

A Work Project, presented as part of the requirements for the Award of a Master's degree in
Business Analytics from the Nova School of Business and Economics.

Drivers and prediction of organic search engine CTR

The impact of differently phrased titles

ERIK FUBEL

Work project carried out under the supervision of:

Qiwei Han, PhD

25-01-2023

Abstract

33% of web traffic in the \$5.7 trillion e-commerce industry originates from organic search engines. Thus, website providers benefit from understanding the drivers of click-through rates (CTR) on organic searches. However, existing literature focuses on position as the primary CTR influence, disregarding other result page characteristics. To solve this problem, we conduct an elaborate data analysis and determine suitable CTR prediction modeling techniques. We discover novel patterns impacting CTR and find tree-based models to outperform state-of-the-art deep-learning models.

Furthermore, we conduct an NLP-based analysis of result titles and show that particular formulation patterns can significantly influence a result's CTR.

Keywords

organic click-through rate, CTR prediction, e-commerce, SERP features, NLP

Acknowledgements

We are very grateful for the continued support of our supervisors, Qiwei Han and Maximilian Kaiser. Additionally, we highly appreciate Grips enabling us to conduct this work by providing the extensive data set used.

This work used infrastructure and resources funded by Fundação para a Ciência e a Tecnologia (UID/ECO/00124/2013, UID/ECO/00124/2019 and Social Sciences DataLab, Project 22209), POR Lisboa (LISBOA-01-0145-FEDER-007722 and Social Sciences DataLab, Project 22209) and POR Norte (Social Sciences DataLab, Project 22209).

1. Introduction

Over the past years, the e-commerce business has seen extraordinary growth, with the Covid-19 pandemic additionally boosting its expansion. In 2022, Statista estimates global e-commerce sales at 5.7 trillion US-Dollars (Chevalier 2022a) which exceeds the 2021 GDP of the world's third-largest economy, Japan (O'Neill 2022). With this immense potential, online businesses are flocking the market, trying to gain a share of the steadily growing pie of e-commerce revenue. This revenue is created through sales on a website leading to the ultimate challenge of generating as much traffic on a website as possible. A prominent way to accomplish this is search engine optimization (SEO), an online marketing strategy with the goal of achieving the highest possible traffic for a website (Ledford 2009).

For e-commerce, 33% of the total web traffic is estimated to originate from organic traffic (Chevalier 2022b). Organic traffic refers to cases in which a user searches a keyword in a search engine and arrives at a website by clicking on one of the non-ad results. Because of the high relevance of this organic traffic, it is critical for website providers to improve their chances of generating clicks through search engines. In general, the clicks to a website on a search result page for a specific keyword can be formulated as

$$(1) \textit{Clicks} = \textit{Impressions} * \textit{CTR}$$

where CTR refers to the click-through-rate of a result (Google 2022a). Accordingly, to achieve higher web traffic, SEO practitioners want to increase at least one of both factors. While the impressions for any keyword can be approximated with the search volume available (Semrush 2022a), the CTR is only available to a website owner for keywords they already rank for. Thus, website providers can easily understand which other keywords are most promising to rank for but usually have very little insight into how much of the volume will

convert into actual traffic to their website. Providing an understanding of what influences CTR can overcome this hurdle and enable companies to receive more website traffic by achieving a higher CTR.

To increase CTR, SEO efforts tend to focus on the position of a website on a result page, as top-ranked websites receive more user attention and clicks (Lewandowski, Sünkler, and Yagci 2020). To rank high on a result page, a website needs to be considered as relevant for the particular keyword. To achieve this, SEO practitioners concentrate their efforts on matching a website's content and metadata to the keyword. (Cui and Hu 2011; Shih et al. 2013; Bala and Verma 2018; Ziakis et al. 2019; Das 2021; Olson et al. 2021)

This approach of focusing only on the ranking of a website on a result page to increase CTR appears one-dimensional, as it ignores how other characteristics of a search page directly influence user behavior leading to the desired clicks.

These characteristics can be grouped into four categories potentially influencing CTR:

1. The *position* of a result on a particular result page.
2. *Keyword* characteristics: These characteristics can relate to a keyword itself (e.g. length or content), the intent of the search, or the accompanying metadata about the result page, such as the search volume or difficulty of ranking high for a keyword.
3. Characteristics of a *result*: Such include URL, title, or description of a result.
4. *SERP features*: These are visual elements containing information from organic search results and aim at making Google result pages more engaging (Google 2022b). These increasingly prominent elements of a search engine result page (SERP) can, for example, take the form of an instant answer to a user's question, a knowledge panel on the right of the result page, or an image next to a result. Appendix I shows an illustration of SERP features.

The categories demonstrate that apart from the position, a website provider can influence various variables to gain users' attention. In order to do so and accurately predict CTR, however, one first needs to understand the characteristics' effects and dynamics. This opens two research questions:

RQ1: How do these characteristics influence CTR, and how do they impact the importance of a website's position?

RQ2: Which result page characteristics and machine learning models can be used to accurately predict the organic CTR of a website on a Google result page?

In this work, we extend the existing analysis of position as the main direct influence on CTR to additional keyword, result, and SERP feature characteristics summarized in a unique dataset. An analysis including these characteristics provides a novelty in CTR research. In addition to that, we will identify practical CTR prediction model techniques to apply to this data. This enables ex-ante evaluation of the effectiveness of SEO efforts by estimating expected website traffic through predicted CTR and commonly known impressions.

In the following, the second chapter will provide an introductory analysis of related publications on CTR influences and CTR prediction models. The third chapter introduces the used dataset and applied methodology in detail. After presenting the findings of an exploratory data analysis in chapter 4, chapter 5 examines which CTR drivers have the most predictive power and what models perform best for the prediction of CTR.

2. Related works

When raising the question of how to predict CTR and which features impact it, one is mainly looking at two related research fields. First, current research into drivers of organic CTR, and

second, recent research on models and data used to predict CTR. Therefore, this section will provide a high-level literature analysis of the former.

2.1 Research into drivers of organic CTR

Search engine optimization (SEO) is arguably one of the most wide-ranging online advertising strategies today (Olson et al. 2021; Kingsnorth 2022). According to the Oxford dictionary, search engine optimization is defined as ‘the process of maximizing the number of visitors to a particular website by ensuring that the site appears high on the list of results returned by a search engine’ (Stevenson 2010). The definition stresses the direct relationship between traffic and a website's position on result pages. This importance of position is also expressed by the vast majority of research on SEO marketing. (Cui and Hu 2011; Shih et al. 2013; Bala and Verma 2018; Ziakis et al. 2019; Das 2021; Olson et al. 2021). SEO achieves higher ranks by precisely tailoring a website's content to a given keyword, for example, by adding new and qualitative blog posts with the aim to appear more relevant to search engines (Das 2021). Other aspects such as keyword, result, and SERP feature characteristics, if considered, are only analyzed regarding their influence on the position, but not their direct influence on CTR.

To conclude, previous research sees the position as the primary driver of CTR and the most effective way to gain more user attention. We contribute to this perspective by introducing additional characteristics of result, keyword, and SERP features to the analysis of the directly influencing factors of CTR. In the following, we put a special emphasis on SERP features as they are one of the main visual elements that users perceive on a result page, next to organic and paid search results. We will do so by first outlining how they have been described in related works and to what criticism they are linked in public.

While the literature focuses on a website's rank, SERP features are receiving relatively little attention, with little literature researching isolated SERP features. An example is Sam-Martins (2020) analysis of the *featured snippet*, which in turn relies on information found on blogs from influential SEO companies like Semrush. While there is little research on the impact of SERP features, there are a number of influential marketing blogs that advertise the importance of appearing in SERP features (Moz 2022; Semrush 2022b; Wheelhouse 2022). Most blogs argue that SERP features benefit websites featured in them and harm the CTR of websites that appear next to SERP features without being featured (Moz 2022, Wheelhouse 2022). Google itself advertises significant improvements in click-through rate, visits and time spent on websites when chosen to be shown in SERP features (Google 2022b). Critics, however, state that by adding additional features to its result pages' and leveraging information originally published on third-party websites, Google reduces reasons to leave its ecosystem (Tober 2022). This can harm the publishers of the original information that rely on traffic to their websites. Zero-click searches, which are searches where a user does not leave Google's ecosystem, amount to 25.6% of all searches, according to a non-academic study (Tober 2022). This shows that by analyzing the effect of SERP features on CTR, this work has the potential to shed light on the intransparent role of SERP features in Google's search engine ecosystem and provide first empirical research on the SERP features' impact on website performance.

2.2 CTR prediction research

This chapter discusses achievements in previous research on CTR prediction and points out the difference in approaches between previous publications and this work. These differences lie in the nature of the prediction, i.e. whether it is a classification of clicks from individual users or a regression for aggregated CTR, and in the nature of clicks, i.e. whether they are organic or advertisements.

The problem formulation addressed by the vast majority of CTR research is: “Will user X click item Y?”. More technically, the underlying problem is a binary classification problem, i.e. predicting yes or no. The item in question could be an ad on a search engine result page but also an article in an online shop (Yang and Zhai 2022).

