

¹ Главен Асистент, доктор. Катедра Информационни технологии и комуникации, УНСС.

Each of these data types has a different structure, purpose, analytical potential, and set of challenges. Their integration and joint use are strategically important for knowledge extraction, precise analysis, and informed decision-making.

The goal of this report is to classify, compare, and analyze these three categories of process data, emphasizing their characteristics, applications, and role in modern analytical practices such as Process Mining, Predictive Analytics, and Business Process Optimization. Furthermore, the paper identifies challenges related to data quality, integration, and usability, which lay the foundation for future research and improvement in BPM.

Data categories in BPM

Process Models (BPMN)

Process models are formal and graphical descriptions of business activities, their sequence, participants, rules, and interrelations. They serve as a basic representation for understanding, analyzing, automating, and optimizing business processes. The most common notation is the Business Process Model and Notation (BPMN) — an international standard maintained by the Object Management Group (OMG)[1].

The main characteristics of this data type are listed in Table 1.

Table 2. Main characteristics of process model data

Characteristic	Description
Data Type	Structured (XML/JSON)
Format/Notation	BPMN, EPC (Event-driven Process Chains), UML Activity Diagrams
Components	Activities, events, branches, flows, roles, objects
Storage/Sources	BPMS (Business Process Management Systems), modeling tools (Camunda, Signavio, Bizagi)
Purpose	Modeling, documentation, automation, simulation
Level of Abstraction	High – describes what should happen, not necessarily how

Although process models primarily serve visualization and communication functions through graphical notations, their underlying structured digital representation — typically in XML or JSON — enables automated processing, analysis, and comparison with other process data sources.

Process Execution Data (Logs)

Process execution data, often called event logs or audit trails, are chronological records of actions, events, and states occurring during the actual execution of business processes within an information system.

They are automatically generated by various software systems such as BPMS, ERP (e.g., SAP, Oracle EBS), CRM, IoT platforms, and other applications that record interactions and transactions.

The main characteristics of this data type are listed in Table 2.

Table 2. Main characteristics of log data

Characteristic	Description
Data Type	Structured (tabular, JSON, CSV, XES – eXtensible Event Stream)
Granularity	High – records actions at the event level
Sources	BPMS, ERP, CRM, IoT devices, Web systems
Purpose	Execution analysis, Process Mining, monitoring, auditing
Issues	Different formats, lack of context

Logs contain structured information, most often using the following format (table 3):

Table 3. Example log data format

Поле	Пример
Case ID	INV 00231
Activity	Approve Invoice
Timestamp	10/29/2025 15:42
Resource	John Smith
Outcome	Approved
Context (opt.)	Amount > 1000, Country=BG

Logs are the main data source for Process Mining techniques, which include:

- Discovery – automatic extraction of process models from data;
- Conformance checking – comparing models with real executions;
- Enhancement – improving existing models with real-world data.

Contextual Business Data

Contextual business data enrich the understanding and analysis of business processes. Unlike process models and logs, which directly relate to process structure and execution, contextual data provide an “external perspective,” reflecting the environment in which processes occur.

They include both internal organizational factors (e.g., customer profiles, product categories, employee roles) and external influences (e.g., market conditions, seasonality, regulatory changes, macroeconomic indicators). The addition of contextual data significantly increases the analytical value of process analysis, especially in predictive and personalized scenarios.

The main characteristics of this data type are listed in Table 4.

Table 4. Main characteristics of contextual business data

Characteristic	Description
Data Type	Structured, semi-structured (JSON), unstructured (text, documents)
Sources	CRM systems, BI platforms, external databases, market research, open data
Purpose	Enriching analysis, forecasting, supporting decision-making
Issues	Heterogeneity, difficult integration, lack of standards, dynamic updates

The full potential of contextual data is realized when semantically linked to process events — for example, connecting a customer ID from a log to a CRM profile. This requires the use of metadata management, ontologies, and master data integration techniques.

Analytical potential and applications of data in BPM

Process Mining: From Execution Data to Real Insights

Process Mining is a data-driven approach that leverages event logs to:

- discover how processes are actually executed (discovery),
- compare real execution with predefined models (conformance checking),
- enhance existing models with insights from execution data (enhancement).

It is one of the most effective methods for identifying bottlenecks, deviations, inefficiencies, and opportunities for improvement[6].

Predictive Analytics: Anticipating Process Outcomes

By combining historical event logs with contextual business data, organizations can build predictive models that:

- estimate process completion time,
- forecast the likelihood of delays or failures,
- assess risk levels in ongoing process instances.

Techniques include machine learning algorithms such as Decision Trees, Random Forests, and Gradient Boosting, often enhanced by contextual features like customer segment, product type, or workload[10].

Prescriptive Analytics & Decision Support: Optimization through Insight

The highest level of analytics in BPM involves integrating all three data types to support[7]:

- scenario analysis and simulation (what if scenarios),
- process optimization (resource allocation, cost/time reduction),
- strategic decision-making via Decision Support Systems (DSS).

Figure 1. illustrates the intersection of three key data sources in Business Process Management: process models, execution logs, and contextual data. Each category contributes distinct value - process models provide structural definitions (e.g., BPMN), logs capture real execution behavior (e.g., XES), and contextual data enriches analysis with business relevance. At their intersection lie powerful capabilities such as process mining, predictive analytics, and simulation, ultimately enabling integrated BPM intelligence. This integration supports more informed decision - making, continuous improvement, and strategic alignment of business processes.

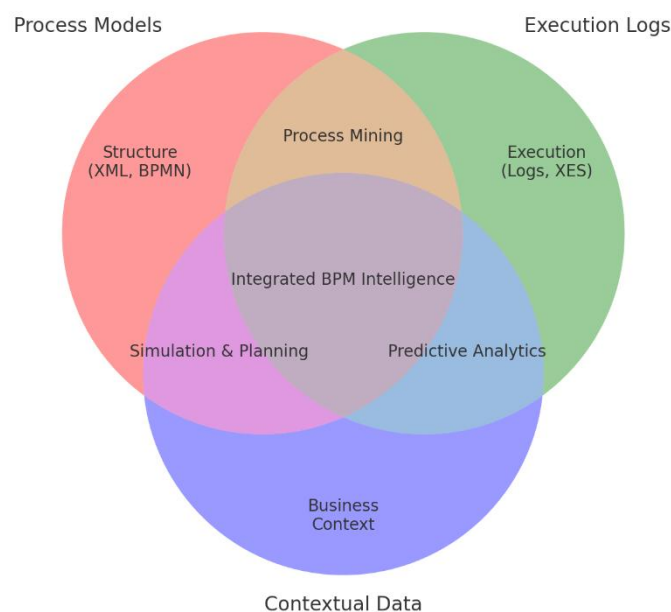


Figure 1: Intersection of BPM Data Types and Analytical use

Conclusion

The three types of BPM data — models, logs, and contextual data — represent complementary sources of knowledge. While each type has its own analytical potential, their integration and combined use yield the highest value for organizations. Major challenges remain: ensuring data quality, overcoming heterogeneity, and building compatible integration platforms. Future research should focus on developing integration frameworks, metadata standards, and AI-based methods for automatic linking between models and real data.

References

1. Belev, I., (2021). Comparing Business Process Management and Case Management, ICAICTSEE–2020, UNWE, Sofia, Bulgaria
2. Miteva, Genka. Murdzheva, A. (2023) Data Preparation Techniques and Platforms in the Context of Machine Learning <https://www.unwe.bg/doi/iited/2023/IITED.2023.21.pdf>
3. Велкова, Ивона. (2024) Намаляване на рекламните измами в дигиталния маркетинг с помощта на изкуствен интелект <https://www.unwe.bg/doi/iited/2023/IITED.2023.11.pdf>
4. Милев, П., Табов, Я. (2023) Влияние на интелигентността на данните върху структурирането на съдържание в уеб публикации. Икономически и социални алтернативи, УНСС, бр. 4, 2023, с.73-85
5. Боянов, Л., Мурджева, А. (2024), Сбъднатите и несбъднатите очаквания от дигитализацията, <https://www.unwe.bg/doi/iited/2024/IITED.2024.15.pdf>
6. Comuzzi, M. (2024). Process Mining: From Descriptive to Predictive and Prescriptive Analysis. In Handbook on BPM.
7. Hasan, E. (2025). Big Data-Driven Business Process Optimization. TechRxiv.
8. Kokala, A. (2024). Leveraging Process Mining to Enhance BPM.
9. Sakr, S. et al. (2018). Business Process Analytics and Big Data Systems: A Roadmap. IEEE Access.
10. Márquez-Chamorro, A. et al. (2017). Predictive Monitoring of Business Processes: A Survey. IEEE Transactions.
11. Mihova, V., Murdjeva, A. (2014). Dynamic administration of database productivity. In Conferences of the department Informatics (No. 1, pp. 430-437).
12. Mihova, V., Marzovanova M. (2023) Modern Techniques and Tools for Generating SQL Queries, 13TH INTERNATIONAL CONFERENCE ON APPLICATION OF INFORMATION AND COMMUNICATION TECHNOLOGY AND STATISTICS IN ECONOMY AND EDUCATION (ICAICTSEE – 2023), UNWE, SOFIA, BULGARIA, <https://icaictsee.unwe.bg/past-conferences/ICAICTSEE-2023.pdf>

ARTIFICIAL INTELLIGENCE AND TAX ADMINISTRATION IN BULGARIA

Georgi Emilov Hristov¹

e-mail: georgi.hristov@unwe.bg

Abstract

This report explores the transformative potential of Artificial Intelligence (AI) technologies in tax administration in Bulgaria, with a focus on bridging the gap between existing tax fraud detection tools and the incorporation of business characteristics, such as abnormal financial status. It underscores the importance of integrating AI technologies to harmonize both structural and financial characteristics in tax fraud detection. The report presents a critical overview of existing tax fraud detection tools, deriving key insights on how such technologies can be improved.

Ключови думи: (до 5 думи). Format: TNR 11 pt.

JEL: H26, C55, C81

Introduction

Artificial Intelligence (AI) has emerged as a transformative force across various sectors, revolutionizing the way businesses and governments operate. In the context of tax administration, AI has the potential to streamline processes, enhance efficiency, and improve compliance, thereby contributing to the economic growth and stability of nations. In recent years, the Bulgarian tax administration bodies, predominantly the National Revenue Agency (NRA), have recognized the significance of digitalization, but the implementation of AI technologies remains a complex process which is still due.

The integration of AI into tax administration not only promises to increase revenue collection but also to simplify complex taxation processes, minimize errors, and foster a climate of transparency and trust between the government and its citizens.

This report aims to provide an overview of how different AI technologies are implemented in different tax administration bodies around the world. It explores the difficulties of such implementation and proposes a tool for enhancing tax fraud detection.

Research gap

This report serves as a critical bridge over an existing gap in the current landscape of tax administration and fraud detection. Presently, the tools and methodologies employed for detecting tax fraud predominantly focus on structural characteristics, such as the organization's structure or transactional patterns. While these aspects are undoubtedly essential for identifying potential irregularities, they often fall short in capturing a more holistic view of tax fraud. This limitation becomes glaringly evident when one considers the inherent dynamism and adaptability of fraudulent activities in the modern economic landscape.

The gap that this report addresses centers around the need to incorporate the business characteristics of tax fraud, specifically factors like abnormal financial status, into the detection process. Traditional methods may not sufficiently recognize signs of tax evasion or fraud that are ***rooted in the financial*** health or behavior of an entity. Such behaviors may include abnormal fluctuations in revenue, expense patterns, or financial ratios that deviate from industry norms and benchmarks.

By delving into this critical gap, this report seeks to shed light on how Artificial Intelligence can be harnessed to bridge the divide between structural and business characteristics in tax fraud detection. AI technologies have the capacity to analyze vast and complex financial datasets, identifying anomalies

¹ Докторант, катедра ИТК, УНСС

and patterns that human auditors might overlook. By taking into account abnormal financial status alongside structural information, AI-powered solutions have the potential to revolutionize tax administration, providing a more comprehensive and accurate assessment of potential fraud.

Literature review

The use of AI in tax administration can have significant benefits. It can improve taxpayer compliance by incorporating information and communication technology (ICT) into the tax administration system (Djafri et al., 2023). The application of e-tax systems, AI, chatbots, and biometric identification can streamline processes such as registration, filing, payment, and identification, making it easier for taxpayers to comply with their obligations (Djafri et al., 2023). Additionally, AI can be used to monitor entrepreneurial activity and identify cases of tax evasion, both by regulatory authorities and by entrepreneurs themselves.

Furthermore, AI can contribute to macroeconomic stability and economic growth. Tax administration reforms, including the integration of AI, can address challenges such as low tax revenue and corruption in the tax bureaucracy (Herbert et al., 2018). By improving tax policy and administration, AI can help generate higher tax revenue and reduce opportunities for corruption, leading to economic stability and growth.

Artificial intelligence (AI) has the potential to revolutionize tax administration by improving efficiency, enhancing tax collection and management, and reducing tax evasion. Several studies have explored the application of AI in tax administration and its impact on various aspects of the tax system.

Zhang (2023) suggests that the construction of an AI-based environmental protection tax system can improve tax collection and management, tax payment service, and tax management (Zhang, 2023). By leveraging AI technology, tax authorities can optimize the design of the tax system and enhance the efficiency of tax administration processes. This can lead to more effective tax collection and a reduction in tax evasion.

