

АНАЛІТИЧНА МОДЕЛЬ МАНУАЛЬНОЇ ІЄРАРХІЧНОЇ КЛАСТЕРИЗАЦІЇ ДАНИХ: СИСТЕМНИЙ ПІДХІД ДО СТРУКТУРИЗАЦІЇ ТА КЛАСИФІКАЦІЇ МИТНИХ ДАНИХ

ANALYTICAL MODEL OF MANUAL HIERARCHICAL DATA CLUSTERING: A SYSTEMATIC APPROACH TO STRUCTURING AND CLASSIFYING CUSTOMS DATA

Аналітична модель Manual Hierarchical Data Clustering пропонує структурований підхід до митної класифікації, інтегруючи методи комп'ютерної лінгвістики, дата-сайєнсу та експертних систем. Він вирішує проблеми фрагментованої документації й динамічних регуляторних вимог, знижуючи ризики помилок. MHDC поєднує лінгвістичні моделі для аналізу текстів, ієрархічну кластеризацію для групування товарів і дерева рішень для присвоєння кодів Гармонізованої системи. Процес включає структуроване вилучення даних, багаторівневе групування та автоматизовану класифікацію, що підвищує точність і швидкість. Емпірична верифікація підтверджує ефективність MHDC у реальних умовах. Модель узгоджується з міжнародними стандартами, адаптуючись до національних вимог, зокрема Митного кодексу України. Рекомендації щодо впровадження охоплюють підготовку фахівців і гібридні моделі для регіонів із обмеженою інфраструктурою, сприяючи гармонізації торговельних процесів.

Ключові слова: MHDC, митна класифікація, ієрархічна кластеризація, Гармонізована система, напівавтоматична обробка.

This article introduces the analytical model of Manual Hierarchical Data Clustering (MHDC), designed as a comprehensive system for structuring and classifying customs data under conditions of fragmented documentation and dynamic regulatory environments. The MHDC approach synthesizes methods from computational linguistics, data science, and rule-based expert systems to enhance the accuracy and transparency of customs classification processes. The model incorporates natural language processing (NLP) techniques to extract structured information from unstructured product descriptions found in commercial documents, enabling standardization and reduction of terminological ambiguities. Hierarchical clustering methods, employing both agglomerative and divisive strategies, are applied to organize goods into multilevel categories based on attribute similarities, supporting a scalable and adaptive data structure that reflects the complexity of customs nomenclature. Further, MHDC utilizes decision tree logic to automate the assignment of Harmonized System (HS) codes. This mechanism operationalizes a sequence of control questions aligned with international classification standards, allowing for a transparent and adaptable coding process that accommodates regulatory updates and exceptions. The model's architecture consists of a standardized taxonomy stage for defining key product attributes, a clustering stage for forming a logical hierarchy of goods, and an automated classification stage ensuring consistent HS code assignment. Optical character recognition (OCR), tokenization, data normalization, and semantic vectorization techniques, such as TF-IDF and Word2Vec, are integral to the data preprocessing pipeline, facilitating efficient information extraction and logical consistency checking. The MHDC model aligns with the standards of the World Customs Organization and the Customs Code of Ukraine, ensuring its compatibility across jurisdictions. It emphasizes the importance of a semi-automated processing approach, preserving human expert involvement where necessary while optimizing time and resource consumption. Special recommendations include the training of customs specialists, the gradual integration of digital tools in regions with limited infrastructure, and the establishment of collaborations with standardization organizations to maintain methodological relevance. In the long term, MHDC aims to serve as a foundation for advancing the digitalization of customs procedures, fostering harmonized trade practices, and strengthening global supply chain integration.

Key words: MHDC, customs classification, hierarchical clustering, Harmonized System, semi-automatic processing.

УДК 004.89:339.543:006.83

DOI: <https://doi.org/10.32782/bses.92-23>

Сорочак О.З.¹

к.т.н., доцент кафедри менеджменту організацій,
Національний університет
«Львівська політехніка»

Сарвас М.М.