This poses a fundamental difference to the problem formulation of this work, as this work deals with a regression problem. It predicts a continuous value based on aggregated data while existing research conducts a classification for a single observation. This difference is also visible in the data used to make predictions. The dataset provided by Grips has high density and low sparsity. Opposingly, the datasets of the existing research, such as the commonly used Criteo dataset, usually contain the columns *user* and *item* with extremely high cardinality for each (Yang and Zhai 2022; Juan et al. 2016; Zhou et al. 2018). To make these columns usable for machine learning models, the values are usually one-hot-encoded, resulting in extremely sparse and high-dimensional feature vectors. Consequently, models are built in such a way that they can incorporate this high dimensionality and sparsity (Zhou et al. 2018).

Recent literature focuses on creating and exploring models suitable for such characteristics. At the beginning of user-item CTR prediction, the research focused on multivariate statistical models like linear regressions enhanced with polynomial features to include feature interaction (Wang, Suphamitmongkol, and Wang 2013; Yan et al. 2014). As these models show difficulties dealing with sparse data, the research progressed to models more suitable for this kind of challenge with a focus on factorization machines and later field-aware factorization machines (Ma et al. 2016; Pan et al. 2018; Yuchin et al. 2016). While these consistently outperform the multivariate models, the literature points out difficulties with generalization (Yang and Zhai 2022). With the rise of convolutional neural networks (CNN), research increasingly attempts to utilize them for CTR prediction (Chen et al. 2016;

Gharibshah et al. 2020). Based on Google's Wide&Deep recommender system architecture (Cheng et al. 2016), the introduction of DeepFM (Guo et al. 2017) marked a new milestone in user-item CTR prediction. DeepFM combines the previously well-working factorization machines with neural networks and lays the foundation for current research. The most recent literature focuses on more accurately modeling feature interactions and produces promising models, such as MaskNet (Wang, She, and Zhang 2021) and DeepLight (Deng et al. 2021), that continuously, though only slightly, outperform DeepFM. Yang and Zhai (2022) provide a detailed overview of the field of CTR prediction in the context of user-item interactions.

Next to the differences between individual click prediction (classification) and aggregated CTR prediction (regression) as described above, a further differentiation lies in the nature of clicks. There are some papers that focus on search engine CTR (Richardson, Dominowska, and Ragno, 2007; Graepel et al. 2010; Zhang et al. 2021). However, only the work of Richardson, Dominowska, and Ragno (2007) focuses on the prediction of aggregated CTR, while the other works predict individual, non-aggregated user interaction. Furthermore, all works focus on paid instead of organic results. This lack of research on organic CTR can be explained by the absence of publicly available data on this topic, as the data is one of the search engines' most significant assets, which therefore has no interest in sharing it.

To summarize, many papers focus on predicting whether an individual user will click on a specific ad and use this information as a recommender system. Additionally, some papers predict aggregated advertisement CTR. However, to the best of our knowledge, there are no publications that predict aggregated organic search engine CTR.

2.3 Implications

The analysis of the research on SEO marketing shows that the position is the most important driver of CTR. Therefore, it promises to play a critical role in predicting organic CTR. In

contrast, the effect of SERP features is theorized but barely analyzed. Therefore, it is essential first to analyze and fully understand the drivers of CTR and potential modeling approaches before predicting CTR. Therefore, in chapter 4, an in-depth exploratory data analysis was conducted to provide fundamental insights into the challenge of CTR prediction on organic search engine result data.

Additionally, while there is a large amount of literature on CTR prediction regarding user-item interactions, there is little research on aggregated search engine CTR prediction. Therefore, while it is insightful to apply the previously mentioned models to the data at hand, it is crucial to consider that the nature of the problem and the available data are different. This is likely to require distinct modeling approaches, which are tested in chapter 5.

3. Dataset & Methodology

This section provides an introductory overview of the used data and methodology applied to answer the research question.

3.1 Dataset

The dataset used for the analysis was provided by Grips, formerly Peekd, a German start-up that aims at ‘shining a light on the blind spots [in e-commerce sales] and create the most comprehensive map of online commerce for retailers and brands’ (Grips 2022). The data comprises historic Google search data for a given URL and keyword. One row can, for example, resemble metrics related to the performance of the URL ‘www.novasbe.unl.pt/en/’ for the searched keyword ‘nova university’. As such, the data is provided in a tabular, heterogenous format. It ranges over 43 domains, collected from US-based desktop searches between May 31st, 2022, and August 18th, 2022. In the raw dataset, 70 features are available for 79,853 data points.

Although the dataset was provided by Grips, it was enriched with features from Semrush, a company specializing in search engine marketing, and Google's Keyword Planner (see the source of each feature in appendix II.). The features cover various aspects, ranging from the searched keyword over metadata about the result page to result-specific information and which SERP features are shown. Finally, the data includes the click-through rate for each result which, in line with the research question, is determined as the target value.

The dataset, with its comprehensive set of features consisting of real-world, not publicly accessible data for dozens of e-commerce stores, signifies a novelty in available data.

3.2 Feature subsets

For a better understanding of the effect of position, keyword, result, and SERP feature characteristics on CTR, the variables in the dataset were grouped accordingly into four feature subsets. If applicable, subsets were enriched with self-generated features. The generated feature subsets form the basic structure for the analysis in chapter 4 and 5. A data dictionary describing all features can be found in appendix II. and a list of all features for each subset in appendix III.

Position - The position subset groups features that are exclusively position related. Those are the position of the result during the measurement, the monthly position average, and the difference in positions to the last measurement. While the position subset is technically also handled as a subset in the following analysis, it is important to note that the feature position plays an overarching role in the problem at hand. The literature review and findings of this work highlight that position is the most important feature for predicting CTR. Therefore, in the analysis, the position subset is often used in conjunction with other subsets.

Keyword - The keyword subset consists of the keyword itself and everything directly resulting from the keyword a user types into the search engine. This includes information about the keyword, such as the length of the keyword, information indirectly related to that keyword, like the competition or the search volume, and finally, the search intent. One generally differentiates between four intents: (i) The informational search, where a consumer is looking for information about a subject; (ii) the navigational intent in which a specific website is searched, e.g. a marketplace for a product; (iii) the commercial intent that implies a purchasing decision but it is not yet decided on the final transaction (e.g., “linux alternative”); (iv) the transactional intent where a consumer has an action in mind such as “buy iPhone”.

To quantify the complexity of each keyword, the Flesch reading ease score (Flesch 1948) was computed. Even though newer scores have been developed, Flesch’s reading ease score is used because it is commonly accepted and easily interpretable. It counts the average number of words per sentence and the average number of syllables per word. A high score corresponds to good readability, characterized by short sentences and short words.

SERP features - The SERP features in the dataset are binary features stating for each entry whether each SERP feature is present (1) or not (0). They can be divided into page and positional SERP features. Page SERP features describe whether a certain feature is present somewhere on the page. Positional SERP features describe whether a result is featured in a particular SERP feature. This can take different forms. For example, a *Review* can be shown beneath a URL, or a *Snippet* from the URL can be displayed in the *Knowledge Panel*. Therefore, if a positional SERP feature is present, the equivalent page SERP feature must be present, but not vice versa. The page SERP features also include information about the kind of ads present on a result page. We aggregated this information to a binary feature describing whether ads are displayed. To add a layer of interpretability, the dataset was additionally enhanced with the total count of SERP features both for page and positional features.

Result - The result subset includes all features related to the characteristics of one URL on a result page. While for the three other subsets the features available in the dataset are very comprehensive, available result-related features only cover the URL of the result but not other relevant aspects such as the title or description. To generate informative value out of the URLs, the associated domain is introduced as a one-hot-encoded feature, and two keyword-related features are computed: It is checked whether at least one word of the keyword appears in (1) the domain and (2) the URL.

3.3 Data preparation

Before looking into the data in more detail, it has to be prepared for further analysis. To do so, the Google data preprocessing guidelines were followed (Google 2022c). Duplicate rows and columns as well as columns that had only one unique value or were redundant were dropped. Since $CTR = \frac{clicks}{impressions}$, both ‘clicks’ and ‘impressions’ are very closely related to the target variable and are consequently dropped, avoiding collinearity and information spillage. Finally, column data types were changed in order to facilitate effective analysis. To minimize noise through observations of result pages that were shown to too few users, an impression threshold of 20 impressions per result was introduced. 58,898 rows remain for analysis. For the modeling, all non-binary features are standard-scaled to a mean value of 0 and standard deviation of 1.

4. Drivers of CTR (EDA)

This chapter aims at identifying drivers of CTR through an exploratory data analysis (EDA). We focus on non-obvious and novel insights as well as characteristics that are relevant for CTR prediction models. Thus, the EDA does not only create a basic understanding of the data but also as the foundation for selecting and examining predictive models in chapter 5.

The analysis is structured along the four feature subsets and dedicates a section to each. In addition to that, we analyze feature interactions and present a summary of the main findings and implications for the modeling.

