

DOI: <https://doi.org/10.36910/6775-2524-0560-2026-64-02>

УДК 004.5:004.89:069

Крамар Тарас Олександрович

<https://orcid.org/0000-0001-8060-0169>

Дуда Олексій Михайлович, к.т.н, доцент

<https://orcid.org/0000-0003-2007-1271>

Тернопільський національний технічний університет, м. Тернопіль, Україна

ВИКОРИСТАННЯ МУЛЬТИМОДАЛЬНИХ LLM ДЛЯ АВТОМАТИЗАЦІЇ ПРОЦЕСІВ ОЦИФРУВАННЯ ДАНИХ ЩОДО ОБ'ЄКТІВ КУЛЬТУРНОЇ СПАДЩИНИ

Крамар Т. О., Дуда О. М. Використання мультимодальних LLM для автоматизації процесів оцифрування даних щодо об'єктів культурної спадщини. У статті розглянуто виклики пов'язані з накопиченням неструктурованих описів об'єктів культурної спадщини та недостатньою семантичною інтероперабельністю традиційних систем зберігання метаданих цих об'єктів, що суттєво обмежує доступність культурної спадщини та ефективність процесів її цифровізації. Метою роботи є розробка комплексного підходу до автоматизованого опрацювання неструктурованих музейних архівів із застосуванням мультимодальних великих мовних моделей. Запропонована клієнт-серверна архітектура поєднує технології системного промптингу з класами онтології CIDOC CRM для трансформації статичних описів експонатів у динамічні граfi знань. Для гарантування машинної зчитуваності та усунення синтаксичного шуму під час генерації застосовано програмно-алгоритмічний механізм обмеженого декодування на рівні API. Забезпечення наукової достовірності та фактологічної точності реалізовано через впровадження гібридної взаємодії «людина в процесі», де музейний експерт здійснює перевірку розмічених LLM даних в інтерактивному інтерфейсі вебформи. Завершальний етап роботи програмно-алгоритмічного комплексу використовує автоматичне перетворення структурованих JSON-об'єктів на SPARQL-запити для збереження у графовій базі GraphDB із динамічною генерацією унікальних ідентифікаторів. Даний підхід суттєво зменшує витрати ресурсів на ручну каталогізацію, забезпечує структурну чистоту інформації та відкриває можливості для інтеграції оцифрованих музейних об'єктів до міжнародних агрегаторів культурної спадщини.

Ключові слова: штучний інтелект, машинне навчання, великі мовні моделі, обмежене декодування, граfi знань, онтологія, CIDOC CRM, цифровізація, культурна спадщина.

Kramar T., Duda O. Application of multimodal LLMs for automating the digitization processes of cultural heritage object data. The article addresses the challenges associated with the accumulation of unstructured descriptions of cultural heritage objects and the insufficient semantic interoperability of traditional metadata storage systems for these objects, which significantly limits the accessibility of cultural heritage and the efficiency of its digitization processes. The aim of the work is to develop a comprehensive approach to automated processing of unstructured museum archives using multimodal large language models. The proposed client-server architecture combines system prompting technologies with CIDOC CRM ontology classes to transform static descriptions of exhibits into dynamic knowledge graphs. To guarantee machine readability and eliminate syntactic noise during generation, a software-algorithmic mechanism of limited decoding at the API level was used. Ensuring scientific reliability and factual accuracy is implemented through the implementation of hybrid «human in the process» interaction, where a museum expert performs verification of marked-up LLM data in the interactive interface of the web form. The final stage of the software-algorithmic complex uses automatic conversion of structured JSON objects into SPARQL queries for storage in the GraphDB graph database with dynamic generation of unique identifiers. This approach significantly reduces resource consumption for manual cataloging, ensures structural purity of information, and opens up opportunities for integrating digitized museum objects into international cultural heritage aggregators.

Keywords: artificial intelligence, machine learning, large language models, constrained decoding, knowledge graphs, ontology, CIDOC CRM, digitalization, cultural heritage.

Постановка наукової проблеми. Установи, що займаються збереженням культурної спадщини, зокрема музеї, стикаються з масштабною кризою накопичення необроблених і неоцифрованих колекцій, які через брак високоефективних інформаційно-технологічних інструментів перетворюються на практично недоступні неоцифровані дані [1]. Традиційне ручне перенесення інформації з неструктурованих історичних документів (наприклад, інвентарних карток) є доволі трудомістким процесом. Каталогізація одного об'єкта культурної спадщини може займати до години, що ускладнює масштабування інформаційних систем та підвищує ризик людських помилок. Цю проблему поглиблює серйозний недолік семантичної інтероперабельності, адже застарілі моделі метаданих фокусуються на фізичних носіях і різних стандартах подання описових даних [1]. Без впровадження сучасних уніфікованих стандартів та онтологій цифрові колекції залишаються розпорошеними, обмежуючи доступність культурної спадщини загалом [2]. Для подолання цих структурних бар'єрів виникає потреба автоматизації процесів за допомогою штучного інтелекту (ШІ) та великих мовних моделей (LLM). Сучасні інтелектуальні технології здатні автоматично аналізувати неструктуровані тексти, видобувати з них іменовані інформаційні сутності та контекстуальні відношення, перетворюючи статичні записи на динамічні граfi знань [1]. Використання ШІ в поєднанні з онтологіями верхнього рівня, як-от CIDOC-CRM, значно

скорочує час обробки, нівелює недоліки ручної каталогізації, та гарантує семантичну сумісність процесів глобальної інтеграції даних [2]. Такий інноваційний технологічний перехід є критично важливим, адже дає можливість оперативно опрацювати музейні колекції повноцінно розкриваючи дослідницький потенціал «великих даних минулого» [3].

Аналіз останніх досліджень і публікацій. У [4] описано оригінальний підхід до моделювання бібліотечних метаданих на основі взаємодії людини та великих мовних моделей, що спрямований на концептуальне впорядкування та усунення складних репрезентативних переплетінь. Робота [1] зосереджена на конвеєрі «людина в процесі» (human-in-the-loop) для перетворення неструктурованих історичних архівів у динамічні граfi знань, де роль традиційного архіваріуса змінюється від каталогізатора, що вводить дані вручну, до стратегічного «семантичного редактора», який перевіряє результати роботи ШІ для забезпечення достовірності даних.

