

Всі вищезгадані дії представлені через зручні елементи інтерфейсу користувача: панелі з логічно згрупованими блоками, випадаючі списки, текстові поля, кнопки [3].

Розроблена система також надає інформацію про те хто видав завдання, хто відповідальний за його виконання та хто з працівників працює над цим завданням.

The screenshot displays a web interface for task management. At the top, there are two dropdown menus: 'Тип завдання' (Task Type) set to 'Build/Destroy' and 'Бізнес пріоритет' (Business Priority) set to 'Середній' (Medium). To the right, a user profile for 'Agent Ure (agent@gmail.com)' is shown with a green 'AU' avatar and the text 'Завдання на перевірку' (Task for review). Below these are two main sections: 'Оцінювання' (Evaluation) and 'Додаткова інформація' (Additional information). The 'Оцінювання' section shows a progress bar for 'Виконано' (Completed) at 42.0 out of 45.0 hours, with a 'Зафіксувати' (Fix) button. The 'Додаткова інформація' section includes a 'Місце' (Location) field with 'effeef', and 'Старт' (Start) and 'Кінець' (End) dates of '28.11.2022' and '30.11.2022' respectively. Below these are two more sections: 'Список підзавдань' (List of sub-tasks) and 'Коментарі' (Comments). The 'Список підзавдань' section shows a 'Додати підзавдання' (Add sub-task) button and a list item 'tech_task1-1' with 'Повернути на виконання' (Return to execution) and 'Закрити' (Close) buttons. The 'Коментарі' section has a 'Текст' (Text) input field and a 'Додати коментар' (Add comment) button.

Рис. 3. Сторінка з детальною інформацією про завдання

Висновки. Розроблено інтелектуальну систему менеджменту підприємницької діяльності для керування персоналом, розмежування доступу працівників, контролю виконання завдань та обліку клієнтів, постачальників і ресурсів.

Список використаних літературних джерел

- [1] Hacking Growth: How Today's Fastest-Growing Companies Drive Breakout Success. / Morgan Brown – 2017 – P. 35 – 40
- [2] The Startup Way / Eric Ries – 2017 – P. 23-25.
- [3] Top 8 React charts libraries [Електронний ресурс] : [Веб-сайт]. – Режим доступу: <https://blog.logrocket.com/top-8-react-chart-libraries/> (дата звернення 12.09.2023). – Назва з екрана.

Andrii Pirko, Iryna Boretska*

postgraduate student, Ukrainian National Forestry University, Lviv, Ukraine.
pirko.andrii@nltu.edu.ua.

* associate professor of CS Department, Ukrainian National Forestry University, Lviv, Ukraine. <https://orcid.org/0000-0002-6767-104X>.
iryna.boretska@gmail.com.

Using machine learning tools and methods for finding similar products according to user

Abstract – The main methods of applying machine learning to search for similar products according to user preferences is described. A method is developed to create a model that can be used to create information systems for searching for relevant products. The results of the study can be used to implement an information system for searching for goods in any field of e-commerce.

Keywords – mathematical model, gradient boosting, top-N algorithm.

Recommender systems are designed to help users find products or goods that best suit their needs in a web context. They collect information about users to determine their interests, apply self-learning algorithms to filter products according to certain criteria, and predict (recommend) things that might be of genuine interest to the user.

There are different approaches to processing information to create recommendations. The content-based (CB) approach recommends products or things similar to those that the user prefers. It evaluates the content of items to generate new recommendations. The collaborative-filtering (CF) approach recommends products to the consumer that people with similar preferences have previously liked, creating a variety of recommendations. It is also possible to combine both approaches to create a new approach – a hybrid approach.

One approach is to identify recommendations that at first glance fully meet user expectations. This approach is called top-N recommendation, in which the algorithm should offer several items that will be of most interest to the user.

The top-N recommendation algorithm is a hot topic for the e-commerce industry [1]. Studies have shown that recommender systems have a positive impact on sales, and identifying several items of interest to an individual user can significantly increase business profits [2]. This work complements ongoing research

by showing that it is not just the similarity or variety of products presented that can arouse users' interest. Shoppers expect to see different types of recommendations depending on the stage of the purchase process they are in.

Gradient boosting is an algorithm that builds simple predictors that improve the objective function. That is, instead of building a complex model at once, many small models are built one by one.

The model training process is as follows: first, the training parameters are parsed by the user. Then the transferred parameters are validated. Next, the data is loaded and a grid is built to quantize the numerical data. This step is necessary to make the learning process faster.

The categorical parameters are processed differently: they are hashed from the very beginning, and then the hashes are numbered from zero to the number of unique category values to make reading the categorical data faster.

Next, the training cycle is launched – the main cycle of machine learning, in which trees are built iteratively. After this step, the model is exported. All steps of the learning process are shown in Fig. 1.