Another study highlights the digitalization of tax administration communication as a result of AI technology (Ihnatišinová, 2021). AI creates new digital communication channels, enabling more efficient and paperless tax administration. This can streamline communication between taxpayers and tax authorities, leading to faster and more accurate processing of tax-related information.

Furthermore, the integration of AI technology and tax risk management can establish an intelligent tax system (Huang & Zhang, 2022). By using AI algorithms such as Bagging and Support Vector Machines (SVM), tax authorities can analyze large volumes of data to identify potential tax risks and anomalies. This can help in detecting tax evasion and improving tax compliance.

Jiang (2022) presents an intelligent tax planning platform based on the current situation analysis, and proposes a new tax planning algorithm based on the advanced technology of big data and artificial intelligence, which helps enterprises to carry out intelligent tax planning under the premise of legal compliance, reduce company labor costs, and increase the freelancers' revenue.

However, there are certain challenges that need to be addressed when implementing such technologies. One such challenge is the need for administrative capacity and accountability. Studies have shown that countries like Bulgaria and Romania face obstacles in terms of administrative capacity and accountability, which can hinder the successful implementation of reforms, including those related to AI (Wegener et al., 2011). Therefore, it is crucial for Bulgaria to address these challenges and build a strong administrative foundation to support the integration of AI in tax administration.

To fully harness the potential of AI in tax administration, Bulgaria should consider the experiences and best practices of other countries. In the following section, we provide several examples of successful implementation of AI technologies in tax administration processes.

Examples of implementation of AI in tax administration

The government of Malta has implemented a system which alerts the authorities when a tax payer's income does not correspond to the value of assets in his possession (Zammit, 2023). While a good example of combating tax fraud and money laundering, this case poses the question of personal rights and security.

The Canadian government has announced TaxGPT (<https://taxgpt.ca/>), a chatbot which assists taxpayers in the process of understanding and filing their income returns. Of course, relying solely on a chatbot to do your taxes is never a good option, but it is still a great support tool.

The Netherlands also implemented a self-learning algorithm which categorizes citizens into risk groups in effort to stop childcare benefits fraud. It is also a good example of how terribly bad such implementation can go if not done right. Thousands of Dutch families have been wrongfully categorized and punished over the years which led to a need for revision of the tool. (Goujard and Manancourt, 2022). Such system has been successfully implemented in Italy.

Australia claims to have identified over \$530 million in unpaid tax bills and prevent \$2.5 billion in fraudulent claims using AI models, including deep learning and natural language models.

China has implemented reforms in tax collection and administration, using AI and blockchain technology to reshape its core processes (Wang & Chen, 2018). Similarly, Russia has utilized AI in the tax sphere to identify cases of tax evasion and ensure compliance.

By studying and adapting the successful approaches, Bulgaria can develop a comprehensive strategy for integrating AI into its tax administration system.

Challenges in implementing AI in the Bulgarian tax administration

There are several non-neglectable challenges in implementing AI technologies in the tax administration process in Bulgaria.

Wegener et al., 2011 presents the need for administrative capacity and accountability so that reforms can be pushed through successfully. As he goes to show, Bulgaria and Romania did not have that capacity then. The country's political and socioeconomic transformations are still unfinished, which hinders the adequate completion of the Europeanization process (Wegener et al., 2011). These challenges can impact the implementation of AI in tax administration, as they require a strong administrative capacity and a robust accountability framework.

Another issue is lack of infrastructure and technology. Upgrading and modernizing existing IT infrastructure to support AI implementation can be costly and time-consuming. This includes procuring powerful hardware, software, and networking resources. Moreover, integrating the complex AI systems with the already in-place legacy systems requires skillful professionals in the field. This leads to the next issues – skill gap.

Building a team with the necessary AI expertise can be challenging. Recruiting and training data scientists, machine learning engineers, and AI specialists may require significant effort. In addition to that, they also have to be skilled in tax, audit, and accounting which makes the task even harder.

As we saw in the aforementioned examples, ethical considerations are always an issue which has to be taken care of. Tax administration AI systems must be designed with ethical considerations in mind, avoiding bias in algorithms and ensuring fairness in decision-making are crucial aspects for people to trust the technology. Building and maintaining public trust in AI-powered tax administration is critical. Ensuring that the public understands how AI is being used and its benefits is important for acceptance.

To successfully implement AI in the Bulgarian tax administration, a comprehensive strategy that addresses these challenges is essential. This strategy should involve collaboration among stakeholders, including government agencies, IT experts, data scientists, and the public, to ensure that AI systems are effective, transparent, and trusted.

Conclusion

The implementation of AI in tax administration in Bulgaria holds immense potential for transforming the efficiency, effectiveness, and transparency of the country's tax system. This report has explored the various facets of integrating AI into tax administration, highlighting both the opportunities and challenges associated with this endeavor.

Bulgaria, like many other countries, faces significant challenges in modernizing its tax administration system. The complexities of tax laws, the need for efficient revenue collection, and the demand for improved taxpayer services all converge to create a pressing need for innovative solutions. AI technologies, including machine learning, data analytics, and automation, offer a promising path forward.

As Bulgaria progresses in its journey toward AI-driven tax administration, it is imperative that the government and tax authorities collaborate with technology experts, data scientists, and other stakeholders. This collaboration will help in building a strong foundation for AI implementation, ensuring that it aligns with the unique needs and requirements of Bulgaria's tax system.

In conclusion, embracing AI in tax administration is not only a technological improvement but also a strategic necessity for Bulgaria. By harnessing the power of AI, Bulgaria can create a more efficient, fair, and taxpayer-friendly tax system that not only bolsters revenue collection but also enhances the overall economic and business environment in the country.

References

1. Ihnatišínová, D. (2021). Digitalization of tax administration communication under the effect of global megatrends of the digital age. SHS Web of Conferences, 92, 02022. <https://doi.org/10.1051/shsconf/20219202022>
2. Joshi, A., Prichard, W., & Heady, C. (2014). Taxing the informal economy: the current state of knowledge and agendas for future research. The Journal of Development Studies, 50(10), 1325-1347. <https://doi.org/10.1080/00220388.2014.940910>
3. Kochanova, A., Hasnain, Z., & Larson, B. (2018). Does e-government improve government capacity? evidence from tax compliance costs, tax revenue, and public procurement competitiveness. The World Bank Economic Review, 34(1), 101-120. <https://doi.org/10.1093/wber/lhx024>
4. Wegener, S., Labar, K., Petrick, M., Marquardt, D., Theesfeld, I., & Buchenrieder, G. (2011). Administering the common agricultural policy in bulgaria and romania: obstacles to accountability and administrative capacity. International Review of Administrative Sciences, 77(3), 583-608. <https://doi.org/10.1177/0020852311407362>
5. Huang, W. and Zhang, J. (2022). Artificial intelligence technology and tax risk management innovation.. <https://doi.org/10.1117/12.2641091>
6. Jiang, C. (2022). Research and implementation of tax planning platform based on big data and artificial intelligence.. <https://doi.org/10.1117/12.2658261>
7. Lu, Y. (2022). The influence of public mental health based on artificial intelligence technology on the teaching effect of business administration major. Journal of Environmental and Public Health, 2022, 1-11. <https://doi.org/10.1155/2022/5353889>
8. Zhang, J. (2023). Optimization of the environmental protection tax system design based on artificial intelligence. Frontiers in Environmental Science, 10. <https://doi.org/10.3389/fenvs.2022.1076158>
9. Zammit, C. (2023). New AI-powered software to be deployed to catch tax cheats. Times of Malta, <https://timesofmalta.com/articles/view/finance-ministry-deploy-ai-powered-software-catch-tax-cheats.1030430>
10. Clothilde, G., Manancourt, V., (2022). Dutch scandal serves as a warning for Europe over risks of using algorithms, <https://www.politico.eu/article/dutch-scandal-serves-as-a-warning-for-europe-over-risks-of-using-algorithms/>

ИЗПОЛЗВАНЕ НА МАШИННО ОБУЧЕНИЕ ЗА АВТОМАТИЗАЦИЯ ПРИ АНАЛИЗ НА НЕХАРМОНИЗИРАНИ ФИРМЕНИ ДАННИ

Using Machine Learning to Automate the Analysis of Unharmonized Company Data

Теодор Тодоров¹,
e-mail: ttodorov131@unwe.bg¹

Абстракт

Значително предизвикателство пред малкия и среден бизнес при извличането на знание се явява осигуряването на качествени хармонизирани данни. В малките и средните предприятия информацията често се съхранява във фрагментирани и слабо структурирани файлове, което прави изграждането на надеждни аналитични и прогностични системи сложно и времеемко. Настоящият доклад разглежда възможностите за използване на методи на машинното обучение при работа с нехармонизирани фирмени данни.

Abstract

A significant challenge for small and medium-sized enterprises (SMEs) in the process of knowledge extraction lies in ensuring the availability of high-quality, harmonized data. In SMEs, information is often stored in fragmented and poorly structured files, which makes the development of reliable analytical and predictive systems both complex and time-consuming. This paper explores the potential applications of machine learning methods for working with non-harmonized corporate data.

Ключови думи: извличане на данни, машинно обучение, хармонизация

JEL: C810

Увод

Дори и през 2025 година, малките и средни предприятия продължават да се сблъскват със съществени предизвикателства при анализа на ключова за техния бизнес информация. Непрекъснатото развитие на ERP и BI системите в наши дни води до наличието на все повече данни, които трябва да се превърнат в знание. На пръв поглед, компаниите могат да се възползват от огромен набор информация, чрез която да оптимизират своите процеси, но практиката показва задълбочаване на фрагментацията ѝ от различните платформи. Наличието на разнообразни пазарни решения с ниска месечна цена не разрешава този казус – на разположение са различни - и ефективни сами по себе си – софтуерни решения за счетоводство, продажби, CRM, логистика, управление на финансите. Въпреки всички предимства и по-ниската относителна цена, те често функционират като изолирани острови от данни, без стандартизирани интерфейси и без възможности за автоматизиран обмен на информация помежду си. В резултат на това, бизнес анализът започва с времеемко преобразуване на множество несвързани файлове с разнообразна структура и формати, което изисква значителен човешки ресурс и може да представлява значително предизвикателство пред малки и средни компании.

¹ Докторант, катедра „Информационни технологии и комуникации“, факултет „Приложна информатика и статистика“, УНСС, гр. София, ORCID 0009-0002-5119-8873

Допълнително предизвикателство се явява периодичността в необходимостта от анализи – с приключването на отчетен период или вътрешен организационен етап, компанията е принудена да възпроизвежда целия процес по интеграция и анализ на данните от различните платформи от самото начало. Това води до по-голямо забавяне и – в крайна сметка – игнориране на част от важната информация, която би могла да се превърне в допълнително знание. Разбира се, през последните години платформи като Microsoft Power BI, Tableau, Google Data Studio и други, които помагат за улесняване достъпа до визуализации, като позволяват на потребителите да обединяват данни от различни източници и да изграждат динамични табла за управление. Въпреки това, ефективността на тези решения зависи в голяма степен от предварителното качество и съгласуваност на данните, както и от ресурсите, които компанията може да отдели, за тяхното хармонизиране. Процесът на свързване, настройка на връзки, създаване на дименсии и осигуряване на актуални данни обикновено изисква значителен ръчен труд и техническа експертиза. По тази причина ще бъдат разгледани възможностите за различни подходи за автоматизирано извличане на знание от нехармонизирани фирмени данни, при които машинното обучение може да заеме ключова роля в ETL процеса.

Ключови положения при анализ на данни в малки и средни предприятия *Подготовка на данни. Хармонизирани и нехармонизирани данни*

В свое определение, изданието Gartner (2023) поставя акцент върху непрекъснатата повтаряемост на процеса по предварителната подготовка на данни. Това е сложна и многопластова дейност, която изисква значителна човешка намеса – и специалисти по данни, и бизнес анализатори, които да дадат възможност за извличане на знания за бизнес нуждите на организацията. Дори при използването на модерни инструменти за автоматизация на визуализациите, в основата на анализа остава ръчната работа по почистване и съгласуване на данните. Именно това обуславя и нуждата от нови подходи, а една възможност е машинното обучение да поеме част от тези задачи чрез автоматично откриване на зависимости, аномалии и липси.