4.1 Position - the most important predictor

There are two main findings regarding the *position* subset: First, the position is more relevant than other subsets in determining CTR. Second, results on position one receive much higher average CTRs than results on any other position.

In line with previous research (see chapter 2 and Iqbal et al., 2022, 4), our data shows that of all features, the position has the strongest impact on CTR. This is based on position and related features having the highest correlations with CTR. The average position per month has a Pearson correlation of 0.5 with CTR. Related features, such as the previous position, have a Pearson correlation of more than 0.4. This is significantly higher than the second most correlated variable with CTR ($\rho = 0.25$)¹.

Our findings show that the CTR on position one is much higher than on other positions, as many search engine users only read and click the first result. The result on the first position receives an average CTR of 0.25. This is 2.5 (3.6) times higher than the average CTR on the second (third) position. Therefore, our data supports the prevailing belief in the academic literature that optimizing the rank is crucial for improving a website's traffic. However, there is a considerable standard deviation of 0.19 for position one. Outliers on the first position reach up to 1.0 CTR, while a CTR of 0.0 is also possible. This can also be seen in the second and third positions, with a CTR range between 0 and 0.8.

This range translates to a significant variance reiterating the need to correctly predict CTRs, as even minor deviations in CTR can have significant click and revenue implications for

¹ ρ denotes the Pearson correlation coefficient

businesses and their SEO efforts. However, the strong variance also underlines the importance of other explanatory variables to explain strong deviations. Lastly, the results highlight position as the single most important feature, which inherits substantial implications for the modeling: We infer that the *position* subset is crucial for accurate predictions and that it can thus serve as baseline subset to use in models to analyze the predictive power of other subsets.

4.2 Keyword characteristics - the role of a keyword's complexity

Among keyword characteristics, the complexity score reveals a particularly noteworthy effect on CTR. For the complexity score introduced in chapter 3, the simplest possible keyword receives a score of 121, with the score decreasing towards a technically infinite negative value for more complex keywords.

Figure 1 shows the average CTR for keywords on position one depending on their complexity. Additionally, it distinguishes between branded² and unbranded keywords. The Figure shows that keyword complexity creates mainly two contrasting effects: For more complex keywords consisting of many and/or longer words³ (score below 80), the average CTR of branded keywords is significantly lower than for unbranded ones. This changes for values with a complexity score of more than 80, where simple branded queries lead to significantly higher average CTRs. For the most simple keywords, branded keyword results, on average, achieve a CTR of more than 50% higher than unbranded keywords. Additionally, a general tendency of decreasing CTR for simpler keywords can be observed for unbranded keywords.

The results imply that optimizing for position one makes sense for branded keywords that are simple and short, where i.e. the brand is likely to be a very prominent component. However, the results also hint that longer and more complex keywords including brand names can even

² A keyword, or search query, is branded if it includes a brand name.

³ To make the score more tangible: A search query consisting of five words with an average number of syllables per word of 1.8 would achieve a reading ease score of 49.5. An example of such a query is “buy cat food online now”.

harm the CTR of results on position one. This implies that users are more explorative for such keywords and tend to be more likely to consider results after position one.

The results also highlight feature interactions between the complexity of a keyword and whether it is branded, implying a need for prediction models to be capable of modeling feature interactions.

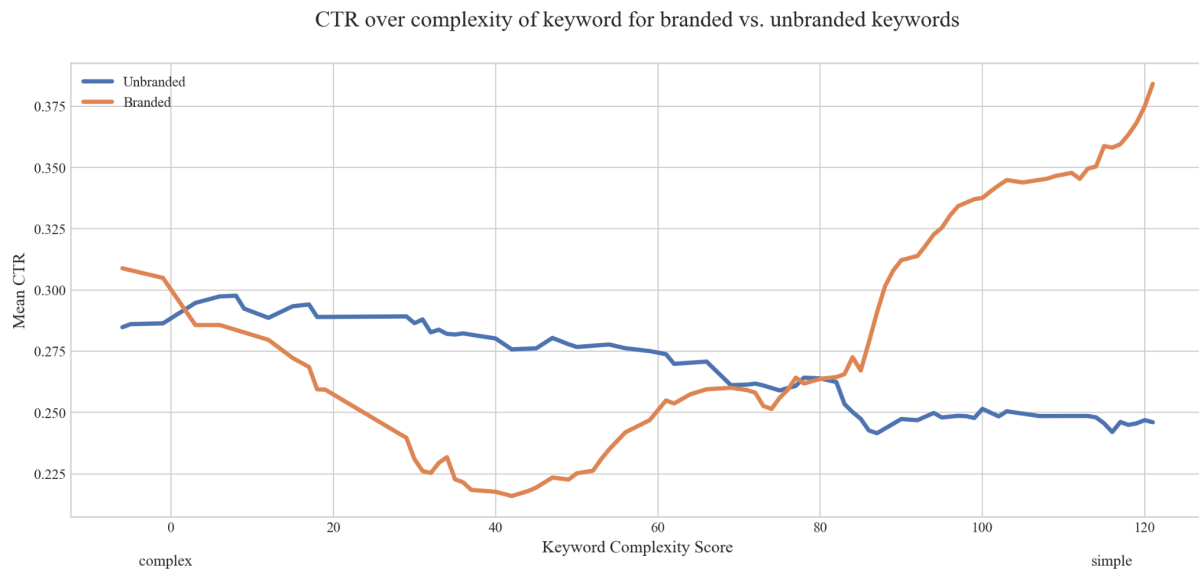


Figure 1: Mean CTR over keyword complexity score for branded and unbranded keywords. Mean CTR is calculated as a centered moving average with a total window size of 30. Keywords are more complex for low values and simpler for high values.

4.3 SERP features - variation in effects over positions

In the following analysis, we examine the impact of SERP features on CTR and outline implications for SEO decision-making concerning SERP features.

As the count of positional SERP features increases, i.e. a single result has more SERP features bound to it, the CTR tends to rise ($\rho = 0.22$), whereas with more page SERP features, there is a less clear trend, with CTR tending to decrease ($\rho = -0.11$).

In chapter 4.1, we have shown the importance of positions. When adding another dimension, namely the presence of a certain SERP feature, noteworthy tendencies arise. While for some SERP features the average CTR per position remains relatively constant regardless of the

presence of the SERP feature, for others, this presence has a big influence. Tolerating minor exceptions, there are two main patterns observable: With the presence of a SERP feature, first, the average CTR declines over all positions ('Lower'), and second, the average CTR declines for up to the first three positions but increases for subsequent positions ('Lower → Higher'). See Figure 2 for an example of each pattern category and SERP features it applies to.

Explanations for why these patterns arise for certain SERP features are not clearly identifiable. However, some features are considerably correlated with intents, e.g. *Image* with transactional search intent ($p = 0.29$). Consequently, when interpreting the results, we must keep in mind that effects might also originate from other feature categories, such as the search intents.

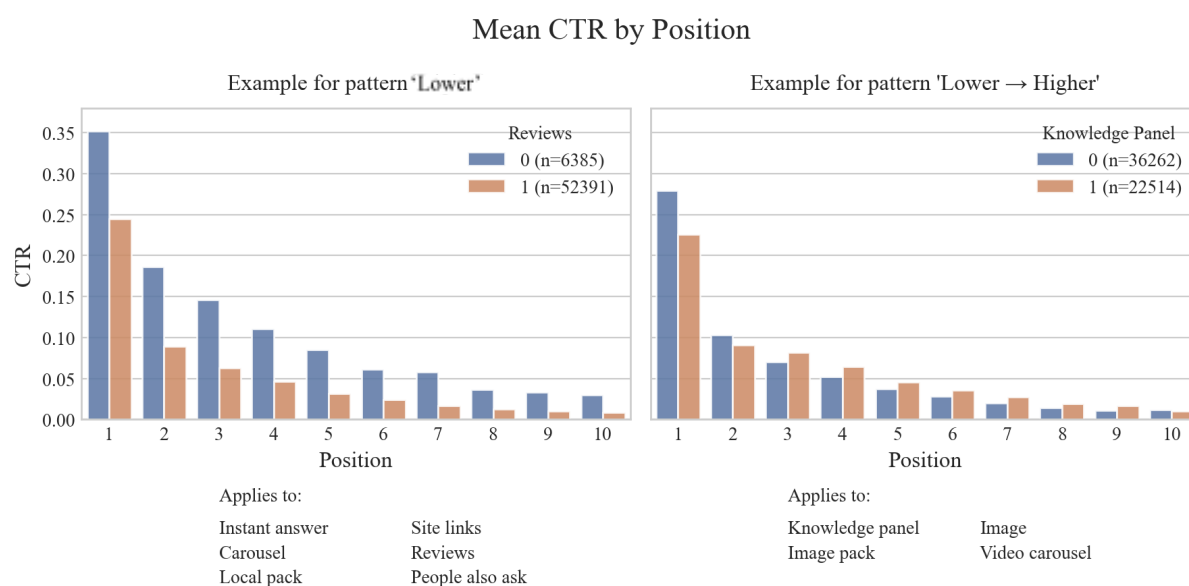


Figure 2: Mean CTR per position if a SERP feature is present (1) compared to when it is not present (0). Left plot shows pattern 'Lower' and right plot pattern 'Lower → Higher'. SERP features without a clear pattern or $n < 1000$ are not listed.