Еволюція технологій розпізнавання історичних текстів демонструє перехід від традиційних поетапних методів до наскрізної мультимодальної обробки. Ранні системи оптичного (OCR) та рукописного (HTR) розпізнавання спиралися на гібридні архітектури згорткових і рекурентних неймереж, які вимагали значних обсягів розмічених даних, були чутливими до інформаційного «шуму» старовинних документів та потребували складної попередньої сегментації сторінок [3], [5]. Натомість сучасний етап характеризується застосуванням мультимодальних мовних моделей, як-от GPT-4o чи Gemini, що здатні безпосередньо аналізувати зображення документів і з високою точністю генерувати транскрипції без необхідності попереднього вирівнювання тексту чи ручного сегментування (див. Рис. 1).

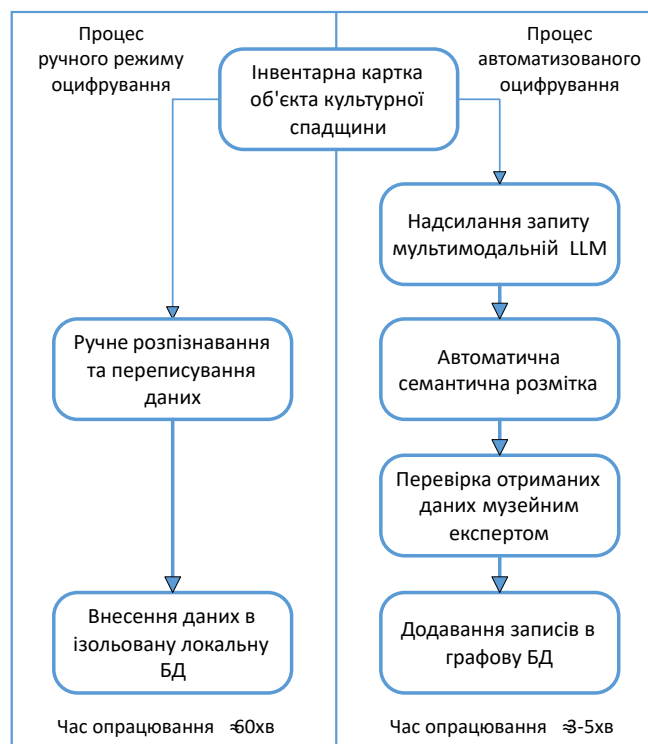


Рис. 1 – Процеси оцифрування інвентарних карток об'єктів культурної спадщини

Для ефективного використання мовних моделей критичним є механізм керування їхньою поведінкою через системні промпти (запити) та інженерію промптів. Дослідники підкреслюють, що застосування спеціалізованих шаблонів промптингу дає змогу структурувати взаємодію та суттєво підвищувати релевантність відповідей [6]. Серед таких методів особливе місце посідає шаблон призначення ролі моделі (single-character role play), коли системний промпт жорстко задає віртуальному агенту певну ідентичність, стиль, тон та межі знань, наприклад, експерта-куратора або дослідника об'єктів культурної спадщини. Це фокусує процес генерації контенту та дає можливість створювати коректні й контекстно-орієнтовані процеси обміну даними [7], [8].

Щоб гарантувати безпечну інтеграцію згенерованих даних в інформаційні системи, застосовуються технології структурованого виводу та обмеженого декодування. Ці програмно-

алгоритмічні засоби примусово формують словник допустимих токенів під час генерації даних засобами мультимодальної великої мовної моделі, гарантуючи сувору відповідність фінального результату заданому формату (наприклад, JSON) без додавання стороннього синтаксичного інформаційного «шуму» [9], [10]. Хоча цей підхід надійно запобігає синтаксичним збоєм, він не усуває «галюцинацій формату» – феномену, коли модель продукує структурно ідеальний вивід, але заповнює його фактологічно хибними значеннями, що робить необхідною додаткову верифікацію коректності даних на рівні окремих полів [11].

У сфері побудови графів знань використання мовних моделей суттєво оптимізує автоматичне мапування неструктурованих текстів на стандартизовані класи онтології CIDOC CRM. Передові методології, пропонують ітеративні конвеєри, у яких великі мовні моделі розпізнають факти, гіпотези та доказову базу з наукового дискурсу і вирівнюють їх за концептуальними шаблонами [12]. Такі автоматизовані підходи дають змогу розпізнавати іменовані сутності та перетворювати їх на RDF-триплети із строгим зіставленням класів (наприклад, E21 Person, E53 Place) та властивостей, що забезпечує глобальну семантичну інтероперабельність [13].

Регуляторний та етичний контекст впровадження штучного інтелекту в музейну практику активно формується на рівні міжнародних інституцій ЮНЕСКО [14] та Europeana [15]. Їхні рекомендації зосереджуються на врегулюванні авторських прав щодо використання оцифрованого надбання для тренування ШІ, забезпеченні прозорості алгоритмів, інклюзивності та запобіганні культурним упередженням. Також наголошується, що цифрові інновації мають посилювати демократизацію доступу до культури без компромісів із достовірністю та етичними принципами збереження спадщини [16].

Метою дослідження є розробка комплексного підходу до автоматизованого опрацювання та семантичного структурування інформації з неструктурованих музейних архівів на основі застосування мультимодальних великих мовних моделей. Зокрема, акцент зроблено на використанні технологій системного промптингу та обмеженого декодування для точно розмічених даних за класами онтології CIDOC CRM, впровадженні гібридної взаємодії «людина в процесі» для забезпечення достовірності метаданих, а також зручності використання вебінтерфейсу для експертної перевірки та збереження інформації працівниками музею.

Виклад основного матеріалу. На основі проведеного аналізу наукових джерел виділимо узагальнені, незалежні від конкретних технологій принципи побудови систем для обробки музейних інвентарних карток та інтеграції ШІ-сервісів.

Принцип декомпозиції. У сучасних інформаційних системах принцип декомпозиції реалізується через розподіл складного неструктурованого завдання на дрібніші, незалежні та легко керовані сегменти [17]. Замість використання єдиної монолітної інформаційної системи, обчислювальна архітектура будується як модульний послідовний конвеєр (наприклад, розпізнавання/оцифрування тексту, видобування іменованих інформаційних сутностей, семантичне зв'язування та генерування метаданих), що дає змогу здійснювати поетапну обробку, ізолювати підзадачі та звести до мінімуму ризик поширення помилок між програмними модулями [12]. Такий покроковий підхід застосовується у складних системах обробки запитів, де процеси розбиваються на ізольовані рівні структурного, семантичного та прагматичного аналізу [13].