кандидатка на ступінь Магістра права,
Юридичний Коледж Чикаго-Кент,
Іллінойський технологічний інститут (ІІТ)

Sorochak Oleh

Lviv Polytechnic National University

Sarvas Marta

Chicago-Kent College of Law,

Illinois Institute of Technology (IIT)

Постановка проблеми. У сучасних умовах глобалізації митна класифікація товарів стикається зі зростаючою складністю через стрімке розширення номенклатури товарів, фрагментованість документації та динамічні зміни в регуляторних вимогах. Неструктуровані текстові описи в рахунках-фактурах чи пакувальних листах ускладнюють точне визначення кодів Гармонізованої системи, що призводить до помилок у класифікації. Такі помилки спричиняють затримки в митному оформленні, фінансові санкції та порушення міжнародних торговельних угод. Водночас ручні методи

обробки даних, які досі застосовуються в багатьох митних системах, не справляються з обсягами інформації та потребують значних часових і людських ресурсів. Ситуацію ускладнює неоднозначність термінології, коли один і той самий товар може мати різні назви чи описи залежно від регіону чи мови. Існуючі автоматизовані системи часто обмежені в адаптивності до специфічних регуляторних вимог чи нових категорій товарів, що знижує їхню ефективність. Отже, виникає потреба в розробці гнучкого, прозорого та відтворюваного підходу до митної класифікації, який би поєднував

¹ ORCID: <https://orcid.org/0000-0003-2889-9284>

обробку неструктурованих даних, ієрархічне структурування та чітку логіку прийняття рішень, мінімізуючи ризики помилок і оптимізуючи митні процедури.

Аналіз останніх досліджень і публікацій.

У сучасних дослідженнях кластеризації даних ієрархічні методи посідають особливе місце завдяки своїй здатності виявляти багаторівневу структуру у складних інформаційних масивах. Структуризація митних даних та визначення кодів Гармонізованої системи (ГС) вимагає особливого підходу, що поєднує автоматизовані обчислення й експертне втручання. У цьому контексті пропонується аналітична модель мануальної ієрархічної кластеризації даних, яка опирається на наукові напрацювання в галузі кластеризації та аналізу великих даних. Саксена та ін. описують ієрархічні методи розглядаючи їх як потужний інструмент для роботи з багаторівневими структурами даних, а також, хоча автоматичні ієрархічні алгоритми здатні виявляти природні групування, вони часто не враховують контекстуальні особливості предметної області, що обмежує їхню ефективність. Це підкреслює необхідність доповнення алгоритмічних рішень людською експертизою [1]. Практичне застосування мануальної ієрархічної кластеризації продемонстроване в дослідженні Еллефсен і Сміт, де автори використовували Bayesian finite mixture models для аналізу геохімічних даних, а кінцеву структуру кластерів формували за допомогою ручної корекції. Такий підхід дозволив забезпечити більш релевантне групування об'єктів завдяки глибокому розумінню особливостей даних [2]. Ідея активного залучення експерта в процес кластеризації розкривається у роботі Янг, де запропоновано інтерактивне управління побудовою дендрограм через візуальні інтерфейси [3]. У роботі Езугву та ін. подано, що ієрархічна кластеризація визнається однією з найбільш надійних технологій для задач, що потребують багаторівневої інтерпретації даних, акцентується на важливості адаптивних методів, що здатні враховувати специфіку даних [4].

Питання вибору методів кластеризації відповідно до завдань аналізу підкреслюють Родрігес та ін., які у своїй статті здійснюють порівняння різних алгоритмів кластеризації за якістю та стабільністю результатів [5]. Закі та Мейра розглядають як базові поняття, так і сучасні підходи до аналізу даних, ручне втручання в структуру ієрархії може стати критично необхідним там, де автоматичні алгоритми стикаються з невизначеністю або складною природою даних [6]. У роботах Філіп Чен та Чжан і Ікегву та ін. також зазначається, що для роботи з Big Data необхідні гібридні моделі, які комбінують автоматизацію і можливість втручання користувача для оптимізації результатів аналізу [7; 8].

Аспект залучення людини в процес обробки даних у митній сфері розглядали Кім та ін. [9]. Роботи Вазар і Бан та Рахманов свідчать про зростання потреби в нових аналітичних інструментах для митниць [11]. Важливість інтеграції великих даних і штучного інтелекту для оптимізації бізнес-процесів розглядають Рашид та ін. [12].

Отож, аналіз літератури підтверджує актуальність МНДС, що поєднує автоматизацію і ручне втручання для високоточної структуризації митних даних і уточнення кодів ГС.