Although we have previously stated that ranking on a top position is of utmost importance for maximizing CTR, the aforementioned findings on SERP features spur the question of whether, in some cases, it is more important to rank on top positions than in other cases. While above, we see the impacts of SERP features isolatedly, in practice, SERP features appear together and thereby likely influence each other. Thus, in the following, we analyze the

importance of positions for each combination of page SERP features. To be able to compare different combinations better, for each combination with more than 200 occurrences, we normalize the CTR to the average value on the first position and then calculate the decay for each subsequent position (Figure 3). As a result, it becomes evident that patterns identified for individual SERP features continue across combinations of them. Among the top three combinations by count, for two, the CTR decays much faster, while for the other one, it decays much slower compared to the entire dataset. For practitioners, this implies that efforts towards ranking on the first position can be of much more value if some SERP features are present, while in the presence of others, a positioning within positions two to five might also suffice. For example, three differing SERP features can already double the maintained CTR on position three (see the difference between the green and blue combination in Figure 3).

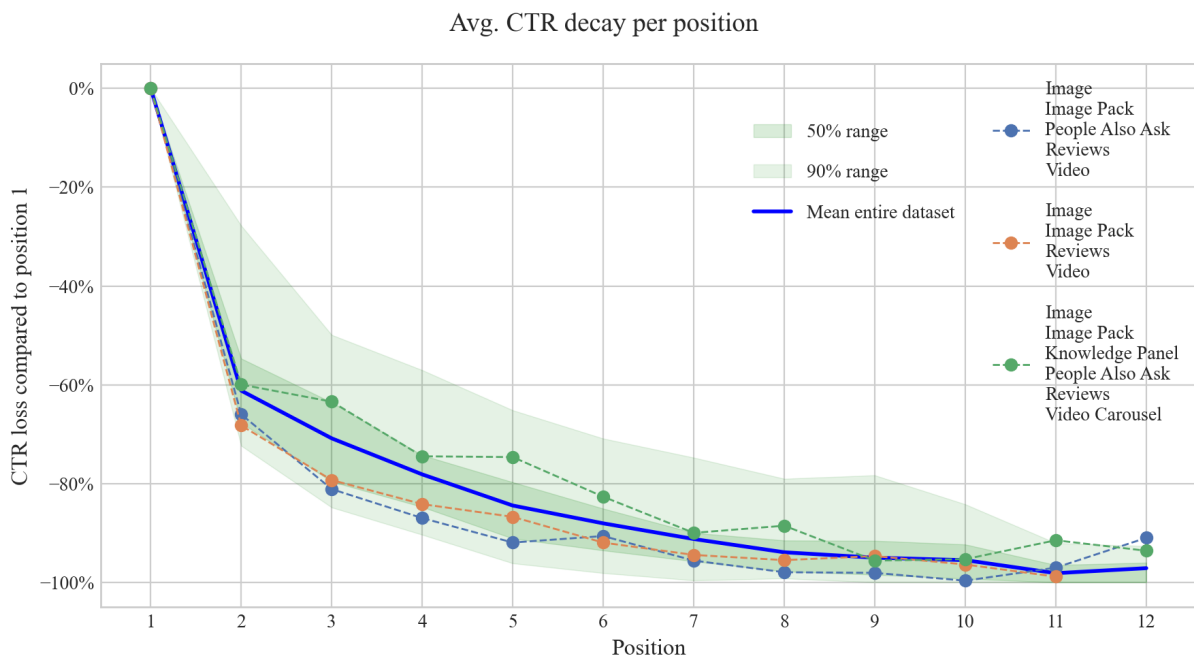


Figure 3: Loss in CTR per position. Values are avg. CTR loss for each position in % of avg. CTR on position 1. Only combinations with > 200 occurrences are considered (n=59). Dashed lines represent the top three most frequent SERP feature combinations. Shaded areas represent the range in which 50% and 90% of all values fall.

Together, the implications for the modeling are twofold. First, we show that impactful patterns can be observed on different levels of aggregation, i.e. when using the count of all SERP

features, looking at each feature separately, or looking at combinations of features. Thus, knowledge about SERP features should help models to more accurately predict CTR. Second, we can observe a high degree of interactions. The interactions go beyond the second degree, like the interaction between position and combinations of SERP features. Consequently, to make use of all information inherited in SERP features, models must be able to incorporate feature interactions higher than second degree.

4.4 Result characteristics - how domains & brand perception influence CTR

In this chapter, we focus on the effect of the domains on CTRs. The domains in the analyzed dataset are exclusively from e-commerce websites. They either belong to a manufacturer or a marketplace specialized in product categories like electronics or apparel. Therefore, the domains either represent a brand name or a marketplace and can often be seen as a placeholder for a product category.

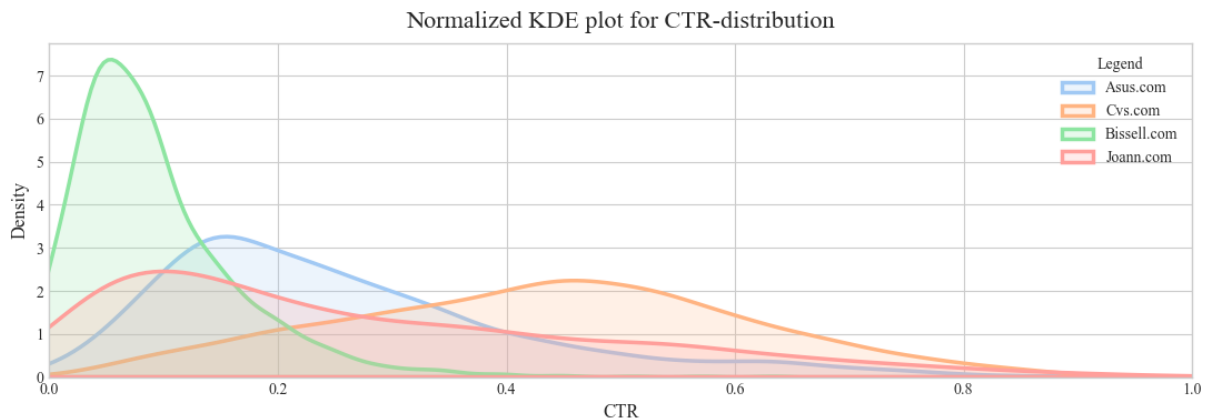


Figure 4: The Kernel Density Estimation (KDE) plot for four of the top five domains by count, normalized to unify the area under the curve.

Figure 4 shows that domains have widely different CTR distributions on position one. CVS, a well-known pharmacy, achieves the highest average CTR among domains that appear at least 1,000 times. While Bissel, a premium vacuum cleaner company, receives the lowest average CTR of these domains. The kernel density estimated distributions of CTR are right-skewed

for domains with low average CTR, whereas domains with high average CTR go along with a more uniform distribution. We find that there are no correlations with other variables that could clearly explain these patterns. This indicates that the brand name is still important in determining CTR to a website's traffic.

This is especially clear when looking at the top and the worst performing domains on position one, CVS and Bissell. Both of the websites are reached from search queries that include their brand name in 87% of cases for Bissel and 77% for CVS. In both cases, 98% of these queries have a transactional (clear purchasing) intent. However, even though both domains appear on position one, in search queries with a purchasing intent for a product of their brand, there is a large difference in CTR. 'cvs.com' achieves an average CTR of 26%, compared to 6% for 'bissel.com'.

A possible explanation could be a certain brand recognition by the consumers. While the pharmacy, CVS, and its online business is a respected marketplace, the premium vacuum cleaner company, Bissel, is known for its product but not as a primary marketplace, as third parties often offer lower prices. When displayed in Google, Bissel's website is in 70% of the cases enriched with the positional SERP feature *Review*. As the *Review* can also show prices, it is likely that third-party retailers show a lower price and drive traffic away. Additionally, searches for Bissel's products are 55% more likely to include ads than the rest of the data set. This presumably leads to more competition for users' attention.

Supporting the thesis for the high relevance of a website's brand, we find that the two domains with the highest average CTR ('hp.com' & 'cvs.com') are the only ones listed in the top 100 most powerful brands index in the US (Tenet Partners 2020), while the others do not appear.

Based on these findings, we conclude that domains are important in determining the CTR of a website and have high relevance for predicting it.

4.5 Feature interactions - SERP features & the brand name in keywords

In our previous analyses, we have uncovered several feature interactions across and within feature categories. In the following, we quantify the strength of feature interactions and reveal the strongest interaction's underlying forces.