Принцип розділення відповідальностей. Фундаментальним принципом побудови обчислювальної архітектури є чітке розмежування між шарами зберігання даних, ШІ-обробки та подання. Наприклад, концептуальні моделі застосовують суворе відокремлення сервісів роботи з даними (бекенду, баз знань та точок доступу до них) від користувацьких інтерфейсів (фронтенду), що дає можливість незалежно оновлювати й масштабувати кожен шар [1]. Для безпечної роботи інноваційних компонентів необхідне впровадження проміжного програмного забезпечення (middleware), яке ізолює ШІ-модулі від застарілих музейних інформаційних інфраструктур і традиційних систем управління колекціями. Це гарантує інтероперабельність, захист первинних даних та збереження суворого кураторського контролю [7].

Принцип оркестрації процесів забезпечує взаємодію програмно-алгоритмічних агентів та людини. Оркестрація процесів у таких системах забезпечує динамічну координацію між розрізненими компонентами через мультиагентні підходи та гібридні робочі процеси:

- **Мультиагентна оркестрація:** впроваджується програмна архітектура, де кожен спеціалізований віртуальний програмно-алгоритмічний агент відповідає за конкретний етап обробки даних (наприклад, один програмно-алгоритмічний агент аналізує об'єкти, інший формує

описи), діючи за принципом «домінантності» над своєю ділянкою для підвищення загальної узгодженості системи [17].

- *Конвеєрна оркестрація та виклик інструментів*: ШІ-компонент взаємодії звертається до зовнішніх баз даних (наприклад, через архітектуру доповненої генерації) для фактологічного обґрунтування результатів на основі перевірених музейних джерел [7].

- *Гібридна оркестрація «людина в процесі» (Human-in-the-Loop)*: Критичним принципом інформаційних систем для оцифрування культурної спадщини є вбудовування механізмів ручної перевірки в конвеєр обробки даних; у цій моделі автоматизовані статистичні обчислення ШІ жорстко контролюються людиною-експертом (семантичним редактором), яка встановлює когнітивні межі системи та керує процесом ітеративного навчання для гарантування абсолютної наукової достовірності продуктованих ШІ даних [1].

Покрокова архітектура системного промту.

Крок 1. Призначення ролі (Role Assignment). Процес починається із задання жорсткого рольового контексту (див. лістинг 1).

Лістинг 1 – Опис рольового контексту системного запиту до LLM

prompt = ""

Ти експерт з музейної справи, архіваріус та спеціаліст із семантичного вебу.

Проаналізуй наданий текст та/або зображення інвентарної картки експоната.

Твоє завдання – витягнути сутності і розподілити їх за 31 класом онтології CIDOC CRM.

Правила заповнення:

- Відповідай виключно українською мовою.
- Уважно зіставляй знайдені факти з відповідними полями (наприклад, техніку – у designOrProcedure, дату – у timeSpan, особу – у person тощо).
- Якщо інформації для певного класу немає в тексті, ОБОВ'ЯЗКОВО поверни порожній рядок "".
- Не вигадуй дані та не пиши "немає".
- В поле 'description' запиши короткий загальний опис об'єкта.

Це дає можливість сфокусувати генерацію, встановити професійний тон, визначити межі знань віртуального агента та адаптувати його поведінку до специфіки конкретного завдання [7]. Використання патерну персони суттєво підвищує релевантність відповідей та забезпечує дотримання стилістичних настанов [6].

Крок 2. Інструкція вилучення сутностей (Task Instruction / Contextual constraints). Інструкція використовується для надання чітких вказівок щодо того, яку саме інформацію необхідно проаналізувати у вхідних даних (input data). На цьому етапі великій мовній моделі надаються концептуальні схеми або цільові списки сутностей та відношень (наприклад, перелік класів згідно зі стандартом CIDOC CRM: E21 Person, E53 Place), які слугують суворим семантичним орієнтиром для аналізу тексту [2]. До цієї інструкції обов'язково додається сам вхідний текст для обробки (див. лістинг 2) [6].

Лістинг 2 – Приклад вхідних даних для LLM.

ПРИКЛАД ДАНИХ ДЛЯ ОПРАЦЮВАННЯ (Few-Shot):

Вхідний текст: "Скульптура святого Миколая. Інвентарний номер: ДІАЗ КН-1397 СК-4. Матеріали: дерево, левкас, поліхромія. Висота: 120 см. Стан: Задовільний, наявні втрати левкасу. Автор: Іоан Георгій Пінзель. Створено: середина 18 століття, м. Збараж (Пізнє бароко). Зберігач: Національний заповідник 'Замки Тернопілля'. Документ: Інвентарна книга №1. Статус: Договір страхування №112/24. Історія: Акт передачі об'єкта до фондів; Тимчасова передача до Львівської галереї; Виставка 'Скарби Галичини', 2024. Реставраційний огляд, 2023. Покриття консерваційним розчином виконав Національний реставраційний центр. Оцифрування скульптури здійснив Крамар Т.О. (Photogrammetry and 3D Gaussian Splatting). 3D-модель святого Миколая (Creative Commons Attribution 4.0). Архівна фотографія 1970-х років. Опис: Дерев'яна поліхромна скульптура..."

Крок 3. Приклади для навчання в контексті (In-Context Learning / Few-shot demonstrations). Інструкція використовується для запровадження до промпту декількох ретельно підібраних прикладів правильного виконання завдання (пар «вхідні дані – очікуваний результат»). Навчання на декількох прикладах (few-shot prompting) скеровує велику мовну модель, покращує точність зіставлення сутностей та допомагає їй «зрозуміти» складну структуру і очікувану логіку генерації (див. лістинг 3) [8]. Використання динамічно підібраних прикладів із баз знань або глосаріїв додатково збагачує промпт перевіреним доменним контекстом [18].