Постановка завдання. Мета цієї статті полягає в розробці та обґрунтуванні методичних рекомендацій для впровадження моделі Manual Hierarchical Data Clustering у процеси митної класифікації. Вона прагне систематизувати підходи до обробки складних, фрагментованих даних, що виникають у митних операціях.

Виклад основного матеріалу дослідження. Аналітична модель мануальної ієрархічної кластеризації даних (МНДС) поєднує теоретичні напрацювання з комп'ютерної лінгвістики, дата-сайєнсу та експертних систем на основі правил, створюючи структурований підхід до вирішення складнощів у митній класифікації. Цей підхід враховує безліч чинників, пов'язаних із фрагментованою документацією, динамічними регуляторними вимогами та високими ризиками помилкової класифікації товарів. Інтеграція лінгвістичних моделей кластеризації становить один із ключових елементів МНДС. У митній діяльності обробка неструктурованого тексту, як-от описи товарів у документах, потребує системного підходу. Моделі обробки природної мови (NLP) адаптуються для вилучення характеристик товарів через токенизацію, розмітку частин мови та виявлення іменованих сутностей, що знижує неоднозначність і сприяє уніфікації маркування [1]. Ієрархічна кластеризація, як другий елемент, застосовує агломеративні або дивізивні методи для групування товарів за схожістю атрибутів. Це дозволяє створювати багаторівневу структуру – від загальних категорій до специфічних підгруп, що адаптується до різноманітності даних [4]. Використання логіки дерев рішень забезпечує прозорий механізм присвоєння кодів Гармонізованої системи (ГС) через послідовність контрольних запитань, що враховує регуляторні вимоги та винятки [5]. Поєднання цих складників забезпечує високу точність і швидкість обчислень при одночасному дотриманні митних стандартів.

Перший етап моделі МНДС зосереджується на створенні стандартизованої таксономії, яка слугує основою для впорядкування митних даних і забезпечує надійний фундамент для всіх наступних етапів обробки. Цей процес розпочинається з ретельного визначення ключових атрибутів, що відображають сутність товарів і є критично важливими для митної класифікації. Назва товару

формується як стислий ідентифікатор, що чітко передає його основну суть, уникаючи двозначностей. Детальний опис включає розгорнуту інформацію про фізичні характеристики, хімічний склад, функціональне призначення чи унікальні особливості товару, що дозволяє глибше зрозуміти його природу. Показники кількості та вартості охоплюють числові дані, такі як вага, об'єм, кількість одиниць, а також ціна в різних валютах, що потребує уніфікації для забезпечення порівнянності. Критерії країни походження фіксують відомості про місце виробництва, складання чи походження сировини, що має значення для застосування тарифних пільг чи обмежень. Тарифні індикатори включають специфічні маркери, які вказують на можливі митні ставки, спеціальні регуляторні вимоги чи пільгові режими, пов'язані з торговельними угодами. Для обробки документів застосовуються алгоритми напівавтоматичної обробки, які поєднують автоматизацію з можливістю ручного втручання для підвищення точності. Технології оптичного розпізнавання символів (OCR) перетворюють відскановані документи, такі як рахунки-фактури чи пакувальні листи, у текстовий формат, що робить їх придатними для подальшого аналізу. Токенізація розбиває текст на окремі слова, фрази чи семантичні одиниці, створюючи структуровану основу для вилучення релевантної інформації. Нормалізація даних усуває неоднорідність, наприклад, шляхом приведення різних форматів чисел, одиниць виміру чи валют до єдиного стандарту, а також виправлення орфографічних помилок чи невідповідностей у термінології. Механізми перевірки даних відіграють ключову роль у забезпеченні якості. Вони виявляють пропущені поля, такі як відсутність даних про країну походження чи вагу, і позначають їх для ручного доопрацювання. Некоректні значення, наприклад, негативна вага чи нереалістична ціна, автоматично коригуються або направляються на додаткову перевірку. Перевірка логічної цілісності оцінює, чи відповідає опис товару заявленій категорії, чи узгоджуються кількісні показники з іншими атрибутами. Цей етап спирається на принципи інформаційної архітектури, які довели свою ефективність у систематизації хаотичних і неструктурованих даних. Завдяки чіткому структуруванню інформації знижується суб'єктивність інтерпретацій, що часто виникає через неоднорідність документації чи різницю в термінології. У результаті ймовірність помилок перед переходом до етапу кластеризації суттєво зменшується, що забезпечує надійну основу для подальшого аналізу [3].