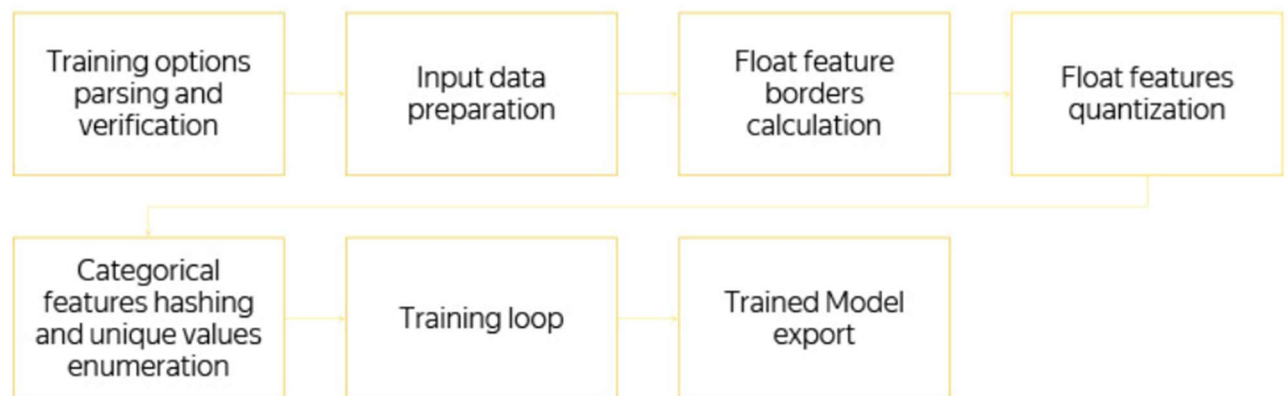


Fig. 1. Steps of model training using the CatBoost algorithm

The training cycle consists of four steps:

1. building one tree;
2. checking what increase or decrease in quality it gives;
3. checking whether the overlearning detector has been triggered;
4. saving the snapshot.

Training a single tree is a cycle through the levels of the tree. In the case of categorical training parameters or ordered boosting, a permutation of the data is randomly selected from the very beginning. After that, each branch of the tree is recalculated. Next, the "good" branches of this tree are selected.

The cycle of selecting tree levels is as follows: from the very beginning, a bootstrap is made - weighing objects, after which only the selected objects will be used to build the tree. The bootstrap can also be recalculated before selecting each branch. Next, the derivatives are aggregated into histograms, and this step is performed for each branching candidate. The histograms are used to estimate the change in the objective function that will occur if a given candidate is selected. The candidate with the highest score is added to the tree. After that, the statistics are calculated using this selected tree, the values in the leaves are updated, the values in the leaves for the model are calculated, and the cycle moves on to the next iteration.

Conclusion. This work has elucidated the fundamental techniques employed in harnessing the power of machine learning to facilitate the search for similar products aligned with user preferences. We have not only described these methodologies but have also developed a novel model that can serve as a foundational building block for the creation of robust information systems dedicated to product discovery. The implications of our research extend beyond the confines of our specific application, as the insights and methodologies presented here can be seamlessly integrated into various domains of e-commerce.

The development of this model opens up exciting opportunities for the implementation of information systems designed to streamline and enhance the search for products in any field of e-commerce. By leveraging machine learning, we can empower users to find relevant items more efficiently and effectively, ultimately improving the overall user experience and potentially boosting conversion rates for e-commerce businesses.

References

- [1] Cremonesi, P., Koren, Y., and Turin, R. (2010). Performance of recommender algorithms on top-n recommendation tasks. In Proceedings of the RecSys 2010.

[2] Hurley, N. and Zhang, M. (2011). Novelty and diversity in top-n recommendation analysis and evaluation. ACM Transactions on Internet Technology.

Дарій Бензак, Юрій Процик*

магістрант кафедри КН, НЛТУ України, Львів, Україна.

* к.ф.-м.н., доцент кафедри КН, НЛТУ України, Львів, Україна.
protsyk@nltu.edu.ua.

Розроблення системи прогнозування результатів матчів зі снукеру з використанням методів машинного навчання

Анотація – В роботі розглядаються питання, пов’язані з розробленням системи прогнозування результатів матчів зі снукеру на мові програмування Python з використанням методів машинного навчання. Описуються особливості реалізації засобів збору інформації про результати матчів, попередня обробка даних, налаштування декількох моделей машинного навчання та оцінка якості роботи відповідних алгоритмів.

Keywords – снукер, система прогнозування, машинне навчання, класифікація, Python.

Вступ. Снукер, як один із найпопулярніших видів більярду, завоював серця мільйонів шанувальників по всьому світу. Ця стратегічна гра, заснована на точності та вмінні прогнозувати ходи, викликає захоплення як у професійних спортсменів, так і у любителів. Розвиток технологій та поширення методів машинного навчання приводять до новаторських підходів у різних сферах, включаючи прогнозування результатів спортивних подій.

Задача прогнозування результатів матчів зі снукеру може бути сформульована як задача бінарної класифікації, мета якої полягає в тому, щоб на основі вхідних ознак про матч (наприклад, статистика гравців, рейтинг, історія зустрічей) визначити ймовірність перемоги одного з гравців або класифікувати матч як перемога чи поразка для конкретного гравця.

Процес бінарної класифікації включає наступні кроки:

Збір даних: зібрати вхідні дані про матчі зі снукеру, включаючи ознаки гравців, їхню статистику, рейтинги, історію зустрічей та інші відповідні дані.