В съвременната литература, терминът хармонизирани данни, *harmonized data*, обозначава съвкупност от различни по смисъл данни, които са съгласувани в „технически“ смисъл – имат еднакъв формат, структура, семантика. В този ред на мисли, различните източници на информация трябва да имат сходна логика на описание, позволяваща анализ и интерпретация (Inmon, 2005; Kimball and Ross, 2013). Хармонизацията има за цел да осигури взаимна съвместимост между системи, така че една и съща бизнес променлива – например приход – да бъде разбрана по еднакъв начин в различния бизнес контекст. В своя разработка от 2021 година, Николай Драгомиров и Любен Боянов (Dragomirov, Boyanov, 2021) изследват важни предизвикателства при анализа на данни от различни източници в сферата на логистиката, като подчертават възможностите за огромно количество потенциално знание. За да бъде достигнато то, обаче, стои необходимостта от уеднаквяване (хармонизация), което се извършва в по-голяма платформа за анализ на големи данни. Чрез хармонизация на данните, може да бъде постигната и „крайната цел“ в сферата – адекватен анализ. Разбира се, от другата страна, или по-точно – в началото на всяко намерение за анализ, стоят т.нар. „нехармонизирани данни“. Те са дефинирани като неструктурирани, несъвпадащи или частично-съвпадащи единици от информация, които се явяват основно предизвикателство пред автоматизираната обработка (Goodfellow, Bengio and Courville, 2016). Различия между данните често възникват още в създаването на отделните „късчета“ информация в ERP системите, тъй като различни

части на платформите могат да използват различни условности или формати за вход-изход. Липсата на автоматична проверка и централизирано управление на първичните данни може да доведе до т.нар. разминаване на значенията, *semantic drift* – в този случай, освен че данните стават трудни за интерпретация, те могат да доведат до възникването на неточни изводи (Ivanova, Iliev, Stoyanov, 2020). В свой труд за интернет на нещата (IOT), четирима български автори отбелязват съществената необходимост от хармонизация на данните като предпоставка за изграждането на интегрирани модели за анализ (Zlatev, Georgieva, Todorov, Stoykova, 2022). В същата разработка се отбелязва и че липсата на хармонизирани данни в повечето български компании може да бъде пряка последица от историческата зависимост към универсални инструменти (като Excel), което в контекста на настоящия доклад е съществена пречка пред автоматизацията анализа на ключови показатели

В следващите части на доклада, ще бъде обърнато внимание именно на някои автоматизирани подходи за хармонизация, при които алгоритми за машинно обучение могат да изпълняват задачи като:

- Автоматизирано сравнение имената на колони с различни наименования, т. нар. *schema matching*;
- Разпознаване и преобразуване на различни мерни единици, *unit conversion*;
- „Приписване на значения“ към данни с липсващи стойности, *data imputation*;
- Откриване и корекция на аномалии, *anomaly correction*.

Тези подходи биха могли да съкратят времето и разходите за хармонизация на данни и биха били особено полезни за малки и средни предприятия, като прехвърлят значителна част от „мисловната“ дейност от човека към компютърната система (Agrawal, 2023).

Съществуващи решения и ограничения.

Клиентски BI и ETL системи с модули с изкуствен интелект

Голяма част от малките и средни компании вече са добре запознати с текущите решения за визуализация на данни на пазара. Освен тази си основна функция, продукти като Power BI и Data Studio предлагат помощ при събиране и трансформация на данни. Добавяйки факта, че става дума за сравнително евтини решения с месечен абонамент (Software as a service, SaaS модел), малкият бизнес често е „изкушен“ от това да започне с някои от изброените решения. Част от възможностите на подобни BI системи са именно връзките с различни източници на данни – Excel файлове, извличане на данни чрез API и други. Основното предизвикателство, обаче, отново се корени във въпроса какво да направи предприятието с данните, преди да ги „подаде“ към BI платформата. Редица автори посочват, че тези платформи не решават основния проблем с качеството и съгласуваността на данните. В свое практическо изследване от 2021, няколко индийски учени изказват тезата, че въпреки изключителната си визуална мощ, BI инструментите остават зависими от предварително обработени и надеждни данни, без които техните изводи губят стойност (Vasudevan, 2021). Ако се опитаме да оборим това твърдение с наличието на все по-напреднали допълнения (add-ons) с изкуствен интелект в решенията на водещите софтуерни компании, ще установим, че дори и днес изкуственият интелект все още има нужда от редица насоки и потвърждения от потребителя, преди да пристъпи към конкретно действие. Добавяме и възможността от

допускане на голям брой еднотипни пропуски и автоматизацията на този процес се оказва подпомогната, но не и постигната, и продължава да бъде налице необходимост от тясно проследяване на всеки аспект от хармонизацията от страна на бизнес организацията. Казано накратко, макар решенията за бизнес анализи да стават все по-мощни и многопластови, те продължават да са последна инстанция от анализа на данни и са най-полезни при визуализацията на вече обработени данни. Дори и при наличие на инструменти с изкуствен интелект, се изисква от потребителя предварително да дефинира цялостната логика зад данните, която той трябва да разбира.

Low-code платформи за интегриране на данни

Агресивното развитие на облачните технологии спомогна и за утвърждаването на корпоративни платформи от тип Low-code. Това са софтуерни решения за ускорено разработване и поддръжка на приложения, използващи инструменти за разработка, базирани на модели; генеративен изкуствен интелект и предварително изградени каталози на компоненти за целия технологичен стек на приложението (Gartner, 2025). Примери за подобни продукти са Talend, Apache NiFi, Airbyte, RapidMiner, като те предоставят възможности за цялостно моделиране на ETL процеси и често се използват в научни и големи корпоративни среди за бързо моделиране на анализи. Low-code и open-source платформите като KNIME, RapidMiner и Talend имат за цел да улеснят потребителите чрез визуални среди за моделиране на процеси и готови конектори. Въпреки това, тяхното ефективно използване изисква концептуално разбиране на моделите на данните и основни умения за интеграция, което често надхвърля капацитета на малките предприятия (Gartner, 2024; Deloitte, 2022). В този смисъл, макар да елиминират нуждата от писане на код, те не елиминират нуждата от експертиза в анализа на данни и логиката на ETL процесите.

Умни ETL решения. AI-Enhanced ETL.

С нарастването на обемите от данни и необходимостта от бърз анализ, през последните години възниква ново поколение технологии, определяни като AI-Enhanced ETL (разширени чрез изкуствен интелект процеси по трансформация на данни, Intelligent Data Preparation Tools). Те намаляват времето за хомогенизиране и интеграция на фирмените данни, като позволяват на бизнес потребителите да използват едновременно различни източници. Освен това, те дават шанс на компаниите да идентифицират аномалии, да изследват и очертават модели, да подобрят качеството на данните от своите открития (Gartner, 2025). Такива инструменти използват изкуствен интелект и машинно обучение, за да автоматизират традиционните фази на извличане (Extract), трансформация (Transform) и зареждане (Load) на данни.

Класическият ETL процес работи по предварително зададени правила – всеки източник на данни трябва да бъде описан, всички полета – съпоставени, трансформациите – дефинирани.

AI-Enhanced ETL добавя **слой от интелигентност**, при който платформата:

- Автоматично **определя и категоризира** данните – открива типове полета, формати и зависимости;

- **Открива аномалии** и липси чрез алгоритми за машинно обучение (напр. чрез Isolation Forest, AutoEncoder);
- Извършва **автоматична импутация** на липсващи стойности на базата на вероятностни зависимости;
- **Прилага трансформации**, използвайки модели, обучени върху предишни операции;
- **Оптимизира последователността** на задачите.

По този начин системата не просто изпълнява инструкции, а се самообучава от предишни трансформации и грешки, намалявайки нуждата от човешка намеса при повторно извличане на данни. AWS Glue DataBrew, Google Cloud Dataprep (Trifacta), Snowflake Cortex и Databricks AutoML, например, прилагат елементи на ML за автоматично профилиране, откриване на аномалии и генериране на препоръки за трансформации. Според Gartner (2024), *„умните“ ETL инструменти се намират във фаза на преход от асистираща към автономна автоматизация*, като повечето решения все още предлагат препоръки, но не ги прилагат самостоятелно. С други думи, AI-Enhanced ETL все още не е напълно самообучаваща се система, но бележи първата стъпка към създаването на процеси, които автоматично се адаптират при промени в структурата или съдържанието на данните.

Въпреки напредъка, прилагането на изкуствен интелект в процесите по трансформация и зареждане на данни, среща няколко предизвикателства:

- Необходимост от вече налични големи обеми исторически данни за обучение на моделите;
- Възможни затруднения при интерпретацията на автоматично извършените трансформации (т.нар. explainability проблем – не е ясно защо „платформата“ взима конкретно решение);
- Повишени изисквания към изчислителните ресурси и сигурността (макар този казус да се разрешава сравнително лесно с използването на облачни ресурси)

В контекста на малките и средни предприятия, тези ограничения често правят внедряването им неефективно от финансова гледна точка. Въпреки това, тенденцията е ясно очертана – бъдещето на ETL процесите е в самообучаващите се модели, които ще позволят анализ върху нехармонизирани данни с минимална намеса на човека.

Направеният обзор показва, че независимо от напредъка в областта на бизнес анализа и автоматизацията на данните, съществуващите решения остават зависими от предварителна подготовка и ръчна намеса. Дори и най-модерните инструменти, които използват машинно обучение, все още не притежават пълна автономност и не могат самостоятелно да обработват нехармонизирани фирмени данни в тяхната първоначална форма. Това създава реална нужда от методология, която да обедини предимствата на ML-базираните алгоритми с практичността на BI инструментите, като по този начин осигури еднократно изграждане и последващо самообновяване на анализа при постъпване на нова информация. Следователно, в следващата глава се предлага рамка за автоматизация на анализа на нехармонизирани фирмени данни, основана на интеграция между алгоритми за машинно обучение и процесите на класическия ETL цикъл.

Анализ на възможностите за автоматизация при хармонизацията на данни чрез машинно обучение.

Още през 1959, Артър Лий Самюъл дефинира машинното обучение като наука, която изучава способността компютрите да се обучават за изпълнение на задача, без да бъдат конкретно програмирани за извършването на същата. След това, през 1986, Том Мичъл развива съвременната дефиниция за машинно обучение, като я представя като зависимост между ефективността за изпълнение на определена задача, и увеличаването на опита, натрупан от изпълнението на същата. Съвременните алгоритми могат да открият закономерности, групи и отклонения в данни с нееднородна структура и да подпомагат ETL процеса. Така например, модели като *k-nearest neighbors (KNN)* и *Multiple Imputation by Chained Equations (MICE)* се използват за импутация на липсващи стойности (Azur, 2011), докато *Isolation Forest* и *AutoEncoder Neural Networks* са ефективни за откриване на аномалии в таблични данни (Sakurada & Yairi, 2014). Тези методи позволяват на системата автоматично да допълва, коригира и валидира данните, намалявайки времето, необходимо за тяхното ръчно преглеждане. В контекста на малкия и среден бизнес прилагането на машинно обучение има стратегическо значение, защото премахва нуждата от дълбока техническа експертиза, характерна за големите организации. Както посочват Gartner (2024), нарастването на обемите от данни и разпространението на визуални инструменти за машинно обучение създават възможности за „демократизация“ на анализа – т.е. достъп до автоматизирани функции без специализирани познания по програмиране. Такива платформи могат да извършват *schema matching* (съпоставяне на несъответстващи колони), *unit normalization* (уеднаквяване на мерни единици и формати) и *semantic labeling* (автоматично приписване на значения към колони) въз основа на вече обучени модели (Wang, 2022). Тези функционалности превръщат машинното обучение в „интелигентен посредник“ между различните източници на информация и позволяват по-бързо достигане до консистентни, готови за анализ данни. Освен в ETL процеса, ML алгоритмите могат да бъдат използвани и за прогнозен анализ на ключови бизнес индикатори. Модели като *XGBoost* позволяват предсказване на тенденции на базата на исторически данни, като се адаптират автоматично към сезонни колебания и липсващи наблюдения – нещо, което традиционните статистически методи трудно постигат. Това е особено полезно за компании, които не разполагат със собствени отдели за анализ, но се нуждаят от бърза и достъпна прогноза за бъдещи продажби, разходи или пазарни рискове.

Процесът на хармонизация чрез машинно обучение обединява няколко взаимосвързани стъпки, които могат да се прилагат последователно или паралелно. Първата стъпка е *профилиране на данните*, при което алгоритми за клъстериране и класификация (напр. *k-means* или *decision trees*) автоматично откриват сходства между колони с различни наименования и определят каква е тяхната семантика — например „Client ID“, „Customer_Number“ и „Клиент“ се разпознават като идентификатори на едно и също понятие. След това се прилагат модели за *нормализиране на стойности*, които уеднаквяват формати и единици, като се обучават върху примери от предишни трансформации (Abdelaal, 2023). В третата стъпка се използват методи за *импутация* и *откриване на аномалии*, при които машинното обучение идентифицира липсващи или нетипични стойности и предлага вероятни корекции на базата на сходни наблюдения (Azur, 2011). Накрая, чрез *self-supervised* подходи, системата може да изгражда собствени представяния на данните (*embeddings*), които улесняват последващи съпоставяния между различни източници, дори когато имената и формите на колоните са напълно различни. По този начин хармонизацията престава да бъде изцяло ръчен

процес и се превръща в непрекъснат цикъл на самоусъвършенстване, при който моделите постепенно „научават“ структурата на корпоративните данни и я прилагат автоматично при нови източници.

Заклучение

В ерата на нарастващи обеми от информация малките и средните предприятия са изправени пред парадокс – разполагат с повече данни от всякога, но трудно ги превръщат в знание. Липсата на хармонизирани източници, разнообразието от формати и ограничените ресурси за поддръжка на сложни ИТ решения правят традиционните подходи за анализ неефективни. В този контекст машинното обучение предлага все по-реалистична и гъвкава алтернатива. В природата на ML е да открива закономерности, да попълва липси и да уеднаквява разнородни структури, като така значителна част от ръчната работа по подготовка и интерпретация на данните може да отпадне. Настоящият доклад разгледа концепцията за използване на алгоритми за машинно обучение при анализа на нехармонизирани фирмени данни, като очерта възможностите за автоматизация на ключови етапи от процеса – профилиране, нормализиране, импутация и откриване на аномалии. Представеният теоретичен модел демонстрира, че интеграцията между съществуващите BI инструменти и ML-базирани алгоритми може да изгради единен и самообновяващ се аналитичен цикъл, който съкращава времето за анализ и повишава надеждността на резултатите. Така малките и средните предприятия могат да използват предимствата на изкуствения интелект без необходимост от мащабна ИТ инфраструктура или сериозни програмни умения.