Другий етап аналітичної моделі Manual Hierarchical Data Clustering (MHDC) зосереджується на формуванні ієрархічної структури даних шляхом групування товарів на основі їхніх характеристик, що є критично важливим для забезпечення точної та ефективної митної класифікації.

Кожному товарному запису призначається багатокомпонентний вектор ознак, який охоплює числові параметри, такі як вага, габарити, кількість одиниць і вартість, а також семантичні вектори, створені за допомогою передових методів обробки тексту. Метод TF-IDF оцінює важливість слів у текстових описах товарів, виділяючи ключові терміни, тоді як Word2Vec перетворює текст у числові вектори, враховуючи контекстуальні зв'язки між словами, що дозволяє виявляти приховані семантичні взаємозв'язки. Ці вектори разом створюють комплексне представлення кожного товару, яке враховує як кількісні, так і якісні аспекти. Для групування застосовуються алгоритми ієрархічної кластеризації, зокрема агломеративний підхід, що починається з розгляду кожного товару як окремого кластера з подальшим їх об'єднанням на основі схожості, або дивізивний підхід, який бере за основу єдину групу всіх товарів і поступово розподіляє її на менші підгрупи. Процес кластеризації розгортається через кілька послідовних фаз. Спочатку відбувається ініціалізація, під час якої кожен товар розглядається як окремий об'єкт. Далі алгоритм ітеративно оцінює схожість між об'єктами, об'єднуючи їх у більші кластери або розподіляючи на менші залежно від обраного підходу. Порогові значення схожості, що визначають межі групування, налаштовуються на основі аналізу історичних даних, що забезпечує адаптацію моделі до специфіки митних операцій, таких як різноманітність товарів чи регуляторні вимоги. Результатом цього процесу є ієрархічне дерево кластерів, у якому верхні рівні представляють широкі категорії як «електроніка» чи «текстиль», тоді як нижні рівні деталізують підгрупи, такі як «смартфони з певними технічними характеристиками» чи «тканини з конкретним складом». Така багаторівнева структура дозволяє митним спеціалістам гнучко переглядати як узагальнені групи товарів для швидкого аналізу, так і деталізовані кластери для точного визначення категорій, що суттєво знижує ризик помилкових узагальнень і підвищує точність класифікації. Крім того, ієрархічна організація даних полегшує виявлення аномалій, таких як товари з нестандартними характеристиками, що можуть вимагати додаткового аналізу. Цей підхід, підкріплений науковими дослідженнями в галузі аналізу даних, забезпечує надійну основу для автоматизації митних процесів, зберігаючи при цьому можливість людського контролю там, де потрібна експертна оцінка [6].

Третій етап реалізує автоматизоване присвоєння кодів Гармонізованої системи (ГС) за допомогою бінарного дерева рішень, яке забезпечує чітку та ефективну класифікацію товарів. Кожен вузол цього дерева перевіряє певну умову, що базується на характеристиках товару чи регуляторних вимогах. Процес побудови дерева включає кілька

ключових етапів. Спочатку формуються контрольні запитання, які створюються на основі стандартів ГС і нормативних вимог. Далі відбувається інтеграція з правилами, що враховують виключення, спеціальні примітки та тарифні преференції, забезпечуючи відповідність класифікації міжнародним і локальним стандартам. Структура дерева регулярно оновлюється, адаптуючись до змін у нормативному полі, що дозволяє системі залишатися актуальною. Важливим елементом є цикл зворотного зв'язку: якщо код ГС присвоєно неправильно, система аналізує помилку та коригує логіку дерева, підвищуючи точність подальших класифікацій. У результаті кожен товар отримує однозначний код ГС, що відповідає його характеристикам і регуляторним вимогам. Прозорість логіки бінарного дерева полегшує проведення аудиту, дозволяючи перевірити правильність присвоєння кодів. Завдяки адаптивності системи до нових товарів і змін у законодавстві забезпечується її довгострокова ефективність і надійність у процесі класифікації. Схематично аналітична модель MHDC представлена на рис. 1.