Лістинг 3 – Очікувана логіка генерації відповіді LLM

```
Очікуваний результат (всі поля обов'язкові, відсутні заповнюються як
""):
{
  "appellation": "",
  "identifier": "ДІАЗ КН-1397 СК-4 (ДІАЗ-12345)",
  "title": "Скульптура святого Миколая",
  "type": "Скульптура",
  "humanMadeObject": "Скульптура святого Миколая",
  "curatedHolding": "Колекція фондів Національного заповідника",
  "material": "Дерево, левкас, поліхромія",
  "dimension": "Висота: 120 см",
  "conditionState": "Задовільний, наявні втрати левкасу",
  "conditionAssessment": "Реставраційний огляд, 2023",
  "production": "Процес створення скульптури Миколая",
  "creation": "Оцифрування скульптури",
  "designOrProcedure": "Photogrammetry and 3D Gaussian Splatting",
  "timeSpan": "Середині 18 століття",
  "place": "м. Збараж",
  "period": "Пізнє бароко",
  "actor": "Національний реставраційний центр",
  "person": "Іоанн Георгій Пінзель",
  "group": "Національний заповідник 'Замки Тернопілля'",
  "informationObject": "3D-модель святого Миколая",
  "visualItem": "Архівна фотографія 1970-х років",
  "linguisticObject": "Дерев'яна поліхромна скульптура святого
Миколая...",
  "language": "Українська",
  "document": "Інвентарна книга №1",
  "legalObject": "Договір страхування №112/24",
  "right": "Creative Commons Attribution 4.0",
  "acquisition": "Акт передачі об'єкта до фондів",
  "transferOfCustody": "Тимчасова передача до Львівської галереї",
  "event": "Виставка 'Скарби Галичини', 2024",
  "activity": "Покриття консерваційним розчином",
  "attributeAssignment": "Запис у базі Data Mesh",
  "description": "Дерев'яна поліхромна скульптура святого
Миколая..."
}
```

Крок 4. Обмеження формату виводу (Output Format Indicator). Фінальна текстова директива у промпті, що вказує моделі на необхідність специфічного форматування результату. Сюди належать інструкції на кшталт «Поверни результат у форматі JSON із ключами 'title' та 'abstract'» [19]. Цей крок необхідний для базового структурування кінцевої відповіді, щоб забезпечити її сумісність із подальшими етапами машинної обробки [6]. Замість виключно текстових інструкцій, системі передається жорсткий шаблон даних, де модель примусово переводиться у режим генерації структурованих даних (див. лістинг 4).

Лістинг 4 – Встановлення шаблону виводу даних

```
response = model.generate_content(
    content_parts,
```

```

generation_config=genai.GenerationConfig(
    response_mime_type="application/json",
    response_schema=ArtifactSchema,
    temperature=0.1,
)
)

```

Різниця між текстовою інструкцією та обмеженим декодуванням на рівні API. Текстова інструкція у промпті (наприклад, «поверни лише JSON») спирається виключно на статистичну здатність великої мовної моделі дотримуватися вказівок та імітувати задану структуру на семантичному рівні. Проте великі мовні моделі від природи схильні до генерування розмовного тексту, тому навіть за наявності чіткої інструкції вони часто додають зайвий синтаксис, пояснення (так звані «вільні радикали») або порушують загальну схему. Це робить непередбачуваними результати роботи LLM і змушує розробників писати складний додатковий код для парсингу та виправлення помилок формату [10].

Натомість *обмежене декодування (constrained decoding)* або ж використання схем відповідей на рівні API діє не на семантичному рівні промпту, а на глибокому технічному етапі генерації токенів (inference stage) безпосередньо у самій моделі [8].

Технічний механізм цього підходу полягає у тому, що під час роботи великої мовної моделі використовуються парсери, скінченні автомати або регулярні вирази, які динамічно обчислюють «маску токенів» на кожному окремому кроці продукування результату [9]. Ця маска накладається на розподіл ймовірностей наступного токена і примусово відфільтровує (виключає зі словника моделі) будь-які варіанти, що порушують задану граматику чи цільову структуру. Таким чином, в моделі є можливість обирати наступний токен виключно з тих варіантів, які гарантовано є синтаксично правильними для цільового результату (див. Рис. 2) [10].