В по-широк план машинното обучение се явява не просто технологична иновация, а методологичен мост между данните и бизнес знанието. То позволява на организациите да преминат от еднократен, ръчно извършван анализ към непрекъснат процес на самоусъвършенстване, при който всяка нова информация автоматично подобрява точността на следващата. Това поставя основата за следващи изследвания, насочени към практическа реализация на предложената рамка и към разработване на достъпни, адаптивни решения за анализ и прогнозиране на данни в реалната бизнес среда на малките и средни предприятия.

Източници

1. Agrawal, R., Majumdar, A. & Kumar, A. et al. (2023). *Integration of Artificial Intelligence in Sustainable Manufacturing: Current Status and Future Opportunities*. *Operations Management Research*
2. Azur, M.J., Stuart, E.A., Frangakis, C. & Leaf, P.J. (2011). *Multiple Imputation by Chained Equations: What Is It and How Does It Work?* *International Journal of Methods in Psychiatric Research*. Available at: <https://pubmed.ncbi.nlm.nih.gov/21499542/> (Accessed: October 2025).
3. Boyanov, L. & Dragomirov, N. (2021). *Supply Chain Management and Logistics Big Data Challenges in Bulgaria*. *Academia.edu*. Available at: https://www.academia.edu/68932279/Supply_Chain_Management_and_Logistics_Big_Data_Challenges_in_Bulgaria (Accessed: October 2025).
4. Chen, T. & Guestrin, C. (2016). *XGBoost: A Scalable Tree Boosting System*. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*.

5. Gartner (2023). *Market Guide for Data Preparation Tools*. Gartner Inc. Available at: <https://www.gartner.com/en/documents/3987296> (Accessed: October 2025).
6. Gartner (2024). *Enterprise Low-Code Application Platforms*. *Gartner Peer Insights*. Available at: <https://www.gartner.com/reviews/market/enterprise-low-code-application-platform> (Accessed: October 2025).
7. Goodfellow, I., Bengio, Y. & Courville, A. (2016). *Deep Learning*. MIT Press. Available at: <https://www.deeplearningbook.org> (Accessed: October 2025).
8. Inmon, W.H. (2005). *Building the Data Warehouse* (4th ed.). Wiley.
9. Ivanova, E., Iliev, T. & Stoyanov, I. (2020). *Modelling and Simulation of Big Data Networks*. *IEEE Xplore*
10. Kimball, R. & Ross, M. (2013). *The Data Warehouse Toolkit: The Definitive Guide to Dimensional Modeling* (3rd ed.). Wiley.
11. Liu, F.T., Ting, K.M. & Zhou, Z.-H. (2012). *Isolation Forest*. *ACM Transactions on Knowledge Discovery from Data*.
12. Mittal, V. (2025). *Leveraging Artificial Intelligence to Optimize ETL Pipelines: Enhancing Efficiency, Accuracy and Scalability*. *ResearchGate preprint*. Available at: https://www.researchgate.net/publication/381925765_Leveraging_AI_to_Optimize_ETL_Pipelines (Accessed: October 2025).
13. Mitchell, T.M. (1997). *Machine Learning*. McGraw-Hill Education.
14. Sakurada, M. & Yairi, T. (2014). *Anomaly Detection Using Autoencoders with Nonlinear Dimensionality Reduction*. *Proceedings of the 2nd Workshop on Machine Learning for Sensory Data Analysis (MLSDA 2014)*.
15. Schrage, M. (2024). *The Future of Strategic Measurement: Enhancing KPIs With AI*. *MIT Sloan Management Review*. Available at: <https://sloanreview.mit.edu/projects/the-future-of-strategic-measurement-enhancing-kpis-with-ai> (Accessed: October 2025).
16. Soon, E., Vasudevan, A., Alshamayleh, A. & Mohammad, S. (2025). *Redefining Realities: AI's Cost-Efficient Revolution in Indian Cinema's VFX Landscape*. *Natural Sciences Publishing*. Available at: www.naturalspublishing.com/files/published/519fn13620io92.pdf (Accessed: October 2025).
17. Taylor, S.J. & Letham, B. (2018). *Forecasting at Scale*. *The American Statistician*.
18. Wang, Z., Li, H. & Wang, C. (2022). *Sudowoodo: Contrastive Self-Supervised Learning for Multi-purpose Data Integration and Preparation*. Available at: <https://arxiv.org/abs/2207.04122> (Accessed: October 2025).
19. Zlatev, Z., Georgieva, Ts., Todorov, A. & Stoykova, V. (2022). *Energy Efficiency of IoT Networks for Environmental Parameters of Bulgarian Cities*.

DIGITAL TWINS IN HEALTHCARE: SIMULATION, DATA, AND THE FUTURE OF PERSONALIZED MEDICINE

Lyuben Zyumbilski¹

e-mail: lzyumbilski@unwe.bg; lzyumbilski@gmail.com

Абстракт

Дигиталните близнаци – виртуални представяния на физически системи, които се актуализират непрекъснато с реални данни – се превръщат в трансформираща технология в здравеопазването. Чрез интеграция на физиологични модели, медицински данни и информация за пациенти в реално време, дигиталните близнаци позволяват симулации, прогнозиране и оптимизация на здравни състояния и резултати от лечението. В статията се разглежда архитектурата, интеграцията на данни, ролята на изкуствения интелект, етичните и правни аспекти, както и бъдещите тенденции.

Abstract

Digital twins-virtual representations of physical systems continuously updated with real-world data-are emerging as a transformative technology in healthcare. By integrating physiological models, medical data, and real-time patient information, digital twins enable simulation, prediction, and optimization of individual health states and treatment outcomes. This paper explores the concept and architecture of digital twins in medicine, focusing on data integration, modeling techniques, and the role of artificial intelligence in maintaining accurate and adaptive simulations. Ethical and privacy challenges are also discussed, along with the potential of digital twin ecosystems for clinical research, preventive medicine, and personalized treatment planning. The study concludes that digital twins represent a paradigm shift from reactive to proactive healthcare, driven by data, computation, and intelligent modeling.

Keywords: Digital Twin, Healthcare, Artificial Intelligence, Data Integration, Simulation

JEL: I10, O33, C88

Introduction

Healthcare systems are increasingly data-intensive, with electronic health records (EHRs), imaging archives, laboratory information systems, and connected medical devices generating high-velocity, high-volume streams of heterogeneous information. Yet clinical decision-making remains largely reactive-responding to events after they occur rather than anticipating and preventing them. Digital twins (DTs) propose an alternative paradigm: continuously updated computational representations that mirror the evolving state of a patient, organ, or clinical process. These representations enable simulation, prediction, and optimization, thereby supporting proactive, personalized care.

While DTs originated in manufacturing and aerospace, where they are used to monitor equipment, optimize maintenance, and enhance safety, their translation to healthcare introduces unique challenges. Biological systems are nonlinear, multi-scale, and subject to substantial inter-individual variability; moreover, healthcare data are fragmented across institutional, regulatory, and technical boundaries. This paper analyzes the foundational components needed to realize healthcare DTs, synthesizes

¹ Докторант към катедра ИТК, УНСС, email: lzyumbilski@unwe.bg; lzyumbilski@gmail.com

implementation patterns, and outlines the policy and ethical safeguards that must accompany deployment.

The contribution of this work is threefold. First, it proposes a clear conceptual architecture for healthcare digital twins that integrates data pipelines, computational models, and intelligent services. Second, it identifies data and algorithmic requirements for multi-scale modeling and continuous learning. Third, it discusses implementation barriers and governance mechanisms necessary to ensure safety, privacy, and equity. The analysis is scoped to clinical and operational use-cases and focuses on practical pathways from pilot to scale.

Concept and Architecture of Digital Twins

A healthcare digital twin can be conceptualized as a layered system that couples a physical entity (the patient or clinical process) with a virtual entity maintained by data-driven and mechanistic models. The coupling is bi-directional: telemetry, measurements, and clinical events update the virtual state, while simulated trajectories and counterfactuals inform decisions in the physical domain. Four architectural layers are commonly distinguished: the physical layer, the data acquisition and integration layer, the modeling and simulation layer, and the intelligence and interaction layer.

The physical layer encompasses sensors and devices (bedside monitors, implantables, wearables), as well as clinical workflows that produce discrete observations (laboratory results, imaging, clinician notes). The integration layer standardizes inputs using interoperability profiles (e.g., HL7 FHIR resources for clinical data and DICOM objects for imaging) and performs data quality management, identity resolution, and semantic harmonization. Data are persisted in secure repositories with auditable access control. Streaming pipelines enable near real-time ingestion for time-sensitive monitoring tasks.

The modeling and simulation layer combines mechanistic physiology (e.g., compartmental or finite-element models) with data-driven components (e.g., state-space models, recurrent neural networks) to form hybrid digital representations. Model calibration uses patient-specific parameters inferred from historical and current data. The intelligence and interaction layer exposes services-risk scoring, anomaly detection, what-if simulations-through decision-support interfaces. Importantly, this layer maintains explainability artifacts (saliency maps, feature attributions, uncertainty estimates) to support clinical validation and trust.

Data Integration and Multi-Scale Modeling

Constructing clinically useful digital twins requires reconciling heterogeneous data with varying cadence, fidelity, and semantics. Time-aligned integration is essential: physiological sensor streams (millisecond to minute resolution) must be synchronized with episodic EHR updates and periodic imaging. Robust timestamping, clock synchronization, and data versioning are necessary to reconstruct a faithful longitudinal state. Missingness patterns-common in routine care-should be modeled explicitly using probabilistic frameworks rather than imputed naively.

From a representational standpoint, graph-based data structures are advantageous. Patients, encounters, observations, and interventions can be modeled as typed nodes and edges, allowing the twin to capture temporal dependencies and causal pathways. Graph neural networks (GNNs) extend this representation with learnable message passing, supporting individualized predictions. At the same time, mechanistic models encode biophysical constraints that stabilize learning and improve extrapolation.

Multi-scale modeling links molecular-level dynamics (e.g., pharmacokinetics/pharmacodynamics), tissue or organ-level behavior (e.g., electrophysiology), and system-level responses (e.g., hemodynamics, metabolism). Coupling across scales can be achieved through surrogate models-reduced-order approximations trained on high-fidelity simulations-that provide tractable components for real-time use. Bayesian calibration harmonizes priors from population studies with patient-specific data to update parameter distributions and quantify uncertainty.

Interoperability remains a practical bottleneck. Adoption of open standards such as HL7 FHIR for resources (Patient, Observation, MedicationStatement), DICOM for imaging, and OMOP CDM for observational research enables reproducible pipelines and cross-institution validation. Privacy-preserving computation-federated learning, secure enclaves, and differential privacy-allows multi-center modeling without centralizing sensitive data.

Artificial Intelligence and Simulation

AI augments digital twins along three axes: perception, prediction, and control. Perception algorithms transform raw signals into semantically rich features-arrhythmia detection from ECG, motion artifacts from accelerometers, phenotype extraction from clinical notes using natural language processing (NLP). Prediction models estimate risk trajectories, treatment response, and counterfactual outcomes under alternative interventions. Control leverages reinforcement learning and model-predictive control to recommend actions subject to safety constraints.

A critical requirement is continual learning. As clinical practice evolves and populations shift, static models decay. Online learning paradigms-bounded by safeguards such as shadow deployment, drift detection, and human-in-the-loop review-enable twins to adapt without compromising safety. Uncertainty quantification (e.g., ensembles, Bayesian neural networks) is essential for calibrated decision-making, with thresholds and alerts tuned to clinical utility rather than raw accuracy metrics.

Explainability must align with the task. Clinicians often prefer concise, mechanistically plausible rationales over opaque saliency maps. Hybrid twins offer a practical route: they fuse mechanistic structures-where cause-effect relations are explicit-with machine-learned residuals that capture idiosyncratic patterns. This design yields interpretable outputs while maintaining predictive performance.

Practical Applications and Case Studies

Patient monitoring twins integrate bedside telemetry and wearable data to detect early physiological deterioration in general wards. By modeling baseline variability and individual set points, the twin suppresses false alarms and elevates clinically meaningful deviations. In perioperative care, twin-guided hemodynamic simulations inform fluid management strategies tailored to patient-specific cardiovascular dynamics.

In chronic disease management, long-term twins synthesize medication adherence, activity profiles, and lab trends to anticipate exacerbations and recommend adjustments. Hospital operations benefit from process twins that simulate patient flow and resource allocation, enabling proactive staffing and bed management during demand surges. Pharmaceutical R&D explores cohort-level twins-population models calibrated to trial data-to optimize protocol design and identify responders.

Early industrial deployments demonstrate feasibility: vendors have piloted DT-enabled monitoring platforms that integrate EHR streams with telemetry for real-time risk stratification, while academic-

industry consortia have produced high-fidelity organ models (e.g., cardiac electrophysiology) validated against clinical outcomes. Although many pilots remain pre-market, they provide a template for translating DTs from research to clinical practice.

Implementation Barriers and Policy Perspectives

Scaling digital twins beyond pilots faces organizational, technical, and regulatory barriers. Data custodianship is fragmented; hospitals, device manufacturers, and software vendors maintain siloed systems with divergent incentives. Robust data-sharing agreements and standardized APIs are prerequisites for end-to-end integration. Economically, sustainable business models must demonstrate value through reduced adverse events, shorter length of stay, or optimized resource utilization.