Friedman and Popescu (2008) introduce the H-statistic to measure interactions of two or more features in an ML model. The H-statistic quantifies the interaction strength of two features by the difference in the explained variance of a decomposed partial dependence (PD) function and the observed PD function. The decomposed PD function assumes no interactions, whereas the observed PD function measures interactions. Thereby, it calculates the explained variance resulting from interactions.

$$(2) \ H_{jk}^2 = \frac{\sum_{i=1}^N [PD_{jk}(x_{ij}, x_{ik}) - PD_j(x_{ij}) - PD_k(x_{ik})]^2}{\sum_{i=1}^N PD_{jk}(x_{ij}, x_{ik})^2}$$

If the resulting score amounts to zero, all variance is explained by the individual 'contributions' of the features. Thus, there is no interaction between them. A higher value indicates that more variance is explained by the PD functions with interactions. This suggests that each individual PD function is constant and the effect on the prediction only comes through the interaction (Friedman and Popescu 2008).

The H-statistic has some limitations that need to be addressed when using it. First, it is non-deterministic, i.e. the same input leads to different results. Therefore, ten iterations were completed to see if it delivers stable results. Second, the H-statistic is very computationally expensive when applied to higher-degree interactions. Thus, we only analyze second-degree interactions. Third, the H-statistic only considers the interaction strength. It does not indicate whether the features have a relevant effect on the target variable or the kind of interaction. For

these reasons, it is necessary to analyze the feature combinations proposed by the H-statistic, to clearly explain their effects. In the following, we point out the most relevant findings.

As Figure 5 reveals, the strongest second-degree interaction of an H-value of around 0.6 appears for the features branded with the count of page SERP features. These features reveal an interesting interaction pattern, which is displayed in Figure 6. Branded keywords perform nearly three times as well as unbranded ones when few SERP features are displayed by Google. This effect reduces on result pages with more SERP features. When more than seven SERP features are present, the effect switches and unbranded keywords perform better than branded ones.

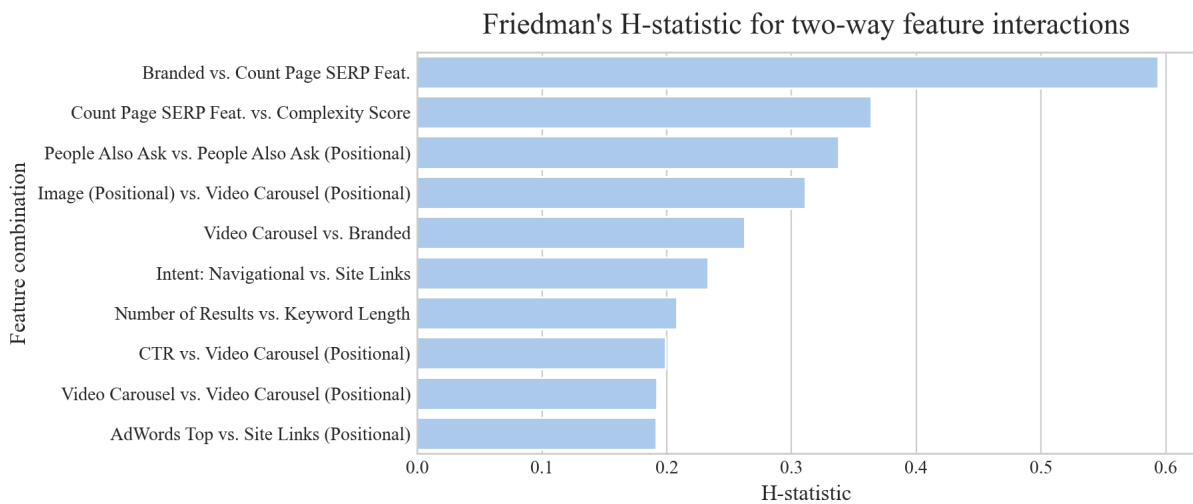


Figure 5: Friedman's H-statistic for two-way feature interactions. Highly correlated features and features with less than 500 occurrences were dropped, as they are likely to bias the H-statistic.

As this effect is not intuitively explainable, it is worth investigating it further. We find that result pages with few SERP features have very different properties, depending on whether the keyword is branded or not. The data also shows that result pages for branded keywords are much more likely to have navigational and transactional intent than those for unbranded ones. In contrast, unbranded result pages with few SERP features are more often informational than

their branded counterparts. To summarize, branded result pages with few SERP features have a high CTR, driven by purchase-oriented users who are looking for a specific (branded) website.

The implications are two-fold. First, branded search queries are more likely to generate website traffic if they lead to result pages with few SERP features than to result pages with many SERP features. Second, for unbranded search queries, the number of SERP features is not similarly relevant and a medium CTR can be achieved regardless of them.

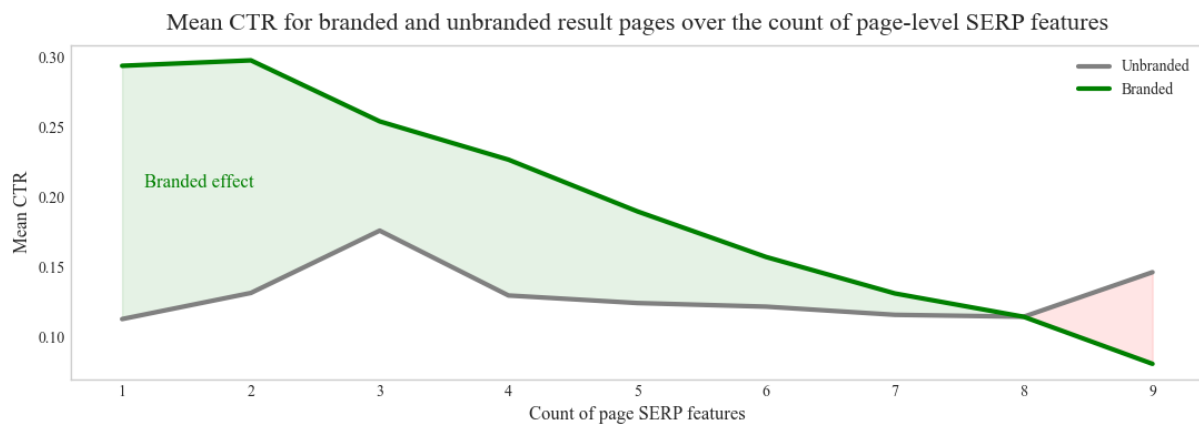


Figure 6: Mean CTR for branded and unbranded result pages over the count of page-level SERP features.

4.6 Findings & implications summary

Next to the detailed analysis of individual drivers of CTR above, there are two main implications based on this data analysis. First, within each subset there are features that have a relevant effect on CTR. Therefore, we will consider all of them as model inputs. Second, we revealed many interactions between features. Most features interact with position, which we either explicitly stated or we analyzed features on only one position, to isolate the feature effects. This shows the high relevance of the position for CTR. Additionally, features within and in between subsets interact with each other. These feature interactions imply that the tested models should either implicitly or explicitly handle feature interactions.

5. Modeling

In the following, different models for the prediction of CTR are tested and the predictive power of the four feature subsets is examined. First, we introduce the tested models. Then, the modeling methodology is outlined. Lastly, we analyze the modeling results.

5.1 Examined models

The different feature subsets are tested over different models. As pointed out in 2.2, the underlying problem is a regression problem. Consequently, the tested models are regression models. Furthermore, these tested models were selected considering their relevance and performance in recent literature as well as their applicability to the dataset and its characteristics as pointed out in the EDA. Additionally, linear regression models are applied as a benchmark due to their simplicity. Three model categories emerge: Linear regression models, neural network regression models, and tree-based regression models.

5.1.1. Linear regression models

Linear regression attempts to model the relationship between two or more variables by fitting a linear equation to observed data (Uyanık and Güler 2013). This allows for fast computation and easy interpretability. However, linear regression models assume the predictor variables to be independent from one another and linearly related with the independent, predicted variable (Su, Yan, and Tsai 2012). Furthermore, linear regression models are prone to outliers and overfitting the training data in the case of a larger number of features (Filzmoser and Nordhausen 2020). As a result, the model can be simplistic and unable to capture real-life complexities in data. Based on the dataset's characteristics outlined in chapter 3, the limitations also apply to the dataset at hand. There are strong feature interactions, i.e., the predictor variables are not independent, and relations with CTR are not always linear.

Regardless of these drawbacks, we will use ordinary least squares regression (OLS), the most popular linear regression model (Su, Yan, and Tsai 2012), as our baseline model. In an OLS model, the squared errors are minimized (De Gryze, Langhans, and Vandebroek 2007). Using this model allows us to compare more elaborate models with a solid, proven statistical framework. Suppose a model performs worse than OLS linear regression, despite the shortcomings and violated assumptions of linear regression described above. In that case, we can assume that it is not suitable for the dataset at hand. We also adapt the OLS model by adding feature interactions as the product of the values of two features (Poly2). As in this way, the number of features grows immensely, we only allow feature interactions with a pearson correlation coefficient > 0.05 in the training data set to be used by the model. Potential overfitting will be addressed using a third variant, a Ridge regression model, which implicitly performs feature selection to avoid overfitting using L2-Norm.