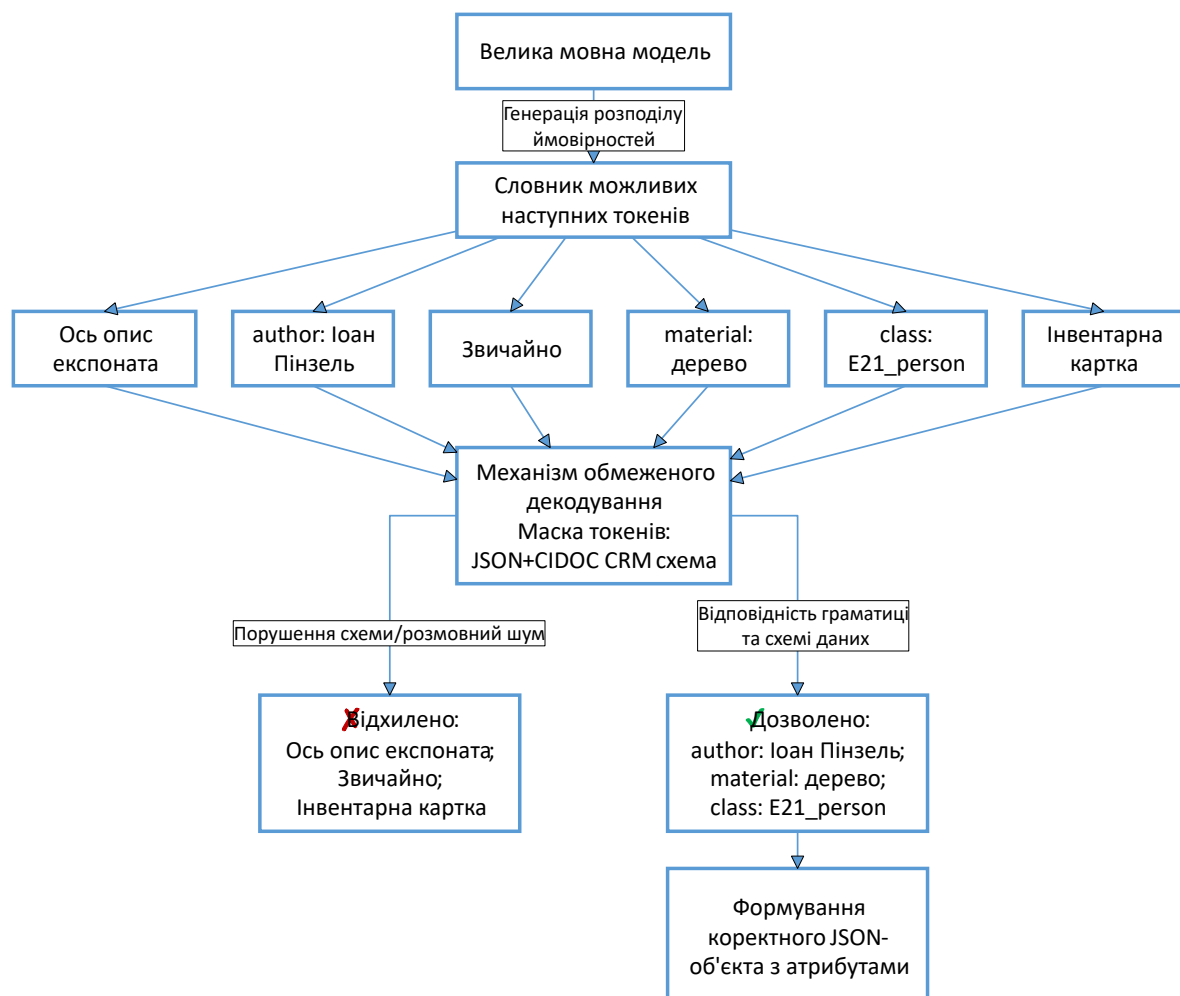


Рис. 2 – Інтеграція механізму обмеженого декодування на рівні генерації токенів

Хоча обмежене декодування повністю розв'язує проблему синтаксичних збоїв і забезпечує коректний формат без «галюцинацій структури», важливо зауважити, що це жодним чином не гарантує фактологічної точності згенерованих значень усередині цього формату [11]. ШІ-системи можуть видавати коректно структурований JSON, у який вписані фактологічно хибні дані, що вимагає впровадження додаткових механізмів семантичної перевірки.

Практична реалізація. Для практичної перевірки запропонованих теоретичних принципів та апрограмих шаблонів було розроблено функціональний прототип програмно-алгоритмічного комплексу для оцифрування інвентарних описів музейних об'єктів, побудований на основі клієнт-серверної архітектури. На рівні подання (frontend) реалізовано динамічний користувацький інтерфейс (див. Рис. 3), який забезпечує взаємодію користувача з інформаційною системою та дає змогу завантажувати текстові файли з описами та цифрові копії інвентарних карток.

Додавання нового об'єкта

Мультимодальний ШІ Асистент

Вставте текст або додайте фото інвентарної картки об'єкта

Оберіть файл або фото

Опис об'єкта...

Опрацювати з допомогою ШІ

Повна специфікація об'єкта

E42 Ідентифікатор

E41 Апеляція (Посилання)

E35 Назва

E55 Тип

E22 Створений людиною об'єкт

E78 Музейна колекція

E57 Матеріал

E54 Вимір

E3 Стан збереженості

E14 Оцінка стану

Рис. 3 – Інтерфейс адміністративної панелі прототипу програмно-алгоритмічного комплексу для оцифрування інвентарних описів музейних об'єктів

Таким чином реалізується наскрізний мультимодальний підхід до збору первинних даних. Серверний шар (backend) виконує функцію безпечного проміжного програмного забезпечення. Він інкапсулює логіку формування запитів до хмарного API великої мовної моделі, захищає ключі авторизації та здійснює первинну маршрутизацію потоків даних, виступаючи мостом між «сирими» вхідними матеріалами та кінцевою графовою базою даних.

Процес автоматизованого виділення сутностей та їхнього семантичного структурування відбувається на серверному рівні через глибоку інтеграцію з API великої мовної моделі. Для подолання проблеми синтаксичних «галюцинацій» і забезпечення безперебійної машинної зчитуваності результатів, у системі застосовано програмно-алгоритмічний механізм обмеженого декодування [8].

Під час програмного звернення до ШІ-агента на рівні бекенду динамічно формується жорстка об'єктна схема, яка концептуально відтворює виділені класи онтології CIDOC CRM. Ця конфігурація охоплює повний спектр CIDOC CRM атрибуції об'єкта культурної спадщини: від базових ідентифікаторів (E42 Identifier, E41 Appellation) і фізичних характеристик (E57 Material, E54

Dimension) до складних історичних подій (E65 Creation, E12 Production), дійових осіб (E39 Actor, E21 Person) та юридичних аспектів (E30 Right) [2].

Передаючи цю конфігуровану інформаційну сутність як обов'язковий формат генерації, розроблена система примушує нейромережу великої мовної моделі не просто аналізувати текст, а безпомилково здійснювати семантичну розмітку знайдених фактів у відповідні вузли [10]. Відповідь великої мовної моделі формується у суворо типізованому форматі JSON із мінімальним параметром креативності моделі, що унеможливлює будь-які відхилення від заданої структури та формує коректний набір даних щодо об'єкта культурної спадщини (див. Рис. 4).



Рис. 4 – Блок-схема алгоритму автоматизованого опрацювання та збереження даних про об'єкти культурної спадщини

Важливим елементом розробленого програмно-алгоритмічного комплексу є реалізація концепту «людина в процесі» [1]. Після того, як ШІ-модель здійснює мультимодальний аналіз вхідних даних та зіставляє виділені сутності з онтологією CIDOC CRM, ці дані не записуються до бази автоматично. Натомість вони відображаються у розширеній інтерактивній вебформі, де всі згенеровані значення вже попередньо розподілені по відповідних категоріях.