From a safety standpoint, DTs implicate software as a medical device (SaMD) regulations. Lifecycle management requires documented verification and validation, real-world performance monitoring, and post-market surveillance. Change control for continuously learning systems must be formalized: updates should pass through staged evaluation with rollback capability and audit trails. Clinical governance needs clear delineation of responsibility between decision-support recommendations and clinician judgment.

Equity considerations demand rigorous bias assessment and inclusivity in training data. Federated studies should ensure representation across demographic groups and care settings to avoid performance gaps. Policymakers can encourage safe innovation through regulatory sandboxes, reference implementations, and public datasets curated to benchmark fairness, robustness, and privacy-preserving efficacy.

Future Trends and Challenges

Three technical trends will shape the next decade of healthcare DTs. First, high-performance and cloud-edge hybrid computing will enable real-time multiscale simulations at the point of care. Second, standardized provenance and model cards will make DT components composable and auditable, supporting reuse across institutions. Third, privacy-enhancing technologies-federated analytics with secure aggregation, homomorphic encryption for selected operations-will allow richer collaboration without exposing raw data.

On the methodological front, digital physiology libraries will expand, providing validated mechanistic building blocks (e.g., circulatory models, pulmonary mechanics) that can be parameterized for individuals. Surrogate modeling and physics-informed neural networks will accelerate complex simulations. Human factors engineering will gain prominence: DT interfaces must integrate seamlessly into clinical workflows, minimize alert fatigue, and present recommendations with appropriate context and uncertainty.

A realistic trajectory is progressive deployment: start with narrow, high-value tasks (e.g., deterioration detection on a specific ward), establish data and governance foundations, and iteratively broaden scope. Success metrics should extend beyond AUC to encompass calibration, clinical impact, equity, and cost-effectiveness. Ultimately, DTs can help shift health systems from episodic care to continuous, preventive, and personalized management.

Conclusion

Digital twins offer a coherent framework that unifies data integration, modeling, and AI-driven decision support for personalized and proactive healthcare. By mirroring patient states and simulating

interventions, DTs can improve timeliness of care, reduce avoidable harm, and optimize resource use. However, realizing this potential requires rigorous engineering, careful governance, and sustained collaboration across clinical, technical, and policy domains.

This paper articulated a reference architecture, clarified data and modeling requirements, and examined practical applications alongside barriers to scale. Future work should prioritize prospective evaluations in diverse health systems, transparent reporting of real-world performance, and open, standards-based ecosystems that enable reproducibility and trust.

References

1. Björnsson, B., Borrebaeck, C., Elander, N. et al. (2020). Digital twins to personalize medicine. *Genome Medicine*, 12(1), 4.
2. Corral-Acero, J., Margara, F., Marciniak, M. et al. (2020). The 'Digital Twin' to enable the vision of precision cardiology. *European Heart Journal*, 41(48), 4556–4564.
3. Fuller, A., Fan, Z., Day, C., & Barlow, C. (2020). Digital Twin: Enabling Technologies, Challenges and Open Research. *IEEE Access*, 8, 108952–108971.
4. Kapteyn, M. G., et al. (2021). Physics-informed machine learning: modeling and simulation. *Annual Review of Fluid Mechanics*, 53, 491–518.
5. Viceconti, M., & Hunter, P. (2016). The Virtual Physiological Human – A Review of Integrative Biomedical Simulation. *Interface Focus*, 6(2), 20150092.
6. Bruynseels, K., Santoni de Sio, F., & van den Hoven, J. (2018). Digital Twins in Health Care: Ethical Implications of an Emerging Engineering Paradigm. *Frontiers in Genetics*, 9, 31.
7. Wong, K. C., & Goh, W. L. (2022). Edge Computing for Medical Digital Twins. *Journal of Biomedical Informatics*, 130, 104086.
8. Topol, E. (2019). High-performance medicine: the convergence of human and artificial intelligence. *Nature Medicine*, 25, 44–56.
9. Rumsfeld, J. S., Joynt, K. E., & Maddox, T. M. (2016). Big Data Analytics to Improve Cardiovascular Care: Promise and Challenges. *Journal of the American College of Cardiology (JACC)*, 69(10), 1187–1199.
10. Rajpurkar, P., Chen, E., Banerjee, O., & Topol, E. J. (2022). AI in health and medicine. *Nature Medicine*, 28, 31–38.
11. Björnsson, B., et al. (2021). Personalized medicine and digital twins: challenges and opportunities. *Frontiers in Genetics*, 12, 626685.
12. European Commission. (2023). DigiTwins Project – Digital Twins for Preventive Health. Horizon Europe Programme. <https://cordis.europa.eu/project/id/777435>
13. WHO (2023). Ethics and Governance of Artificial Intelligence for Health. Geneva: World Health Organization. <https://www.who.int/publications/i/item/9789240029200>
14. Hunter, P., Viceconti, M., & Coveney, P. (2022). Computational biomedicine and the digital twin: The road ahead. *Medical Engineering & Physics*, 104, 103812.
15. Tolk, A., Diallo, S. Y., & Turnitsa, C. (2020). Modeling and Simulation Support for System of Systems Engineering Applications. John Wiley & Sons.

ARTIFICIAL INTELLIGENCE IN MEDICAL DIAGNOSTICS: ALGORITHMS, DATA, AND CHALLENGES IN PRACTICAL IMPLEMENTATION

Lyuben Zyumbilski¹

e-mail: lzyumbilski@unwe.bg; l.zyumbilski@gmail.com¹

Абстракт

В доклада се разглеждат алгоритмите, данните и инженерните предизвикателства при прилагането на изкуствен интелект (ИИ) в медицинската диагностика. Представя се структуриран преглед на типовете медицински данни (образни, текстови, физиологични), моделите за машинно и дълбоко обучение, жизнения цикъл на разработване на AI системи и метриките за оценка. Анализирани са етичните, правните и организационните аспекти (GDPR, прозрачност, пристрастия), както и практически препятствия при внедряване в клинична среда. Целта е да се формулират технически принципи за надеждни, обясними и безопасни AI системи, които подпомагат диагностиката, без да заменят клинициста.

Abstract

Artificial intelligence (AI) is reshaping medical diagnostics by transforming heterogeneous data-imaging, clinical text, and physiological signals-into actionable predictions that support clinicians. This paper surveys core algorithmic approaches (supervised learning, deep learning, self-supervised and foundation models), data management requirements, and the development lifecycle for diagnostic AI systems. We emphasize validation methodologies, calibration and generalization across sites, as well as workflow integration and human-in-the-loop oversight. Ethical, legal, and organizational challenges are discussed with reference to GDPR, transparency, bias, and accountability. The paper distills engineering principles for building reliable, explainable, and safe AI tools that augment, rather than replace, clinical expertise.

Keywords: Medical diagnostics, Artificial intelligence, Machine learning, Medical imaging, Model validation

JEL: I10, O33, C88

Introduction

Medical diagnostics relies on interpreting complex evidence under uncertainty-radiology images, laboratory values, clinical narratives, and vital-sign trajectories. Over the past decade, advances in machine learning (ML) and deep learning (DL) have enabled algorithms to detect patterns that are difficult for humans to perceive, especially in high-dimensional imaging and time-series data. Despite rapid progress in research benchmarks, translating AI into trustworthy clinical tools requires careful engineering, rigorous validation, and governance. This paper presents a practice-oriented synthesis for readers with an informatics background.

We structure the discussion around five pillars: (i) data types and pipelines; (ii) algorithmic methods for perception and prediction; (iii) development lifecycle and evaluation; (iv) workflow integration and

¹ Докторант към катедра ИТК, УНСС, email: lzyumbilski@unwe.bg; l.zyumbilski@gmail.com

human oversight; and (v) ethics, regulation, and deployment barriers. The scope is diagnostic support; we deliberately avoid therapeutic recommendations or clinical management guidance.

Medical Data for Diagnostic AI

Diagnostic AI consumes heterogeneous data. Imaging modalities-X-ray, CT, MRI, ultrasound, digital pathology-provide high-dimensional visual inputs with varying resolution, noise characteristics, and acquisition protocols. Textual data from electronic health records (EHRs)-problem lists, radiology reports, discharge summaries-capture clinical context but are unstructured and institution-specific. Physiological signals (ECG, PPG, EEG, respiratory impedance) combine high temporal resolution with device-dependent artifacts.

Data governance and quality are decisive. Label provenance should be explicit: image-level labels derived from expert reports, region-level annotations for localization tasks, and patient-level outcomes for prognostic models. Weak supervision (report-derived labels) scales annotation but introduces noise; consensus reading and adjudication reduce bias. Data splitting must be patient-wise and time-aware to prevent leakage. Cross-site external validation is essential to estimate generalization under domain shift.

Interoperability standards facilitate robust pipelines: DICOM for imaging objects and metadata; HL7 FHIR resources for problems, observations, procedures; and controlled vocabularies such as SNOMED CT and LOINC. Privacy-preserving collaboration-federated learning and secure aggregation-enables multi-center models without centralizing sensitive data.

Algorithms for Diagnostic Tasks

Supervised learning remains the workhorse for classification, detection, and segmentation. Convolutional neural networks (CNNs) dominate image tasks, while transformers and hybrid CNN-transformer architectures increasingly match or surpass CNN baselines. For text, transformer-based language models fine-tuned on clinical corpora perform named-entity recognition, report generation, and entailment tasks. For signals, temporal CNNs and recurrent architectures model rhythm and morphology.

Self-supervised learning (SSL) and foundation models mitigate label scarcity by pretraining on large unlabeled datasets, then fine-tuning for specific tasks. Multi-modal models combine images, text reports, and tabular context to improve discrimination and reduce spurious correlations. Calibration techniques (temperature scaling, isotonic regression) convert raw scores into well-calibrated probabilities suitable for decision thresholds.

Interpretability is task-dependent. Saliency and attribution maps (Grad-CAM, integrated gradients) provide visual rationales in imaging; prototype-based networks and concept bottlenecks offer human-interpretable intermediate spaces. Nonetheless, explanation artifacts must be validated, as misleading heatmaps can arise from confounders.

Development Lifecycle and Evaluation

A disciplined lifecycle spans problem definition, data curation, model development, validation, and post-deployment monitoring. Problem statements should specify intended use, target population, input constraints, and acceptable trade-offs between sensitivity and specificity. Dataset documentation (data sheets, model cards) records provenance, inclusion criteria, pre-processing, and known limitations.

Evaluation must go beyond internal cross-validation. Hold-out test sets from the same site estimate in-distribution performance; external validation on geographically or temporally distinct cohorts probes robustness. Metrics should align with clinical utility: area under the ROC curve (AUC) is insufficient alone; sensitivity at fixed specificity, positive and negative predictive values at plausible prevalence, decision curve analysis, and calibration error (ECE) are critical. For detection/segmentation, mean average precision and Dice/F1 scores are standard.

Uncertainty quantification and safety are integral. Techniques include deep ensembles, Monte-Carlo dropout, and conformal prediction to bound error rates. Dataset shift monitoring and periodic re-calibration maintain reliability as practice changes. Human-in-the-loop review policies define when and how clinicians can override or seek additional data.

Workflow Integration and Human Oversight

Embedding AI into clinical workflows requires interoperability with existing systems (PACS/RIS, EHR) and careful human factors design. Alerting and worklist triage should prioritize actionable findings without inducing alarm fatigue. User interfaces must present concise rationales, uncertainty indicators, and links to source data for rapid verification.

Prospective evaluations-silent mode shadow deployment, randomized workflow studies, and stepped-wedge designs-quantify operational impact on turnaround time, recall rates, and downstream testing. Governance boards should oversee updates, monitor performance across subgroups, and ensure traceable audit logs. Importantly, AI outputs are decision support; final judgments remain with clinicians.

Ethical, Legal, and Regulatory Considerations

Diagnostic AI engages sensitive personal data, triggering obligations under GDPR for lawful basis, data minimization, purpose limitation, and data-subject rights. De-identification and pseudonymization reduce risk but do not eliminate re-identification concerns with rich imaging and text. Transparency requirements imply documenting data use and model behavior.

Fairness requires systematic subgroup analysis by sex, age, ethnicity, and scanner/site characteristics. When disparities appear, mitigation strategies include re-sampling, re-weighting, domain adaptation, and targeted fine-tuning. Explainability should be fit-for-purpose-sufficient for clinician understanding without implying unwarranted certainty.

Regulatory pathways classify many diagnostic AI tools as Software as a Medical Device (SaMD). Verification and validation, post-market surveillance, change control for learning systems, and real-world performance monitoring are essential elements of compliant lifecycle management.

Future Directions

Research is converging on multi-modal, self-supervised, and foundation models that transfer across tasks and institutions with minimal labeled data. Federated learning with secure aggregation will enable cross-hospital collaboration while preserving privacy. Advances in uncertainty estimation and causal representation learning promise better generalization and more reliable decision support.

Operationally, the most sustainable deployments target narrow, high-impact tasks-standardized image quality checks, structured report extraction, or risk triage for defined populations. Success depends on clear intended use, transparent reporting, and continuous monitoring, not solely on benchmark accuracy.