5.1.2. Tree-based regression models

Decision tree models are based on a series of if-then-else decision rules, resulting in tree-like structures, and are already used since the 1960s (Morgan and Sonquist 1963; Messenger and Mandell 1971; Quinlan 1986). Decision trees are enhanced in state-of-the-art models that use ensembles, i.e. a combination of decision trees to make a joint prediction (Omer and Rokach 2018). The considered tree-based models in this work are all ensemble models based on decision trees. The advantages of such models lie in them not requiring specific distributions in the data, being fast without preprocessing, and they are capable of modeling feature interactions (Agarwal et al. 2022). In addition, on medium and small-sized datasets, tree-based ensemble models often outperform more elaborate state-of-the-art models such as neural networks (Grinsztajn et al., 2022; Hancock and Khoshgoftaar, 2020).

Hence, tree-based ensemble models can address both the interactions and limited sample size

of the available dataset, making them very promising candidates.

We will test the top three performing tree-based models for tabular data, found in a recent comparative benchmark by Grinsztajn et al. (2022), namely Random Forest (Breiman 2001), Gradient Boosting Decision Trees (GBDT) (Friedman 2001) and XGBoost (Chen and Guestrin 2016), as well as CatBoost, another tree-based model achieving high scores in benchmarks (Prokhorenkova et al. 2018).

Random Forests make use of bagging, an approach where the results of multiple independent models, each trained on a subset of the data, are combined (Breiman 1996). Applied to Random Forests, this means that to get a prediction, the results of several randomized decision trees are aggregated through averaging (Breiman 2001). This model is continuously praised for its versatility, speed, and ease to use (Biau and Scornet 2016).

Contrary to bagging models, boosting models build decision trees sequentially and learn from the residuals of their predecessors (Freund and Schapire 1996). A variant of this are *gradient* boosting models, in which each tree, instead of predicting the target itself, predicts the error of its predecessor. Friedman (2001) laid the foundation for this variant with GBDT. Empirical studies indicate that gradient boosting models perform comparatively well on heterogeneous data, but do poorly on homogenous data (Hancock & Khoshgoftaar 2020).

Based on the concept of gradient boosting, Chen and Guestrin (2016) proposed XGBoost, which makes use of advanced regularization techniques (L1 & L2), improving its generalization capabilities. In addition, the training of XGBoost is parallelizable, making it much faster. XGBoost remains the go-to tool for most practitioners and data science competitions (Kossen et al. 2021).

Compared to the other examined gradient boosting models, CatBoost optimizes decision trees for categorical variables through two ways. First, the use of ordered target statistics allows for

efficient handling of categorical features with high cardinality. Second, ordered boosting prevents prediction shift, referring to a phenomenon where what the model learns in the training set is not reflected in the testing set (Prokhorenkova et al. 2018).

5.1.3. Neural network regression models

While Deep Neural Networks (DNN) achieve unprecedented performance on tasks related to homogeneous data (e.g. image, audio, and text data), heterogeneous, tabular data is still deemed an “unconquered castle” for DNNs (Borisov et al. 2022; Kadra et al. 2021; Hancock and Khoshgoftaar 2020). However, a recent comparison between the performance of traditional machine learning models and DNNs on tabular data by Borisov et al. (2022), suggests that while for most cases boosting models deliver best performance, on a dataset with more than 10 million samples, DNNs can achieve similar or even better performance. Furthermore, as outlined in chapter 2.2, the best-performing models in recent CTR prediction research incorporate neural networks. Based on these findings, we test a neural network structure especially designed for tabular data, TabNet (Arik & Pfister 2021), as well as Wide&Deep (Cheng et al. 2016) and DeepFM (Guo et al. 2017), neural networks suggested by recent CTR prediction research.

TabNet has been designed for tabular input data and uses the concept of sequential attention to perform row-wise feature selection. This enables it to use its learning capacity only for the most relevant features. Furthermore, TabNet offers a better interpretability than boosting models (Arik and Pfister 2021).

As highlighted in 2.2, CTR prediction research has evolved to currently suggest models that combine a wide and a deep component. These models usually combine the outputs of a deep neural network (deep component) with those of a linear model (wide component) in a common activation function. This model architecture was first proposed by Google as

Wide&Deep (Cheng et al. 2016) and its core strengths lie in its ability to both memorize and generalize (Jais et al. 2019).

State-of-the-art CTR research extends the idea of wide and deep models and adapts them to the specific needs of ad-CTR prediction. Although being outperformed in specific domains, as for example by MaskNet (Wang et al. 2021) on the Criteo dataset, DeepFM (Guo et al. 2017) is still one of the most commonly applied best-performing models in ad-CTR prediction. DeepFM embeds sparse features and interprets the deep component as a neural network while the wide component incorporates a factorization machine. This allows it to capture both two-dimensional feature interactions in the wide component and higher-dimensional interactions in the deep component. As a result, its advantages lie in the flexibility to, in a unified framework, capture low-level feature interactions explicitly and high-level feature interactions implicitly (Guo et al. 2017). However, it also lacks interpretability and imposes high computational complexity (Yang and Zhai 2022).

The implications for organic CTR predictions on dense search engine data are two-fold: On the one hand, DeepFM is capable of capturing feature interactions very well, which is beneficial. On the other hand, it is highly specialized in dealing with sparse feature vectors, which is not necessarily required in the case of result page data and might harm predictions. Furthermore, neural networks often rely on large samples for model training to generate quality predictions (Alwosheel et al. 2018), which is not the case for the available dataset (only medium-sized; ~60k samples).

5.2 Methodology

5.2.1 Subset combinations

To assess each subset's predictive power, different subsets were tested on the models. Testing a model on a subset combination means that all features of the combined subsets are used for

model prediction. Based on both the literature review in 2.1 and our analysis highlighting the position as the single most important feature, the *position* subset serves as a baseline subset. It is then combined with the other subsets.

First, the *position* subset is combined solely with each of the other subsets such that the additional predictive power of each subset can be assessed isolatedly. The *result* subset, however, inherits a structural limitation: Because it mainly consists of one-hot-encoded domains, it is very specialized to the underlying dataset. As such, if a prediction for a new domain was to be made in a production environment, the one-hot encoded domains could not contribute to making a more accurate prediction. Thus, the *result* subset has very limited generalizability, and scores need to be interpreted cautiously.

Based on this limitation, we first check the predictive power of all generalizable subsets combined, i.e. *position* + *keyword* + *SERP features*. To get an understanding of the full potential, we lastly also include the *result* subset to combine all available meaningful features. As a result, 6 subset combinations are tested on each model presented in 5.1.

5.2.2 Evaluation metric

In order to compare the effectiveness of different feature subsets and models, we need a common evaluation metric. All models are evaluated using the Root Mean Squared Error (RMSE), defined as

$$(3) \text{ RMSE} = \sqrt{\frac{1}{N} \sum_{i=1}^N (y_i - \hat{y}_i)^2}.$$

This metric allows for the penalization of larger errors which is desirable in the present business case as larger errors can lead to a complete miss investment of SEO efforts. As companies are likely to optimize their SEO marketing for the keyword with the highest

expected clicks, it would be a particular financial loss if this estimation is wrong. In addition to that, the RMSE can be used as a quick optimization metric across all model categories, as opposed to, for example, the mean absolute error (MAE).

5.2.3 Hyperparameter tuning

To examine each model's full potential, we conducted hyperparameter tuning where reasonable. The hyperparameters for all linear models were kept constant. Tree-based models were tuned based on the hyperparameter sets suggested by Gorishniy et al. (2021) and Grinsztajn et al. (2022). As the tuning algorithm, Bayesian Optimization (Snoek et al. 2012) with 100 iterations is used since it is reported to perform superior compared to random search (Turner et al. 2021). For the comparison of scores of different hyperparameters, 5-fold cross-validation on the training set was used. Please find the full documentation on tested parameter ranges in appendix IV and the resulting parameters in appendix V.

The training of neural networks was limited to 200 epochs with early stopping after 10 epochs without improvement of the RMSE of the validation set. During the training and evaluation of neural networks, we noticed a large variance in the resulting test scores. To compensate for this, we averaged the scores of 10 different experiments to receive the final test score per neural network based model. However, by doing so, an extremely large number of experiments would need to be performed during hyperparameter tuning, resulting in extraordinarily large computation times⁴. Due to this, we could not conduct in-depth hyperparameter tuning for neural networks. Nevertheless, we performed some manual experiments on hyperparameters based on the domain knowledge gained during the EDA. While for most combinations of features subsets, the altered parameters did not lead to significant improvements, scores for the combination of all subsets were drastically improved

⁴ Example: For setup [iterations = 100; No. optimized subsets = 6; CV = 5; experiments per score = 5; training time = 30 sec], the optimization would take 5.2 days per model

(see appendix VI). Thus, for all neural networks, the default hyperparameters were used for all combinations of feature subsets except the combination of all subsets.

5.3 Results & interpretation

In the following, we analyze the model results (Table 1), first concerning how different models perform, and second considering the predictive power each subset incorporates.