Музейний працівник отримує можливість безпосередньо порівняти оригінальний текст або фотографію з результатами машинного аналізу, здійснити експертну перевірку кожної сутності,

скоригувати можливі неточності записів та додати власні коментарі. Такий підхід значно знижує навантаження на кураторів, усуваючи потребу в ручному опрацюванні великих обсягів тексту. Також він зберігає за людиною функцію фінального контролю якості, що гарантує наукову достовірність інформації щодо об'єкта культурної спадщини перед її остаточним збереженням до глобального графу знань. Завершальним етапом роботи прототипу програмно-алгоритмічного комплексу є автоматизоване перетворення семантично розміченого та перевіреного музейним експертом масиву даних (структурованого JSON-об'єкта) у формат динамічних SPARQL-запитів для безпосередньої інтеграції у графову базу даних [12]. Цей процес реалізується на рівні серверного проміжного програмного забезпечення і передбачає перетворення ієрархічної JSON-структури на складну багатовимірну мережу RDF-триплетів (суб'єкт – предикат – об'єкт), що суворо відповідає семантичним специфікаціям онтології CIDOC CRM. Важливим завданням на цьому етапі є забезпечення глобальної унікальності кожного ресурсного ідентифікатора (URI) в межах спільного графа [1]. Для запобігання семантичним конфліктам, дублюванню або випадковому перезапису існуючих сутностей (наприклад, при співпадінні імен історичних осіб чи назв локацій), в бекенд-шар імплементовано програмно-алгоритмічний механізм динамічної генерації унікальних хеш-ідентифікаторів. Цей підхід гарантує, що кожен новий експонат та всі концептуально пов'язані з ним абстрактні сутності отримують власні, ізольовані та унікальні цифрові адреси в базі даних.

Формування транзакційного SPARQL-запиту типу INSERT DATA відбувається за принципом вибіркової генерації (selective triplet formation). Алгоритм програмно ітерує класи схеми даних, проте до фінального запиту входять виключно ті класи та властивості, які містять реальні текстові значення, що були успішно розпізнані великою мовною моделлю та остаточно підтверджені куратором у вебінтерфейсі. Такий алгоритмічний фільтр має принципове значення для підтримки онтологічної чистоти – він повністю унеможливорює створення «порожніх» вузлів або розірваних зв'язків у графі, гарантуючи високу структурну чистоту, цілісність та інформаційну щільність бази знань. Після перевірки та генерації, динамічний запит автоматично виконується на сервері GraphDB. У результаті цієї операції ізольований текстовий опис музейного предмета трансформується на органічну частину глобального машинозчитуваного графа. Візуалізація запису про об'єкт культурної спадщини в GraphDB подана на рисунку 5.

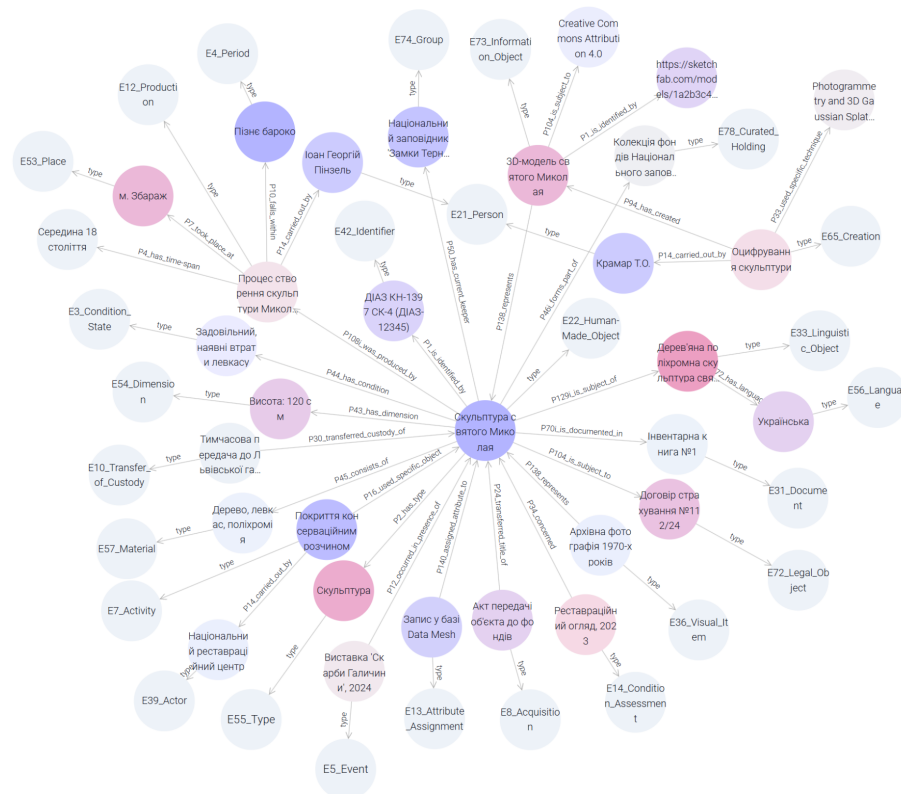


Рис. 5 – Візуалізація даних про об'єкт культурної спадщини в GraphDB

Це не лише забезпечує надійне децентралізоване збереження інформації, але й відкриває принципово нові можливості для складного семантичного пошуку, виявлення прихованих історичних крос-зв'язків між експонатами з різних колекцій та прямої інтеграції оцифрованого надбання до міжнародних агрегаторів культурної спадщини

Висновки та перспективи подальших досліджень. У дослідженні розроблено та обґрунтовано комплексний підхід до автоматизованого опрацювання неструктурованих описів об'єктів культурної спадщини за допомогою мультимодальних великих мовних моделей. Продемонстровано, що використання методів системного промптингу в поєднанні з класами онтології CIDOC CRM дає змогу розв'язати проблему семантичної інтероперабельності та перетворювати статичні тексти на машинозчитувані графи знань.

Ключовим елементом розробленої системи є застосування механізму обмеженого декодування на рівні програмного інтерфейсу. Цей програмно-алгоритмічний підхід обчислює маску токенів та відфільтровує лексичний шум безпосередньо під час генерації. Таким чином забезпечується стабільне формування коректної JSON-структури для подальшої автоматизованої обробки даних.

Важливою складовою запропонованої архітектури є імплементація гібридної взаємодії «людина в процесі». Розмічені за допомогою LLM сутності виводяться у клієнтському вебінтерфейсі для експертної перевірки музейним куратором, що суттєво оптимізує витрати ресурсів на ручну каталогізацію.