Conclusion

AI for medical diagnostics holds significant promise when engineered and governed rigorously. This paper synthesized data pipelines, algorithmic methods, evaluation practices, and deployment considerations with an emphasis on safety, calibration, and human oversight. Properly designed AI systems can enhance diagnostic accuracy and efficiency while respecting privacy, fairness, and accountability.

For practitioners with an informatics background, the path to impactful diagnostic AI lies in disciplined problem definition, high-quality data engineering, robust validation, and thoughtful workflow integration. These principles form the foundation for translating research into reliable clinical decision support.

References

1. Rajpurkar, P., Chen, E., Banerjee, O., & Topol, E. J. (2022). AI in health and medicine. *Nature Medicine*, 28, 31–38. <https://doi.org/10.1038/s41591-021-01614-0>
2. McKinney, S. M., Sieniek, M., Godbole, V. et al. (2020). International evaluation of an AI system for breast cancer screening. *Nature*, 577, 89–94. <https://doi.org/10.1038/s41586-019-1799-6>
3. Esteva, A., Kuprel, B., Novoa, R. A. et al. (2017). Dermatologist-level classification of skin cancer with deep neural networks. *Nature*, 542, 115–118. <https://doi.org/10.1038/nature21056>
4. Litjens, G., Kooi, T., Bejnordi, B. E. et al. (2017). A survey on deep learning in medical image analysis. *Medical Image Analysis*, 42, 60–88. <https://doi.org/10.1016/j.media.2017.07.005>
5. Lundervold, A. S., & Lundervold, A. (2019). An overview of deep learning in medical imaging focusing on MRI. *Frontiers in Neuroscience*, 13, 530. <https://doi.org/10.3389/fnins.2019.00537>
6. Azizi, S., Mustafa, B., Ryan, F. et al. (2021). Big self-supervised models advance medical image classification. *ICCV*, 3478–3488. <https://doi.org/10.1109/ICCV48922.2021.00348>
7. De Fauw, J., Ledsam, J. R., Romera-Paredes, B. et al. (2018). Clinically applicable deep learning for diagnosis and referral in retinal disease. *Nature Medicine*, 24, 1342–1350. <https://doi.org/10.1038/s41591-018-0107-6>
8. Kelly, C. J., Karthikesalingam, A., Suleyman, M., Corrado, G., & King, D. (2019). Key challenges for delivering clinical impact with AI. *Nature Medicine*, 25, 1167–1176. <https://doi.org/10.1038/s41591-019-0546-4>
9. WHO (2021). Ethics and governance of artificial intelligence for health. World Health Organization. <https://www.who.int/publications/i/item/9789240029200>
10. European Union (2016). General Data Protection Regulation (GDPR) - Regulation (EU) 2016/679.
11. Boag, W., Wacome, K., Naumann, T., & Szolovits, P. (2020). From notes to knowledge: Clinical NLP for EHRs. *Patterns*, 1(9), 100171. <https://doi.org/10.1016/j.patter.2020.100171>
12. Irvin, J., Rajpurkar, P., Ko, M. et al. (2019). CheXpert: A large chest radiograph dataset with uncertainty labels and expert comparison. *AAAI*, 590–597.
13. Johnson, A. E. W., Pollard, T. J., Shen, L. et al. (2016). MIMIC-III, a freely accessible critical care database. *Scientific Data*, 3, 160035. <https://doi.org/10.1038/sdata.2016.35>

14. Reyes, M., Meier, R., Pereira, S. et al. (2020). On the interpretability of AI in radiology: Challenges and opportunities. *Radiology: AI*, 2(3), e190043. <https://doi.org/10.1148/ryai.2020190043>
15. Oakden-Rayner, L. (2020). The revolution will not be supervised: self-supervised learning in medical imaging. *arXiv:2006.12313*

МЕТОДОЛОГИЧНИ ПОДХОДИ В МУЛТИМОДАЛНИЯ СЕНТИМЕНТ АНАЛИЗ

Methodological Approaches in Multimodal Sentiment Analysis

Станимира Йорданова¹

e-mail: syordanova@unwe.bg¹

Абстракт

Мултимодалният сентимент анализ се утвърждава като изследователска област, която обхваща обработката на естествен език, компютърното зрение и обработката на реч. За разлика от традиционните подходи, които разчитат единствено на текст, мултимодалните подходи интегрират хетерогенни източници на информация като език, визуални сигнали и аудио характеристики, за да се улови сложността на човешките емоции и мнения. Докладът представя преглед на набори от данни за мултимодалния сентимент анализ и извежда методологични архитектурни подходи, използвани в анализа. В докладът се обсъждат важните предизвикателства като синхронизацията между модалностите, дисбалансът в данните и зависимостта от контекста. Докладът има за цел да допринесе за развитието на устойчиви, мащабируеми и контекстно осъзнати модели за приложение в реални среди на сентимент анализ.

Abstract

Multimodal sentiment analysis is a research area that encompasses natural language processing, computer vision, and speech processing. Unlike traditional approaches that rely solely on text, multimodal methods integrate heterogeneous sources of information such as language, visual, and audio features to capture the complexity of human emotions and opinions. The paper presents popular datasets and methodological approaches used in the multimodal sentiment analysis architecture. Challenges such as synchronization between modalities, data imbalances, and contextual dependency are discussed. The paper aims to contribute to the development of more sustainable, scalable, and context-aware models for application in real-world sentiment analytics environments.

Ключови думи: multimodal sentiment analysis, NLP, aspect-based sentiment analysis

Въведение

Сентимент анализът е задача в обработката на естествен език (НЛП), която има за цел да се извлекат човешките емоции и мнения в комуникацията от текстови източници като отзиви за продукти и публикации в социалните медии. Тъй като общуването между хората е сложен и динамичен процес, който съчетава вербални и невербални знаци, традиционните системи за сентимент анализ, използвайки само текстови източници, често не успяват да уловят нюансите на емоциите, които се изразяват в общуването чрез взаимодействие на езиково съдържание, интонация на речта и изражения на лицето [1]. Тези ограничения неизменно водят до необходимостта от развитие на мултимодалния сентимент анализ, който има за цел да интегрира субективната информация от различните източници като текст, разговори и видеоклипове. Източниците на комуникация са богати на текстови, аудио и визуални данни и предоставят възможност за правилно определяне и извличане на мненията (чувства и емоции), изразени от хората за обекти и техните характеристики. Сложността на мултимодалния сентимент анализ произтича от необходимостта не само да се идентифицират конкретни обекти или характеристики (аспекти) в съдържанието, но и точно да се определят настроеността и мненията,

¹ Гл. ас. д-р към катедра ИТК, УНСС, email: syordanova@unwe.bg

свързани с всеки от тях, като се използват контекстуални знаци от текстова, визуална и аудио информация, както и да се постигне точно разбиране в нюансите на настроеността и мненията, осигурявайки цялостен поглед върху човешката комуникация. Ефективното извличане на настроеността и мненията чрез мултимодален сентимент анализ зависи до голяма степен от висококачествени, анотирани мултимодални набори от данни, като обработката им поставя предизвикателства пред способността на изкуствения интелект да интерпретира човешка комуникация.

Докладът има за цел да представи методологичните подходи в мултимодалния сентимент анализ от гледна точка на задачите в извличането на настроеността и мнения от различни типове източници. В доклада се изследват характеристиките на необходимите данни и изискванията към тях, като се подчертава значението на процеса на анотиране на данни и постигането на представителност, количество и качество на данните. Представена е логическа архитектура на мултимодалния сентимент анализ, изследвани са методологични рамки за мултимодален сентимент анализ и са изведени предизвикателства, свързани със синхронизацията между модалностите, дисбаланса в данните, зависимостта от контекста и изискванията към изчислителната ефективност.

Мултимодален сентимент анализ

Традиционният подход за сентимент анализ, който има за цел да извлича мнение от текстов източник, класифицира мнението на три различни нива – *документ, изречение и аспект на обекта*. За разлика от сентимент анализа, който може да класифицира цял документ или изречение като положително, отрицателно или неутрално, мултимодалният сентимент анализ на ниво аспект цели идентифициране на конкретни обекти или характеристики на обекти (аспекти) в съдържанието и след това да определи точната полярност на мнението (напр. положителна, отрицателна или неутрална), свързана с извлечените конкретни аспекти. Този детайлен анализ позволява много по-богато и приложимо разбиране на мненията и емоциите, защото базира анализа на информация от различни типове източници.

Развитието на методологичните подходи за мултимодален сентимент анализ претърпява еволюция, насочена към осъвършенстване на методите и алгоритмите за извличане на структура от различните типове данни с цел точното определяне на мнението. С сентимент анализът преминава през няколко етапа, които се влияят от нарастването на данните в интернет пространството, развитието на методите за подготовка на данните и машинно обучение, появата на съвременни методи за класификация и развитието на изкуствения интелект. В ранния етап на своето развитие (преди 2010г.), сентимент анализът е фокусиран върху определяне на полярността на мнението на ниво документ, след което преминава към сентимент анализ на ниво изречение и ниво аспект. В този период се прилагат класически методи от машинно обучение, методи, използващи лексикони, както и хибридни методи. С появата на въвеждането на думи и дълбокото обучение, се развиват методологичните подходи за откриване и извличане на аспекти и определяне на мнението по отношение на извлечените аспекти. След 2010г. се достига до признанието, че текстът сам по себе си е недостатъчен за сентимент анализ. Изображенията, аудио и видеото могат да добавят богат контекст, което подтиква появата на техники за мултимодален синтез или интегриране на данни за сентимент анализ от различни типове източници. [1]. През 2020г. напредъкът в развитието на трансформър архитектурите, обучени модели като BERT, RoBERTa и мултимодални трансформър архитектури (VisualBERT), прави възможно свързването на текст с изображение и извличане на съответствие на аспектите от текста с визуалните знаци от изображенията. След 2023г., въвеждането на големите езикови модели прави възможна интеграцията им с мултимодални източници (текст, изображение, аудио) и дава насока на научната дейност към мултимодален сентимент анализ на ниво аспект, използващ промптове и инструкции към генеративните езиковите модели. Съвременните тенденции в развитието на мултимодалния сентимент анализ са насочени към използване на големи езикови модели за обяснимост, липса на езикови, звукови ресурси и изображения, както и приложения за конкретни области (напр. здравеопазване, електронна търговия, мониторинг на социалните медии).

Източници на данни за мултимодален sentiment анализ

В. Liu през 2012г. [2], определя мнението (sentiment, opinion) като съвкупност от четири елемента *обект/аспект на мнение, изразено чувство (feeling) за обекта/аспект на обекта, автор на мнението и датата, когато е изразено мнението*. В комуникацията си, човекът може да изразява *емоции и чувства*. От психологическа гледна точка *емоциите* обикновено са краткосрочни състояния, които се задействат от конкретни събития и включват последователност от когнитивни оценки, телесни реакции и изразения като движения на лицето или гласови знаци. Например моментен поглед на изненада или мимолетно изразяване на гняв. *Чувството* се определя като по-дълбоко, по-трайно разположение или установено мнение, често описвано като позиция към конкретна същност или концепция. Положителното чувство на човек към обект (например продукт или политическа партия) е гледна точка, която може да бъде идентифицирана чрез последователни афективни взаимодействия с обекта. Докато емоциите могат да бъдат преходни, чувствата са по-стабилни и могат да бъдат проследени чрез повтарящи се емоционални реакции. Това разграничение е важно, тъй като системите за мултимодален sentiment анализ могат да бъдат проектирани да идентифицират краткосрочни емоционални състояния или дългосрочни чувства (мнения).

Източниците на данни за мултимодален sentiment анализ са **текст, изображение и аудио**.

- *Текстовите източници* предоставят съдържание и контекстуална информация, която е в основата на идентифицирането на аспекти и мнение. Въпреки това, в текста може да бъде изразено двусмислие, сарказъм или ирония, които често се пропускат в моделите за sentiment анализ от текстови източник.
- *Аудио източниците* предоставят информация за знаци, които липсват в писмения текст като емоционален тон, предаван чрез речеви характеристики - височина, ритъм и сила на звука. Тонът на говорещия може да подсили или да противоречи на текстовото съдържание, осигурявайки важна информация за точно класифициране на мненията.
- *Визуалните източници* предлагат невербална информация от изразения на лицето, жестовите и езика на тялото, която подпомага анализирането на човешкото състояние и поведение.

Обучението на ефективни модели за мултимодален sentiment анализ изисква разнообразни и представителни набори от данни. Висококачествените данни са от съществено значение за точното и ефективно обучение на моделите, тъй като позволяват на моделите да обобщават добре нови данни, като гарантират тяхната приложимост в реални сценарии. Данни с ниско качество, които могат да съдържат неточности, шум или несъответствия, могат да доведат до ненадеждни и по-малко точни прогнозни модели, които се учат от тези недостатъци. Некачествените данни поставят предизвикателства в обработката на данните и в обучението, свързани с шум, липсващи данни или отклонения, въведени по време на събирането на данни или анотацията. Придобиването на голям обем висококачествени, анотирани мултимодални данни е основна пречка за обучение на стабилни модели за sentiment класификация, защото извличането и обработката на мултимодални данни са свързани с различни ограничения, включително проблеми с неприкосновеността на личния живот, строги регулаторни рамки и значителните разходи и специализиран опит за последователна анотация на данните.

За целите на изследователската дейност различни набори от данни са извлечени, анотирани и предоставени за използване, като се превръщат във важни еталони за сравнение в областта на мултимодалния аспекти-базиран анализ на мнения [3]. В таблица 1 е направено сравнение на набори от данни, извлечени и използвани за мултимодален sentiment анализ.