5.3.1 Comparison of model performance

Clear patterns arise both between each model and the categories they are grouped in. Apart from summarizing these patterns, we examine whether the assumptions regarding the applicability of models, made in chapter 5.2, hold true.

On average, linear models perform worse than other categories of models, as expected. As described in 5.1.1., we tested the Poly2 model to check the assumption that model performance improves when we allow interactions between features. This assumption holds true as the Poly2 model significantly outperforms the OLS/baseline model. An exception to this is the case when all feature subsets are used, where the error for Poly2 increases drastically. A potential explanation for this is that with more features, the number of interactions grows quadratically, thus leading to so many features that the linear regression model heavily overfits. Ridge only performs marginally better than OLS, indicating that the OLS model works well with the number of features and overfitting is not an issue.

Although mostly outperforming linear models, models based on a neural network architecture cannot keep up with the performance of tree-based models. While scoring very closely to tree-based models for datasets with few features, such as only position, the differences enlarge with the presence of more features and accompanied complexity (see Table 1).

		RMSE of subset combinations					
		Position	Position + Keyword	Position + SERP Feat	Position + Result	Position + Keyword + SERP Feat	Position + Keyword + SERP Feat + Result
Linear Regression	OLS	0.153	0.145	0.145	0.132	0.137	0.125
	Poly2	0.143	0.129	0.131	0.119	0.122	2.554
	Ridge	0.152	0.144	0.144	0.131	0.137	0.123
Tree Based	Random Forest	0.137	0.108	0.120	0.108	0.100	<u>0.088</u>
	GBDT	<u>0.134</u>	0.109	<u>0.122</u>	0.109	0.108	0.091
	XGBoost	0.133	<u>0.104</u>	0.120	<u>0.106</u>	0.098	0.083
	CatBoost	<u>0.134</u>	0.103	0.120	0.105	<u>0.099</u>	0.083
Neural Network	TabNet	0.137	0.121	0.130	0.111	0.118	0.102
	Wide&Deep	0.137	0.123	0.152	0.110	0.128	0.109
	DeepFM	0.136	0.135	0.137	0.113	0.133	0.114

Table 1: Benchmark results on different feature subsets. The top results for each feature subset are **bold**. We also underline the second-best results.

Furthermore, TabNet, a model resulting from research focusing on using neural networks for inference on tabular data, performs better than models resulting from CTR research (Wide&Deep and DeepFM). Additionally, the strength of CTR research models on sparse data and weakness on tabular data is reflected in the scores. For sparse subset combinations (*position + result*) their error is only 4% above the error of tree-based models, which increases to 19% for dense subset combinations (*position + keyword*).

It is important to note that even without hyperparameter tuning, tree-based models achieved consistently lower errors than the neural networks (see appendix VI). Based on this, we assume that the considerably worse performance of neural networks is mainly rooted in the model’s architecture instead of the lack of hyperparameter tuning.

Together, this confirms the previously highlighted inapplicability of models resulting from

CTR prediction research to the dataset due to lack of sparsity. Additionally, the neural networks likely suffer from a relatively small number of observations.

Tree-based models perform clearly the best. It is noticeable that the boosting models mostly perform better than the bagging model. While the random forest performs better than GBDT in half of the feature subsets, it is worse than XGBoost and CatBoost in every subset. XGBoost and CatBoost are the two top-performing models for every feature subset. For two subset combinations, they achieve the same RMSE, and for the remaining subset combinations, one outperforms the other just by 0.001 RMSE. These very similar performances spur the question of which model to use under secondary, non-performance-based factors. XGBoost is better documented and more established than CatBoost, hence facilitating the interpretation, tuning, and understanding of the applied model. For all these reasons, we suggest using XGBoost to predict the organic CTR of websites on Google result pages.

The general optimality of tree-based algorithms is in line with current research. The excellent performance of tree-based models indicates that they are able to capture feature interactions on tabular, heterogeneous data well, while also being able to learn from relatively few samples.

Overall, we observe four main patterns arising across the different models. First, models capturing interactions perform better than those that do not. Second, models designed for the use on tabular data make better predictions than models designed for user-item-specific CTR prediction. Third, the performance differences between models amplify with the presence of more features. Fourth, tree-based models clearly outperform all other model categories.

5.3.2 Explanatory power of feature subsets

By training the models on combinations of feature subsets, it is possible to evaluate the explanatory power of subsets and derive implications on the overall relevance of subsets for CTR.

When comparing the improvement in RMSE for additional feature subsets, we focus on the results of XGBoost. Thereby, the comparison is not diluted by outliers created through badly performing models. Additionally, we validate these comparisons by running paired t-tests on the results of the, on average, five best models⁵. Thus, we can determine if the difference in RMSE of feature subsets is statistically significant.

We find that RMSE scores improve when adding each of the other subsets to the *position* subset. Adding the *keyword* and *result* subsets notably improves the RMSE, reducing it by 22.6% and 21.1%, respectively. These subsets offer more explanatory power than the *SERP features* subset which only improves the RMSE by 9.8%. When adding *SERP features* to *position + keyword*, the previously best combination, we only see a slight improvement of 4.9%. The combination of all feature subsets results in the lowest RMSE. It reduces the error of the baseline *position* subset by 37.6% and also improves the error of the *position + keyword + SERP features* subset by 16.2%. All of the mentioned differences are statistically significant according to the paired t-test.

Together, this implies that all feature subsets are relevant, as they decrease error when being added and improve the baseline model that solely relies on position. Additionally, the larger improvements in RMSE for the *keyword* and *result* subset than the *SERP feature* subset, indicate that they have a larger effect on CTRs. However, this observation is only an

⁵ The T-test was conducted with RMSE scores from the following models: CatBoost, XGB, Random Forest, Gradient Boosting and TabNet. The alpha used is 5%.

indication as feature importances are not quantified from a model perspective. Individual part B solves this, by applying interpretation techniques to XGBoost, which analyze feature importances and interactions to understand the relevance of the subsets for CTR.

Increasing the number of combined feature subsets to three and four subsets improves the model performance with each added subset. However, improved results including the *result* subset must be interpreted with caution due to its highlighted lack of generalizability. To increase generalizability individual part C generates a variety of *result* features based on a result's title, that are generally applicable. This helps further understand the predictive power of results.

6. Conclusion

By using a dataset that is novel in its comprehensiveness, this work analyzed drivers of CTR and CTR prediction models. It extended existing SEO research by not only considering the position as a driver of CTR, but also aspects grouped into keyword, result, and SERP feature characteristics. Ultimately, the research questions can be answered as follows:

The analysis confirmed the prevailing SEO literature assumption that position is the most important influence on CTR. However, it was found that also keyword, result, and SERP feature characteristics have a significant impact on CTR. Nonetheless, these impacts are highly dependent on the position of a result and subject to feature interactions.

Furthermore, it was shown that state-of-the-art CTR prediction models are inapplicable for predicting average CTR on Google result pages. Instead, the model analysis revealed that tree-based models, specifically CatBoost and XGBoost, are best suited for the underlying prediction task forecasting CTR. The analysis of model results also revealed that keyword and result characteristics tend to have more predictive power than SERP features.

However, several limitations, which restrict the general applicability of the results, need to be acknowledged. The dataset only covered e-commerce stores in the US, thus being limited to a rather niche segment of the entire Google search spectrum. In addition to that, the analyzed dataset contained only 59k rows. This did not only lead to small sample sizes for some feature combinations so that no statistically significant conclusion could be drawn but also restricted the predictive power of neural networks, which can require large amounts of data to work best. Lastly, limited computation resources did not allow for extensive hyperparameter tuning of neural networks. This potentially prevented those networks from performing better in the model comparison.

Future work could tackle these limitations by applying this methodology to similar data from sectors other than e-commerce. Additionally, building on the analysis above, we recommend future research to attempt confirming the findings on a larger dataset which could also provide a solid base for analyses with neural networks. Furthermore, an enhancement of the data with time series information would allow an analysis of consumer behavior and the dynamics of effects over time. Finally, a similar analysis over different nationalities could provide interesting insights into differences in user behavior and support globally operating website providers.

To extend this analysis the individual parts provide an in-depth study on SERP feature effects on a result level and how these can improve the number of clicks to a website. Additionally, model interpretation techniques are applied to assess the model's applicability, reveal patterns in the data, and quantify the subsets' importance. In addition to that, the analysis of results is extended to also incorporate titles of results.

C - The impact of differently phrased titles

C.1. Introduction

Every regular Google search result features a URL, title, and description. Web page administrators can usually directly control these properties, unlike other aspects of result pages, such as the search volume, SERP features, or rankings. As such, they can influence some primary ways a user determines whether a result is relevant to their search query (Google 2021). Thus, an ideal title and description can considerably impact the CTR of a search result. As aforementioned in 1., increasing this CTR can largely contribute to claiming larger parts of the multi trillion-dollar e-commerce industry.

However, academic research is limited to how titles and descriptions can positively influence the ranking of a web page. Thus, it only indirectly analyzes the impact of the title on CTR but disregards the direct influence of the title on CTR on a given position for a particular result page. Furthermore, the first part of this work examined only the influence of result page properties and certain result domains but not the title and description of the results.