Оцінка ефективності моделі MHDC базується на ретельному порівнянні результатів митної класифікації до та після її впровадження, що дозволяє виявити конкретні покращення в обробці даних. Для цього варто застосовувати статистичні метрики, які забезпечують об'єктивну оцінку продуктивності системи. Зокрема F1-міру, що відображає баланс між точністю та повнотою, і слугує ключовим показником здатності MHDC досягати оптимальних результатів у складних умовах. Точність, яка вимірює відсоток правильно класифікованих товарів, підкреслює надійність моделі в уникненні помилок, що можуть призвести до фінансових чи регуляторних санкцій. Повнота, своєю чергою, демонструє, наскільки система здатна охопити всі релевантні випадки, включно з нетиповими товарами чи нестандартними описами, що є частим викликом у митних операціях. Додатково варто оцінити такі параметри, як час обробки однієї декларації та

частота звернень до митних органів, що дозволяє кількісно підтвердити економію ресурсів і підвищення ефективності. Завдяки модульній структурі та гнучкості MHDC гармонізується з міжнародними стандартами, зокрема з Гармонізованою системою Всесвітньої митної організації та Митним кодексом України, що робить його універсальним інструментом для різних юрисдикцій. Впровадження запропонованої аналітичної моделі MHDC потребує системного підходу, який включає підготовку фахівців із глобальної відповідності торговельним правилам через спеціалізовані навчальні програми, що підвищують кваліфікацію митних брокерів і спеціалістів із зовнішньої торгівлі. Співпраця з організаціями зі стандартизації, такими як Всесвітня митна організація чи регіональні митні союзи, сприятиме обміну досвідом і вдосконаленню методології, забезпечуючи її актуальність у контексті змін у регуляторному полі. Для регіонів із обмеженою цифровою інфраструктурою пропонується використання гібридних моделей, які поєднують цифрові технології з паперовими процесами, дозволяючи поступово модернізувати застарілі системи без значних фінансових витрат. Ці заходи сприяють гармонізації торговельних процесів, підвищенню прозорості митних процедур і зниженню ризиків, пов'язаних із помилковою класифікацією. У довгостроковій перспективі MHDC може стати основою для розробки нових стандартів автоматизації митних операцій, зміцнюючи інтеграцію глобальних торговельних ланцюгів і відповідаючи сучасним викликам цифровізації.

Висновки. Розроблена аналітична модель мануальної ієрархічної кластеризації даних (MHDC) становить інноваційний підхід до структуризації та класифікації митних даних, що успішно вирішує проблеми, пов'язані з обробкою фрагментованої документації та динамічними змінами в регуляторному середовищі. Поєднання методів комп'ютерної лінгвістики, дата-сайенсу та експертних систем створює надійну основу для підвищення точності митної класифікації та скорочення часу обробки даних.



Рис. 1. Аналітична модель мануальної ієрархічної кластеризації даних

Джерело: сформовано автором

Наукова цінність роботи полягає в систематизації та інтеграції різних методів аналізу даних, зокрема лінгвістичних моделей, ієрархічної класифікації та логіки дерев рішень, у єдину комплексну систему, що враховує специфіку митної діяльності. Запропонована модель дозволяє ефективно структурувати неоднорідні дані, зменшуючи суб'єктивність інтерпретації та підвищуючи прозорість процесу класифікації. Використання багатокомпонентних векторів ознак, що враховують як кількісні параметри, так і семантичні характеристики товарів, суттєво покращує точність групування та наступної класифікації.

Практична цінність МНДС виявляється в її здатності адаптуватися до різних юрисдикцій і нормативних вимог, що робить її універсальним інструментом для митних органів різних країн. Модель забезпечує зниження частоти помилок при класифікації товарів, зменшення ризиків фінансових санкцій і скорочення часу, необхідного для митного оформлення. Напівавтоматичний характер обробки даних дозволяє зберегти експертний контроль над процесом класифікації при одночасному підвищенні ефективності та продуктивності праці митних спеціалістів. Гнучкість моделі уможливіло її впровадження в регіонах з різним рівнем цифрової інфраструктури, що особливо важливо для країн, що розвиваються.

Перспективи подальших досліджень пов'язані з розширенням функціональності МНДС шляхом інтеграції більш просунутих методів машинного навчання для автоматичного виявлення нових категорій товарів і адаптації до змін у законодавстві. Важливим напрямом є вдосконалення механізмів обробки природної мови для підвищення точності аналізу описів товарів різними мовами, що посилить міжнародну універсальність моделі. Додаткового дослідження потребує розробка методів оцінки невизначеності та ризиків при класифікації, що дозволить ідентифікувати потенційно проблемні випадки і спрямовувати їх на додатковий експертний аналіз. Особливу увагу варто приділити інтеграції МНДС з існуючими митними інформаційними системами для забезпечення безшовного обміну даними та підтримки прийняття рішень уздовж усього ланцюга митних операцій. Ці напрями досліджень сприятимуть подальшому розвитку та вдосконаленню запропонованої моделі, що в підсумку призведе до гармонізації торговельних процесів і зміцнення міжнародної співпраці в митній сфері.