Таблица 3: Сравнение на набори от данни за мултимодален сентимент анализ

Набор от данни	Година	Обем	Модалности	Задача/Област на приложение
MOUD [4]	2013	498 изказвания	Текст, аудио, видео	Полярност на мнения Отзиви за продукти, испански език
Twitter-2015 [5]	2015	5,338 туитове	Текст, изображения	Мнение за аспект (3 класа), Туитър данни, английски език
Twitter-2017 [5]	2017	5,972 туитове	Текст, изображения	Мнение за аспект (3 класа), Туитър данни, английски език
CMU-MOSI (Multimodal Opinion-Level Sentiment Intensity) [6]	2016	~2199 сегмента на мнението	Текст, аудио, видео	Интензивност на настроеността YouTube отзиви
CMU-MOSEI (Multimodal Opinion Sentiment and Emotion Intensity) [7]	2018	~23,453 изречения	Текст, аудио, видео	Мнение и емоции Видео реч
MELD (Multimodal EmotionLines Dataset) [8]	2019	~13,000 изказвания	Текст, аудио, видео	Емоции и чувства Диалози (ТВ предаване "Приятелите")
Multi-ZOL [9]	~2019	5,288 отзиви	Текст, изображения	Мнения за аспект (1–10) (6 аспекта) Отзиви за телефони
TumEmo [10]	~2020/2021	~190,000 примера	Текст, изображения	Класификация на емоциите, Tumblr данни (изображение и надпис)
MASAD (Multimodal Aspect-Based Sentiment Analysis dataset) [11]	2021	~38,000 двойки изображение–текст	Текст и изображения	Мултимодален аспектно-базиран сентимент анализ с 57 категории на аспекти. Множество домейни
CH-SIMS v2.0 [12]	2022	~4,406 маркирани сегменти	Текст, аудио, видео	Мултимодален сентимент анализ, с допълнителни необозначени данни (~10,161) за невербалните знаци. Видеоклипове за подчертаване на акустичните/визуалните знаци.
MuSe 2022 (Passau-SFCH, Hume-Reaction, Ulm-TSST) [13]	2022	Множество набори от данни	Текст, аудио, видео	Множество задачи: откриване на хумор, емоционални реакции, стрес,

				спонтанно поведение, емоции/чувства.
MSED (Multi-modal Sentiment & Emotion Desire dataset) [14]	~2022	9,190 двойки текст-изображение от социалните медии	Текст и изображения	Идентификация на емоции, желания, мнения. Социални медии
M-EMD/ MEMD-ABSA [15]	2023	~20 000 изречения от отзиви; ~30 000 анотирани с явни и имплицитни аспекти/мнения в 5 области.	Текст	Извличане на мнения за аспекти, разглежда явни и имплицитни аспекти. Сравняване на приноса на модалностите.
Multimodal Dataset for Sentiment Analysis and Classification [16]	2024	600 видеоклипа от различни източници; етикетирани за 7 емоции.	Текст, видео, аудио	Класификация на емоциите, изразени чрез видео, английски език
UniC (Dataset for emotion analysis of videos with unimodal and multimodal labels [17])	2025	Данни за емоции както с едномодални, така и с мултимодални етикети.	Текст, аудио, видео	Моделиране на емоции и чувства, поддържа едномодални/мултимодални етикети. YouTube данни, сравнение между мултимодалният анализ и едномодалния анализ.
EmotionTalk [18]	2025	19 250 изказвания (~23,6 часа реч))	Текст, аудио, видео	7 категории емоции, китайски език, богата анотация. Полезно за диалог/емоция, sentiment, стил.
EmoSign [19]	2025	200 видеоклипа (жестомимичен език) с етикети за сантименти и емоции; плюс анотации на емоционални знаци.	Текст, видео	Класификация на чувства и емоции, разбиране в контекста на жестомимичен език, информация за лицето и жестовете.
M-ABSA (Multilingual ABSA) [20]	2025	Голям набор от данни, обхващащ 7 домейна, 21 езика	Текст	Анализ на мнения на ниво аспект. Фокус върху извличането на аспектен термин, категория и полярност. Многоезичен анализ на мнения на ниво аспект.

Новите набори от данни се отличават с богати, нюансирани анотации като адресират недостатъчно проучени аспекти на sentiment анализа: невербалните знаци; липсващи модалности на повече езици (вкл. жестомимични езици); как чувствата и емоциите взаимодействат с желанията или стила на речта. Използването на такива публично достъпни набори от данни за sentiment анализ позволяват на изследователите да подобряват методите на анализ и да сравняват модели за класификация.

Основните предизвикателствата пред системите за мултимодален sentiment анализ са свързани със събиране на данни от източници и с ефективно им интегриране, с цел да улавят сложните връзки, които определят човешкото изразяване. Настоящите изследвания в областта [21] [22] показват, че все още данните от текстовите източници заемат централно място в системите за мултимодален sentiment анализ, а аудио и визуалните данни се използват за разрешаване на неясноти или добавяне на емоционална дълбочина в определянето на мнението. Затова съвременната насока на научните изследвания е да се предлагат архитектури, които използват данните от трите вида източници, като всеки източник допринася еднакво с уникална информация, която подобрява цялостното разбиране на емоционалното състояние или изразеното мнение.

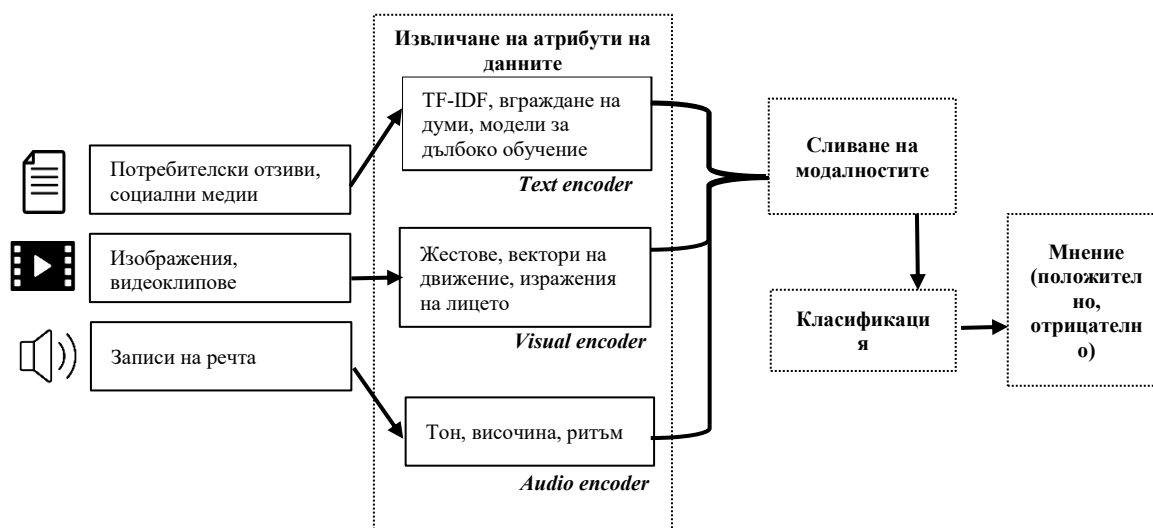
Архитектурни подходи за мултимодален sentiment анализ

Процесът на мултимодален sentiment анализ включва няколко взаимосвързани подзадачи, които допринасят за цялостното разбиране на настроението/мненията в мултимодалното съдържание, позволявайки специализирана обработка и обучение.

- Първата подзадача е *извличане на аспекти от различните типове данни*. Основната цел е точно да се идентифицират и извлекат всички релевантни аспекти термини, независимо дали са изрично посочени или имплицитно загатнати от текстовите данни. Релевантните аспекти термини са конкретните обекти и техните характеристики, за които се изразява чувство или мнение.
- Следващата подзадача е *мултимодалната аспектно-ориентирана класификация* на мнението, която се определя полярността на мнението (напр. положителна, отрицателна, неутрална), свързана с всеки идентифициран аспект термин, като използва информация, получена както от текстовите, така и от другите два източника. Важно е да се отбележи, че мнение към един аспект, може да се различава значително от цялостното мнение, изразено в публикацията или документ.

При мултимодалния sentiment анализ, извличането на аспекти термини и класификацията на мненията към тях са интегрирани в единна рамка и обхваща извличане на двойки аспект-мнение, от които моделите научават и използват присъщите зависимости между тях. Този многозадачен подход води до стабилни и последователни модели и всеобхватна и точна система за sentiment анализ.

Архитектурата на система за мултимодален sentiment анализ интегрира данни от три източника (текст, аудио и изображения) и обхваща пет слоя – извличане на данни от източниците, обработка на данните чрез извличане на атрибути от данните, кодиране на данните, сливане на данните и класификация с цел определяне на чувството/мнението (фиг. 1).



Фигура 1: Архитектура на система за мултимодален sentiment анализ

Входните данни са сурови, необработени текстови данни, изображения и гласови записи. Текстовите данни предоставят субективна информация като може да съдържат споменати аспекти на обекта. Изображенията (свързани снимки или продуктови изображения) или поредица от видео кадри, които съдържат визуални доказателства за обекта или аспекта на обекта (например снимка на камера или екран) или изражения на лицето, движения на главата и жестове на тялото. Гласови записи (например отзиви на клиенти, подкасти или разговори) трябва да включват емоционални характеристики (тон, интонация, темпо на речта).

В слоя *Извличане на атрибути на данните*, суровите данни се изчистват, подготвят (разпознаване на обекти в изображенията, извличане на аудио характеристики) и трансформират в смислени структури (вектори), които улавят значението на текста, визуалните елементи и емоциите в гласа. Добрата обработка и трансформация на данните осигурява ефективност и високо качество на моделите за класификация на чувствата/мненията. В процеса на извличане на атрибути на данните се прилагат енкодери, които са специализирани инструменти, предназначени да превеждат разбираеми за човека данни (език, звук, картина) в машинно разбираеми числови представяния, които се използват от приложения с изкуствен интелект.

- Подходите за подготовка и трансформация на *текстовите данни* могат да бъдат традиционни като Bag-of-Words, TF-IDF, използване на речници (VADER¹[23]) или контекстуализирани вграждане на думи от модели за дълбоко обучение като BERT, RoBERTa, които улавят семантично значение и контекст. Входното изречение се разделя на по-малки единици (токени). Всеки токен се преобразува в предварително обучен вектор с помощта на таблица за вграждане. Векторите са проектирани така, че думите със сходни значения да са близо една до друга във векторното пространство. Енкодер моделите обработват последователността от векторите чрез механизъм, наречен Вниманиe (Attention), за да се претегли важността на всяка дума спрямо всички останали. Енкодерът произвежда един вектор за цялата входна последователност.
- Извличането на атрибути от *видео и изображения* се извършва чрез видео енкодер, който взема поредица от видео кадри (изображения) и придружаващото ги аудио и създава представяне, което улавя както пространствената (в рамките на кадър), така и времевата (през кадрите) информация. Видео енкодерите често комбинират техники от обработка на изображения и текст. Извличането на пространствени характеристики като обекти и сцени от всеки отделен кадър, се извършва чрез конволюционна невронна мрежа (CNN). Извличане на времеви характеристики са важни за разбиране на движението и историята. Последователността от векторите на атрибутите от всеки кадър се подава в модел, който разбира последователностите - повтарящи се невронни мрежи (RNN) и трансформари (с

¹ VADER (Valence Aware Dictionary and sEntiment Reasoner)

механизъм на внимание). Видео трансформърът прилага самовнимание във всички кадри, което му позволява да разбира дългосрочни зависимости и взаимоотношения.

- В извличането на атрибути от *аудио записи* се прилага аудио енкодер, който преобразува необработени аудио сигнали в смислено представяне. Вълната често се преобразува в спектрограма (визуално представяне на спектъра от честоти във времето), след което се извличат характеристики, идентифицират се фонемите и други акустични елементи. Аудио сигналът е последователност, така че модели като RNN или трансформъри могат да го обработват директно, за да уловят времевия контекст, което е от решаващо значение за разбирането на речта и музиката.

Мултимодално сливане, следващият слой в архитектурата, е най-изследваният компонент в системите за мултимодален sentiment анализ. Сливането на модалностите има за цел да комбинира отделните представяния на модалностите [T, A, V] в съвместно представяне Z, което ефективно улавя както интрамодалните, така и кросмодалните взаимодействия [24]. Z е единен векторен изход, който обхваща емоционалната информация от всички участващи модалности (текст, аудио, изображение).

В методологичните подходи преобладават невронни архитектури, които могат да бъдат категоризирани в три поколения: (1) обучение за унимодално представяне с късно сливане, (2) обучение за съвместно представяне с кросмодални взаимодействия и (3) съвременни архитектури за сливане, използващи предварително обучени модели.

В първоначалните подходи в мултимодалното обучение доминира методологичният подход на късното сливане. Всяка модалност (напр. текст, изображение) се обработва от собствена, независима невронна мрежа (напр. CNN за изображения, RNN/LSTM за текст или аудио), които извличат унимодални характеристики. [25]. Сливането на извлечената информация - вектори на характеристиките [T; A; V] - се извършва в последния етап, обикновено чрез конкатениране на получените унимодални представяния, последвано от класификатор. Подходът позволява използването на добре разработени унимодални модели, но не успява да улови сложните корелации и взаимодействия между модалностите (напр. разпознаването на сарказъм), тъй като взаимодействието е минимално и се случва твърде късно в процеса.