Кінцева інтеграція внесення об'єктів до бази GraphDB відбувається через автоматичне формування SPARQL-запитів. Програмно-алгоритмічні засоби вибірково створюють RDF-триплети лише для заповнених полів і динамічно генерують унікальні ідентифікатори, запобігаючи виникненню порожніх вузлів у графі. У перспективі розроблений програмно-алгоритмічний комплекс створює підґрунтя для розширеного семантичного пошуку, аналізу прихованих історичних зв'язків та інтеграції оцифрованих колекцій до міжнародних інформаційних систем агрегації культурної спадщини.

Список бібліографічного опису

1. Usta, Gökhan and Özarslantürk, Akin and Kuluman, Firat (2025) From Bibliographic Islands to Semantic Networks: A Human-in-the-Loop Knowledge Graph Pipeline for the ITU Engineering History Collection. <http://dx.doi.org/10.2139/ssrn.5965914>
2. Wang, Y., & Zhang, M. (2025). CIDOC CRM-Based Knowledge Graph Construction for Cultural Heritage Using Large Language Models. *Applied Sciences*, 15(22), 12063. <https://doi.org/10.3390/app152212063>
3. Maud Ehrmann, Ahmed Hamdi, Elvys Linhares Pontes, Matteo Romanello & Antoine Doucet (2024) Named Entity Recognition and Classification in Historical Documents: A Survey. *ACM Comput. Surv.* 56(2), Стаття 27. ACM. <https://doi.org/10.1145/3604931>
4. Bagchi, M. (2025). Toward Generative AI-Driven Metadata Modeling: A Human-Large Language Model Collaborative Approach. *Library Trends* 73(3), C. 297-322. <https://dx.doi.org/10.1353/lib.2025.a961196>
5. Koch, P., Nuñez, G. V., Arias, E. G., Heumann, C., Schöffel, M., Häberlin, A., & Assenmacher, M. (2023). A tailored handwritten-text-recognition system for medieval Latin. In *Proceedings of the Ancient Language Processing Workshop*, C. 103-110.
6. Marvin, G., Hellen, N., Jjingo, D., Nakatumba-Nabende, J. (2024). Prompt Engineering in Large Language Models. In: Jacob, I.J., Piramuthu, S., Falkowski-Gilski, P. (eds) *Data Intelligence and Cognitive Informatics. ICDICI 2023. Algorithms for Intelligent Systems*. Springer, Singapore. https://doi.org/10.1007/978-981-99-7962-2_30
7. Fan Bu, Zhan Li, Yu Jiang, et al. (2026) Reframing Museum Intelligence: A Survey of Large Language Models in Museums. Authorea. <https://doi.org/10.22541/au.177188048.80779493/v2>
8. Luca Beurer-Kellner, Marc Fischer & Martin Vechev (2023) Prompting Is Programming: A Query Language for Large Language Models. *Proc. ACM Program. Lang.* 7, PLDI, Стаття 186. ACM. <https://doi.org/10.1145/3591300>
9. Luca Beurer-Kellner, Marc Fischer & Martin Vechev (2024) Guiding LLMs the right way: fast, non-invasive constrained generation. *У Proceedings of the 41st International Conference on Machine Learning (ICML'24)*, T. 235, C. 3658-3673.
10. Michael Xieyang Liu, Frederick Liu, Alexander J. Fiannaca, Terry Koo, Lucas Dixon, Michael Terry & Carrie J. Cai (2024) "We Need Structured Output": Towards User-centered Constraints on Large Language Model Output. In *Extended Abstracts of the CHI Conference on Human Factors in Computing Systems*, Стаття 10, C. 1-9. <https://doi.org/10.1145/3613905.3650756>
11. Singh, Abhinav Kumar, Khurdula, Harsha Vardhan, Khemlani, Yoeven D & Agarwal, Vineet (2026) The Structured Output Benchmark: A Multi-Source Benchmark for Evaluating Structured Output Quality in Large Language Models. <https://doi.org/10.48550/arXiv.2604.25359>
12. Schimmenti, A., Pasqual, V., Vitali, F. & van Erp, M. (2026) Knowledge graphs generation from cultural heritage texts: combining LLMs and ontological engineering for scholarly debates. *Journal of Documentation*, T. 82 № 7, C. 206-250. <https://doi.org/10.1108/JD-07-2025-0203>
13. Davide Varagnolo, Dora Melo & Irene Pimenta Rodrigues (2025) Translating Natural Language Questions into

- CIDOC-CRM SPARQL Queries to Access Cultural Heritage Knowledge Bases. J. Comput. Cult. Herit. 18, 2, Стаття 21. ACM. <https://doi.org/10.1145/3715156>
14. UNESCO. (2022). Recommendation on the ethics of artificial intelligence. URL: <https://unesdoc.unesco.org/ark:/48223/pf0000381137> (дата доступу: 30.07.2026)
15. Europeana Foundation, & Open Future Foundation. (2025). Publishing cultural heritage data in the age of AI. Europeana PRO. URL: https://pro.europeana.eu/files/Europeana_Professional/Publications/251202PublishingCulturalHeritageDataInTheAgeOfAI.pdf (дата доступу: 30.07.2026)
16. Pasikowska-Schnass, M., & Lim, Y.-S. (2023). Artificial intelligence in the context of cultural heritage and museums: Complex challenges and new opportunities (PE 747.120). European Parliamentary Research Service. URL: [https://www.europarl.europa.eu/RegData/etudes/BRIE/2023/747120/EPRS_BRI\(2023\)747120_EN.pdf](https://www.europarl.europa.eu/RegData/etudes/BRIE/2023/747120/EPRS_BRI(2023)747120_EN.pdf) (дата доступу: 30.07.2026)
17. Chanjin Zheng, Zengyi Yu, Yilin Jiang, Mingzi Zhang, Xunuo Lu, Jing Jin & Liteng Gao (2025) ArtMentor: AI-Assisted Evaluation of Artworks to Explore Multimodal Large Language Models Capabilities. In Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems (CHI '25), Стаття 659, С. 1-18. <https://doi.org/10.1145/3706598.3713274>
18. Singh, Mayank & Kumar, Abhijeet & Donaparthi, Sasidhar & Karambelkar, Gayatri. (2025). Leveraging Retrieval Augmented Generative LLMs For Automated Metadata Description Generation to Enhance Data Catalogs. <https://doi.org/10.48550/arXiv.2503.09003>
19. Nair, A., Goh, Z. R., Liu, T., & Huang, A. Y. (2025). Web archives metadata generation with gpt-4o: Challenges and insights. Information Technology and Libraries, 44(2).