БІБЛІОГРАФІЧНИЙ СПИСОК:

1. Saxena A., Prasad M., Gupta A., Bharill N., Patel O. P., Tiwari A., Er M.J., Ding W., Lin C.-T. (2017). A review of clustering techniques and developments. *Neurocomputing*. 2017. Vol. 267. P. 664–681. DOI: <https://doi.org/10.1016/j.neucom.2017.06.053>

2. Ellefsen K.J., Smith D.B. Manual hierarchical clustering of regional geochemical data using a Bayesian finite mixture model. *Applied Geochemistry*. 2016. Vol. 75. P. 200–210. DOI: <https://doi.org/10.1016/j.apgeochem.2016.05.016>

3. Yang W., Wang X., Lu J., Dou W., Liu S. Interactive Steering of Hierarchical Clustering. *IEEE Transactions on Visualization and Computer Graphics*. 2021. Vol. 27 (10). P. 3953–3967. DOI: <https://doi.org/10.1109/tvcg.2020.2995100>

4. Ezugwu A.E., Ikotun A.M., Oyelade O. O., Abualigah L., Agushaka J.O., Ek, C.I., Akinyelu A.A. A comprehensive survey of clustering algorithms: State-of-the-art machine learning applications, taxonomy, challenges, and future research prospects. *Engineering Applications of Artificial Intelligence*. 2022. Vol. 110. P. 104743. DOI: <https://doi.org/10.1016/j.engappai.2022.104743>

5. Rodriguez M.Z., Comin C.H., Casanova D., Bruno O.M., Amancio D.R., Costa L. da F., Rodrigues F.A. Clustering algorithms: A comparative approach. *PLOS ONE*. 2019. Vol. 14 (1), e0210236. DOI: <https://doi.org/10.1371/journal.pone.0210236>

6. Zaki M.J., Jr M.W. Data Mining and Analysis: Fundamental Concepts and Algorithms. Cambridge University Press, 2018.

7. Philip Chen C.L., Zhang C.-Y. Data-intensive applications, challenges, techniques and technologies: A survey on Big Data. *Information Sciences*. 2014. Vol. 275. P. 314–347. DOI: <https://doi.org/10.1016/j.ins.2014.01.015>

8. Ikegwu A.C., Nweke H.F., Anikwe C.V., Alo U.R., Okonkw O.R. (2022). Big data analytics for data-driven industry: a review of data sources, tools, challenges, solutions, and research directions. *Cluster Computing*. 2022. Vol. 25 (5). P. 3343–3387. DOI: <https://doi.org/10.1007/s10586-022-03568-5>

9. Kim S., Mai T.-D., Han S., Park S., Thi Nguyen D. K., So J., Singh K., Cha M. (2023). Active Learning for Human-in-the-Loop Customs Inspection. *IEEE Transactions on Knowledge and Data Engineering*. 2023. Vol. 35 (12). P. 12039–12052. DOI: <https://doi.org/10.1109/tkde.2022.3144299>

10. Vozár A., Bán E. Digitalisation, digital transformation – In the practice of the National Tax and Customs Office. *Pénzügyi Szemle = Public Finance Quarterly*. 2024. Vol. 70. No. 1. P. 85–109. DOI: https://doi.org/10.35551/pfq_2024_1_5

11. Jaloliddin R. Digitalization in Global Trade: Opportunities and Challenges for Investment. *Global Trade and Customs Journal*. 2023. Vol. 18, Issue 10. P. 391–395. DOI: <https://doi.org/10.54648/gtcj2023043>

12. Rashid A., Baloch N., Rasheed R., Ngh A. H. (2024). Big data analytics-artificial intelligence and sustainable performance through green supply chain practices in manufacturing firms of a developing country. *Journal of Science and Technology Policy Management*. 2024. Vol. 16 (1). P. 42–67. DOI: <https://doi.org/10.1108/jstpm-04-2023-0050>