Второто поколение архитектури се фокусира върху преодоляването на ограниченията на късното сливане чрез въвеждане на по-ранни и експлицитни механизми за взаимодействие между модалностите. Целта е да се научи съвместно представяне (joint representation) на данните като се използват трансформър-базирани модели (като BERT) и механизми за кръстосано внимание (Cross-Attention). Механизъм на кръстосаното внимание позволява на една модалност (напр. текст) да се обвърже с най-релевантните характеристики от друга модалност (напр. изображение) и обратно, създавайки силни кросмодални връзки между тях [26]. Например, думата "щастлив" в текста може да се обвърже с усмихнато лице във визуалния поток и развълнуван тон в аудио потока. При този модулен подход всяка модалност има свой енкодер, но между тях се вмъкват слоеве за експлицитно сливане и взаимодействие [27].

Най-съвременните мултимодални архитектури се характеризират с мащабно предварително обучение върху огромни масиви от мултимодални данни и стремеж към унифицирано представяне и целево сливане. Използват се модели, базирани на трансформър архитектури, обучени по методи на самоконтролирано обучение върху голям обем двойки данни (напр. изображение-текст). Тези архитектури често постигат ранно сливане като превръщат различните модалности в единен поток от токени още на входа на модела, позволявайки на един трансформър енкодер да обработва всички модалности съвместно. Мултимодални големи езикови модели използват предварително обучени големи езикови модели (LLMs) и енкодери за изображения/аудио, свързани чрез интерфейс за модалности. Те имат способността да приемат, да разсъждават и да извеждат мултимодална информация, което ги прави подходящи за сложен емоционален анализ, включващ разсъждение. Най-новите архитектурни подходи се насочват към пълна интеграция и унификация, използвайки предимствата на предварително обучените модели и сложните механизми за внимание, за да постигнат разбиране, което е близко до човешкото възприятие.

След мултимодалното сливане, класификацията на мнение/чувство се извършва, като обединеното представяне Z на всички модалности се подава към слой за класификация. При късното сливане векторът е агрегирана форма на индивидуалните изходи от класификаторите на всяка модалност. При ранното сливане векторът е получен директно от последния скрит слой на общия трансформър модел или след слоя за кросмодално внимание. Обединеният векторен изход Z се подава на един или повече напълно свързани слоя, които изпълняват задачата по класификация. По време на обучението, класификационният резултат се сравнява с реалния клас, използвайки функция на загуба (Loss Function), която ръководи актуализирането на теглата на цялата мултимодална архитектура. Чрез минимизиране на тази загуба, моделът се научава да създава обединени представяния, които най-добре разграничават класовете, като същевременно използват допълваща информация от всички модалности.

Предизвикателства пред мултимодалния сентимент анализ

Въпреки значителния прогрес в архитектурните и методологични подходи, извличането на мнения от различни типове източници среща предизвикателства, които произтичат от присъщата сложност на човешкия език, разнообразната природа на мултимодалните данни и субективната същност на чувствата.

Едно от основните предизвикателства е семантичната сложност. Изреченията често съдържат множество аспекти, всеки от които потенциално се отнася до различни обекти или понятия, както в текста, така и в придружаващо изображение. Освен това тези различни аспекти и съответните им области на изображението могат да носят различни чувства. Точното улавяне на нюансирани чувства в една двойка текст-изображение изисква моделът да разграничава и определя чувствата на детайлно ниво, откривайки сложните семантични връзки.

Разбирането на чувствата и контекстуалните вариации се моделира трудно поради двусмисленост на езика, при която една дума може да притежава множество значения в зависимост от контекста си. Двусмислеността води до напълно различни интерпретации на чувствата. Културните и контекстуалните различия също значително влияят върху начина, по който настроенятията се изразяват и интерпретират, което изисква алгоритмите да бъдат обучени на различни набори от данни, за да се преодолеят тези вариации.

Друг проблем, който често съпътства данните е свързан с несъответствие между текста и изображението, придружаващо текста. Например изображението не съдържа обектите/аспектите, посочени в текста. Такова несъответствие може сериозно да влоши качеството на модела, като доведе до неточно крос-модално подравняване. Затова съвременните модели са насочвани към идентифициране и директно справяне със специфични, често срещани несъответствия в мултимодалните данни от реалния свят.

Променливостта на данните, субективността и пристрастията създават значителни препятствия пред обучението на стабилни модели, които изискват разнообразни и представителни набори от данни. Процесът на аотиране на данни включва лична преценка, която може да доведе до субективност и пристрастия. В тази връзка е необходимо установяване на ясни правила за аотация и стратегии за управление на пристрастията с цел да се гарантира създаването на надеждни аотирани набори от данни.

Заклучение

Въпреки значителния напредък, постигнат чрез дълбокото обучение и трансформър-базираните архитектури, мултимодалният сентимент анализ все още се сблъсква с важни предизвикателства, свързани с интеграция на данните от различните източници, синхронизацията, интерпретируемостта и недостига на аотирани данни. Съществуващите модели често показват ограничена способност за обобщаване между различни домейни, езици и културни контексти поради пристрастия и малки набори от данни. Преодоляването на тези проблеми изисква разработването на прозрачни и адаптивни мултимодални архитектури,

способни да работят с непълни или данни със шумове и да предоставят интерпретируеми резултати.

Бъдещите изследвания е необходимо да бъдат фокусирани върху разширяване на многоезичните ресурси, нискоресурсното обучение и етичната употреба на данни, за да се постигнат по-надеждни, обясними и контекстуално чувствителни системи за анализ на мнения.

Литература / References

1. Ankita Gandhi, Kinjal Adhvaryu, Soujanya Poria, Erik Cambria, Amir Hussain (2023). Multimodal sentiment analysis: A systematic review of history, datasets, multimodal fusion methods, applications, challenges and future directions, *Information Fusion*, Volume 91, 2023, Pages 424-444, ISSN 1566-2535, <https://doi.org/10.1016/j.inffus.2022.09.025>
2. Liu, B. (2012). *Sentiment Analysis and Opinion Mining*. United States: Morgan & Claypool.
3. Yang, H., Zhao, Y., Wu, Y., Wang, S., Zheng, T., Zhang, H., ... & Qin, B. (2024). Large language models meet text-centric multimodal sentiment analysis: A survey. *arXiv preprint arXiv:2406.08068*.
4. Pérez-Rosas, V., Mihalcea, R., & Narayan, S. (2013). *Utterance-level Multimodal Sentiment Analysis* (MOUD dataset). (paper & dataset page), <https://aclanthology.org/P13-1096.pdf>
5. Yu, J., & Jiang, J. (2019). Adapting BERT for Target-Oriented Multimodal Sentiment Classification. *International Joint Conference on Artificial Intelligence*, pp. 5408-5414
6. Zadeh, A., Zellers, R., Pincus, E., & Morency, L. P. (2016). Mosi: multimodal corpus of sentiment intensity and subjectivity analysis in online opinion videos. *arXiv preprint arXiv:1606.06259*.
7. AmirAli Bagher Zadeh, Paul Pu Liang, Soujanya Poria, Erik Cambria, and Louis-Philippe Morency. 2018. [Multimodal Language Analysis in the Wild: CMU-MOSEI Dataset and Interpretable Dynamic Fusion Graph](#). In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 2236–2246, Melbourne, Australia. Association for Computational Linguistics.
8. Soujanya Poria, Devamanyu Hazarika, Navonil Majumder, Gautam Naik, Erik Cambria, and Rada Mihalcea. 2019. [MELD: A Multimodal Multi-Party Dataset for Emotion Recognition in Conversations](#). In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 527–536, Florence, Italy. Association for Computational Linguistics.
9. Xu, N., Mao, W., & Chen, G. (2019). Multi-Interactive Memory Network for Aspect Based Multimodal Sentiment Analysis. *Proceedings of the AAAI Conference on Artificial Intelligence*, 33(01), 371-378. <https://doi.org/10.1609/aaai.v33i01.3301371>
10. Xiaocui, Yang & Feng, Shi & Wang, Daling & Zhang, Yifei. (2020). Image-Text Multimodal Emotion Classification via Multi-View Attentional Network. *IEEE Transactions on Multimedia*. PP. 1-1. 10.1109/TMM.2020.3035277.
11. Jie Zhou, Jiabao Zhao, Jimmy Xiangji Huang, Qinmin Vivian Hu, Liang He, MASAD: A large-scale dataset for multimodal aspect-based sentiment analysis, *Neurocomputing*, Volume 455, 2021, Pages 47-58, ISSN 0925-2312, <https://doi.org/10.1016/j.neucom.2021.05.040>.
12. Liu, Y., Yuan, Z., Mao, H., Liang, Z., Yang, W., Qiu, Y., ... & Gao, K. (2022, November). Make acoustic and visual cues matter: Ch-sims v2. 0 dataset and av-mixup consistent module. In *Proceedings of the 2022 international conference on multimodal interaction* (pp. 247-258).

13. Christ, L., Amiriparian, S., Baird, A., Tzirakis, P., Kathan, A., Müller, N., ... & Schuller, B. W. (2022, October). The muse 2022 multimodal sentiment analysis challenge: humor, emotional reactions, and stress. In *Proceedings of the 3rd International on Multimodal Sentiment Analysis Workshop and Challenge* (pp. 5-14).
14. Ao Jia, Yu He, Yazhou Zhang, Sagar Uprety, Dawei Song, and Christina Lioma. 2022. [Beyond Emotion: A Multi-Modal Dataset for Human Desire Understanding](#). In *Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 1512–1522, Seattle, United States. Association for Computational Linguistics.
15. Cai, H., Song, N., Wang, Z., Xie, Q., Zhao, Q., Li, K., ... & Xia, R. (2025). Memd-absa: A multi-element multi-domain dataset for aspect-based sentiment analysis. *Language Resources and Evaluation*, 1-29.
16. MALI, SHANKAR (2024), “Multimodal Dataset for Sentiment Analysis and Classification”, Mendeley Data, V2, doi: 10.17632/b7z68nykxt.2
17. Du, Q., Labat, S., Demeester, T. *et al.* UniC: a dataset for emotion analysis of videos with multimodal and unimodal labels. (2025) *Lang Resources & Evaluation* **59**, 2857–2892. <https://doi.org/10.1007/s10579-025-09837-0>
18. Sun, H., Wang, X., Zhao, J., Zhao, S., Zhou, J., Wang, H., ... & Qin, Y. (2025). EmotionTalk: An Interactive Chinese Multimodal Emotion Dataset With Rich Annotations. *arXiv preprint arXiv:2505.23018*.
19. Chua, P., Fang, C. M., Ohkawa, T., Kushalnagar, R., Nanayakkara, S., & Maes, P. (2025). EmoSign: A Multimodal Dataset for Understanding Emotions in American Sign Language. *arXiv preprint arXiv:2505.17090*.
20. Wu, C., Ma, B., Liu, Y., Zhang, Z., Deng, N., Li, Y., ... & Plank, B. (2025). M-ABSA: A Multilingual Dataset for Aspect-Based Sentiment Analysis. *arXiv preprint arXiv:2502.11824*.
21. Zhang, H. (2024). A comprehensive survey on multimodal sentiment analysis: Techniques, models, and applications. *Advances in Engineering Innovation*, 12, 47-52. <https://www.ewadirect.com/journal/aei/article/view/16349#>
22. Tianyu Zhao, Ling-ang Meng, Dawei Song (2024). Multimodal Aspect-Based Sentiment Analysis: A survey of tasks, methods, challenges and future directions, *Information Fusion*, Volume 112, 2024, 102552, ISSN 1566-2535, <https://doi.org/10.1016/j.inffus.2024.102552>
23. Barik, K., Misra, S. Analysis of customer reviews with an improved VADER lexicon classifier. *J Big Data* 11, 10 (2024). <https://doi.org/10.1186/s40537-023-00861-x>
24. Lai, S., Hu, X., Xu, H., Ren, Z., & Liu, Z. (2023). Multimodal sentiment analysis: A survey. *Displays*, 80, 102563. <https://arxiv.org/abs/2305.07611>
25. Jing Gao, Peng Li, Zhikui Chen, Jianing Zhang; A Survey on Deep Learning for Multimodal Data Fusion. *Neural Comput* 2020; 32 (5): 829–864. doi: https://doi.org/10.1162/neco_a_01273
26. Mowbray, T. (2025). A Survey of Deep Learning Architectures in Modern Machine Learning Systems: From CNNs to Transformers. <https://doi.org/10.5281/zenodo.17035434>
27. Wadekar, S. N., Chaurasia, A., Chadha, A., & Culurciello, E. (2024). The evolution of multimodal model architectures. *arXiv preprint arXiv:2405.17927*.

**ИНОВАТИВНИ ИНФОРМАЦИОННИ ТЕХНОЛОГИИ
ЗА ДИГИТАЛИЗАЦИЯ НА ИКОНОМИКАТА**

Публикацията съдържа резултати от научен форум, съфинансиран със средства
от целева субсидия за НИД на УНСС по договор № НИД НФ – 29/2025

Сборник с доклади

Колектив

Дизайн на корицата: Емилия Лозанова

Даден за печат: 5.12.2025 г.

Формат 8/60/84

ISSN 3033-0432 (print)

ISSN 3033-0467 (online)

ПЕПАЧНИЦА УНСС