References

1. Usta, Gökhan and Özarslantürk, Akın and Kuluman, Firat (2025) From Bibliographic Islands to Semantic Networks: A Human-in-the-Loop Knowledge Graph Pipeline for the ITU Engineering History Collection. <http://dx.doi.org/10.2139/ssrn.5965914>
2. Wang, Y., & Zhang, M. (2025). CIDOC CRM-Based Knowledge Graph Construction for Cultural Heritage Using Large Language Models. Applied Sciences, 15(22), 12063. <https://doi.org/10.3390/app152212063>
3. Maud Ehrmann, Ahmed Hamdi, Elvys Linhares Pontes, Matteo Romanello & Antoine Doucet (2024) Named Entity Recognition and Classification in Historical Documents: A Survey. ACM Comput. Surv. 56(2), Article 27. ACM. <https://doi.org/10.1145/3604931>
4. Bagchi, M. (2025). Toward Generative AI-Driven Metadata Modeling: A Human-Large Language Model Collaborative Approach. Library Trends 73(3), 297-322. <https://dx.doi.org/10.1353/lib.2025.a961196>
5. Koch, P., Nuñez, G. V., Arias, E. G., Heumann, C., Schöffel, M., Häberlin, A., & Assenmacher, M. (2023). A tailored handwritten-text-recognition system for medieval Latin. In Proceedings of the Ancient Language Processing Workshop, P 103-110.
6. Marvin, G., Hellen, N., Jjingo, D., Nakatumba-Nabende, J. (2024). Prompt Engineering in Large Language Models. In: Jacob, I.J., Piramuthu, S., Falkowski-Gilski, P. (eds) Data Intelligence and Cognitive Informatics. ICDICI 2023. Algorithms for Intelligent Systems. Springer, Singapore. https://doi.org/10.1007/978-981-99-7962-2_30
7. Fan Bu, Zhan Li, Yu Jiang, et al. (2026) Reframing Museum Intelligence: A Survey of Large Language Models in Museums. Authorea. <https://doi.org/10.22541/au.177188048.80779493/v2>
8. Luca Beurer-Kellner, Marc Fischer & Martin Vechev (2023) Prompting Is Programming: A Query Language for Large Language Models. Proc. ACM Program. Lang. 7, PLDI, Article 186. ACM. <https://doi.org/10.1145/3591300>
9. Luca Beurer-Kellner, Marc Fischer & Martin Vechev (2024) Guiding LLMs the right way: fast, non-invasive constrained generation. In Proceedings of the 41st International Conference on Machine Learning (ICML'24), Vol. 235, P. 3658-3673.
10. Michael Xieyang Liu, Frederick Liu, Alexander J. Fiannaca, Terry Koo, Lucas Dixon, Michael Terry & Carrie J. Cai (2024) "We Need Structured Output": Towards User-centered Constraints on Large Language Model Output. In Extended Abstracts of the CHI Conference on Human Factors in Computing Systems, Article 10, P. 1-9. <https://doi.org/10.1145/3613905.3650756>
11. Singh, Abhinav Kumar, Khurdula, Harsha Vardhan, Khemlani, Yoeven D & Agarwal, Vineet (2026) The Structured Output Benchmark: A Multi-Source Benchmark for Evaluating Structured Output Quality in Large Language Models. <https://doi.org/10.48550/arXiv.2604.25359>
12. Schimmenti, A., Pasqual, V., Vitali, F. & van Erp, M. (2026) Knowledge graphs generation from cultural heritage texts: combining LLMs and ontological engineering for scholarly debates. Journal of Documentation, Vol. 82 No. 7, P. 206-250. <https://doi.org/10.1108/JD-07-2025-0203>
13. Davide Varagnolo, Dora Melo & Irene Pimenta Rodrigues (2025) Translating Natural Language Questions into CIDOC-CRM SPARQL Queries to Access Cultural Heritage Knowledge Bases. J. Comput. Cult. Herit. 18, 2, Article 21. ACM. <https://doi.org/10.1145/3715156>
14. UNESCO. (2022). Recommendation on the ethics of artificial intelligence. <https://unesdoc.unesco.org/ark:/48223/pf0000381137> (date of access: 30.07.2026)
15. Europeana Foundation, & Open Future Foundation. (2025). Publishing cultural heritage data in the age of AI. Europeana PRO. https://pro.europeana.eu/files/Europeana_Professional/Publications/251202PublishingCulturalHeritageDataInTheAgeOfAI.pdf (date of access: 30.07.2026)
16. Pasikowska-Schnass, M., & Lim, Y.-S. (2023). Artificial intelligence in the context of cultural heritage and museums: Complex challenges and new opportunities (PE 747.120). European Parliamentary Research Service. [https://www.europarl.europa.eu/RegData/etudes/BRIE/2023/747120/EPRS_BRI\(2023\)747120_EN.pdf](https://www.europarl.europa.eu/RegData/etudes/BRIE/2023/747120/EPRS_BRI(2023)747120_EN.pdf) (date of access: 30.07.2026)

17. Chanjin Zheng, Zengyi Yu, Yilin Jiang, Mingzi Zhang, Xunuo Lu, Jing Jin & Liteng Gao (2025) ArtMentor: AI-Assisted Evaluation of Artworks to Explore Multimodal Large Language Models Capabilities. In Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems (CHI '25), Article 659, P. 1-18. <https://doi.org/10.1145/3706598.3713274>
18. Singh, Mayank & Kumar, Abhijeet & Donaparthi, Sasidhar & Karambelkar, Gayatri. (2025). Leveraging Retrieval Augmented Generative LLMs For Automated Metadata Description Generation to Enhance Data Catalogs. <https://doi.org/10.48550/arXiv.2503.09003>
19. Nair, A., Goh, Z. R., Liu, T., & Huang, A. Y. (2025). Web archives metadata generation with gpt-4o: Challenges and insights. *Information Technology and Libraries*, 44(2).

Історія статті:

Отримано: 29.05.2026 Доопрацьовано: 12.08.2026 Прийнято до друку: 23.09.2026 Опубліковано: 29.09.2026