To address this knowledge gap and help SEO practitioners better understand how to formulate ideal titles, this work examines the impact of different engagement, content, and format of titles on CTR. The analysis is focused on different effects per intent and position but does not aim at reverse engineering which titles lead to a higher ranking.

First, background on titles and descriptions in Google result pages alongside a broad literature overview is provided. After explaining the dataset and methodology, the main findings from research and SEO guides are introduced and translated into features for analysis. Then, findings from an analysis of the impact of features are presented. Lastly, the effect of the generated features on the models from chapter 5 is examined.

C.2. Background and related works

C.2.1 Displayed titles and descriptions

For each search result, Google displays a title of up to 55-60 characters (Taher 2022) and a description of up to 150-160 characters⁶ (Le 2022). These titles and descriptions for web pages can be set using meta tags in the page's HTML and thus are in the website owner's control. Google states to preferably use these meta tags for the displayed title and description on the result page. However, in some cases, Google chooses the displayed text itself, which can vary from the meta tag (Google 2022d).

As a consequence, when analyzing the title and especially the description of a result, one can not simply infer the actually displayed text from the meta tag of the HTML. Chapter C.3. shows that while for the title in most cases the HTML title is also the title displayed on the result page, for descriptions only in very few cases the content of the meta tag is also actually displayed. Based on this limitation, the following analysis disregards descriptions and only focuses on titles.

C.2.2 Related works

A large amount of academic literature on SEO exists. Based on several works pointing out the position as the most important feature, most research focuses on factors that influence the position of a web page (Matošević, Dobša, and Mladenović 2021; Gupta, Agrawal, and Gupta 2016; Ziakis et al. 2019; Luh, Yang, and Huang 2016). Consequently, findings on how to adapt titles are focused on how to rank higher. These adaptations mostly solely concern variations of featuring the targeted keyword in a title as well as the formatting (Krrabaj, Baxhaku, and Sadrijaj 2017; Ziakis et al. 2019). This work extends the analysis to forms of engagement and additional content features of titles and analyzes the direct influence of titles on CTR.

⁶ Character ranges result from Google in practice cutting off contents after a specific amount of pixels

Apart from academic works on the influence of titles, a plethora of web pages and blogs give advice on optimizing titles. This advice is rarely backed by data and if it is, the sample size is usually small and limited to a single domain (Moseley 2021; von der Osten 2022). In the following, the validity of various claims is tested on an extensive dataset ranging over hundreds of domains.

Lastly, works from related fields covering similar issues exist. Two insightful studies are those of Biyani, Tsioutsoulis, and Blackmer (2016) who discover properties of clickbait news titles, and Pernice, Whinton, and Nielsen (2020) who conduct an extensive eye-tracking study on people's reading habits on the web. Methodology, variables, and findings serve as references for generated features in this work.

During the explanation of generated features in chapter C.4., the most relevant findings and claims from each field will be mentioned and translated into features for analysis and predictive machine learning models.

C.3. Dataset & Data Acquisition Methodology

The used dataset consists of the dataset used in the first part of this work, supplemented by a similar dataset with more than 800k data points on Google search results in the UK, totaling around 902k rows. Since both datasets were missing information about the title of results, as the first part of this work, this information needed to be acquired. For 81.7% of all unique URLs, a valid title and description could be obtained (find a detailed description of the data acquisition process in appendix C.I.). Subsequently, the similarity of HTML titles and descriptions to the content displayed on the result page was validated (find a detailed description of the comparison process in appendix C.II.).

By removing domains where the text similarity of HTML and Google result titles was frequently highly unequal, the expected share of titles that accurately resemble the shown

Google result titles could be raised to 89.4%. Although this share allows for high confidence in the representativeness of the HTML titles, a small amount of noise remains in the data. This can slightly harm the meaningfulness of findings.

Contrary to the high similarity of titles are the results obtained from comparing Google result descriptions to HTML descriptions. In only 31.6% of all cases, both descriptions were sufficiently similar. Furthermore, no frequently differing domains whose removal could increase the expected similarity, such as for titles, could be identified. Thus, it cannot be assumed that the HTML description is also the one shown on the Google result page, while scraping sufficient Google result page samples is impractical. As a consequence, and as mentioned above, descriptions were not considered in the analysis.

Similar to prior chapters, records with less than 20 impressions were dropped. After reducing the dataset to only rows with sufficient informative value, 346,807 rows remained for analysis. The data cleaning process is illustrated in appendix C.III.

C.4. Feature Generation

To examine the effect of different titles on the CTR, features of three main categories were generated: Engagement, content, and format. The choice of generated features is based on frequently suggested title optimization techniques in the related literature and dataset characteristics. For a technical definition of all features mentioned in the following, please refer to the title features data dictionary in appendix C.IV.

C.4.1 Engagement

Many resources suggest using patterns that aim to increase the appeal to read, visual prominence, and interest invoked by a title.

Morris (2021) suggests the use of a call to action, i.e., verbs as the first word of a title. The feature *verb at start* detects whether a title starts with a verb. Furthermore, an attempt to make titles more appealing is the use of exclamation or question marks and superlatives (Biyani, Tsioutsoulouklis, and Blackmer 2016). The features *has question mark*, *has exclamation mark*, and *count superlatives* capture these aspects. Alliterations are suggested as an engaging element by Adames (2022) and considered in the feature *has alliteration*. Although Google (2022f) considers capitalized titles “common in spam and untrustworthy ads”, some blogs still recommend the use of selected capitalized words for emphasis (Hardwick 2020). *Count capitalized* tracks the number of capitalized words in a title.

C.4.2 Content

Content features refer to the chosen content to display.

Mifsud (2019) recommends the use of acronyms, i.e., abbreviations formed from the initial letters of other words (‘United States’ vs. ‘US’), for product titles to save space and be more specific. *Has acronym* represents the existence of an acronym.

In an eye-tracking study, Pernice, Whitenton, and Nielsen (2020) show that readers on the web pay more attention to numeric characters as compared to their written counterparts (‘1’ vs. ‘one’). The feature *has number* checks whether there is at least one number in the title. *Has unit specification* and *has serial number* drill the existence of numbers further down to understand what exactly is quantified.

To capture the general technicality of the wording and syntax of a title, the formality measure suggested by Heylighen and Dewaele (1999) was adapted. The *technicality score* captures the share of nouns, adjectives, and numbers as opposed to more wordy syntactical elements.

One of the most common suggestions in both SEO guides and academic literature is to use the targeted keywords in the title (Constant Contact 2022; Lyons 2022; Ziakis et al. 2019). The

feature *KW / title overlap* measures the share of words in the keyword that overlap with the title.

Another commonly recommended practice is to feature the brand name in the title (Hardwick 2020; Morningscore 2022). To validate a positive effect of this measure, the feature *brand in title* is computed.

To enable further analysis of titles based on their logical meaning, semantic word clusters were built. The aim of such is to group logically related words such as ‘pharmacy’, ‘paracetamol’, and ‘ibuprofen’ together. In order to do so, first, all unique singularized words were extracted from the titles. For a clearer analysis, only English words were considered for the clustering. For that, words for which English is not among the two most likely predicted languages based on predictions of the fastText language identifier (Joulin et al. 2017; Joulin et al. 2016), were removed. The remaining words were vectorized using 300-dimensional word vectors, based on fastText (Bojanowski et al. 2017) trained on Common Crawl (600B tokens). To reduce complexity during clustering, the dimensionality of the resulting word vectors was reduced to 128 dimensions using principal component analysis (73% variance is retained). With the resulting dataset, KMeans clustering (MacQueen 1967; Arthur and Vassilvitskii 2006) was conducted. Using the elbow method (Marutho et al. 2018), $k = 32$ was identified as the best number of clusters (see elbow visualization in appendix C.VI.). As a final step, for each title, the *count of words belonging to each of the 32 clusters* was generated. A process overview of the clustering can be found in appendix C.V. and a 2-dimensional visualization of the resulting clusters in appendix C.VII.

C.4.3 Format

Title format features concern any components of titles that, with a given title content, change the appearance of it.

Many SEO guides recommend avoiding title truncation (Morningscore 2022; Lyons 2022).

The feature *truncated* captures all titles with more than 60 characters.

To separate categories and hierarchies, dashes (‘ - ’) and pipes (‘ | ’) are commonly used. A sample resulting structure is ‘<product name> | <category> | <brand>’. To understand whether the use of pipes and dashes is beneficial and whether there is potential overuse of them, the features *count pipes* and *count dashes* are computed. Furthermore, this allows for analyzing which of the two separators works better, as many SEO guides disagree between the recommended use of pipes vs. dashes (Moseley 2021; Montti 2020).

C.5. Analysis

In this chapter, the main findings with practical implications of patterns for each category of features are presented.

The analysis in chapter 4 highlighted the importance of positions and the need to differentiate between positions when examining patterns. To incorporate this, the following analysis averages the impact of a feature on a per-position level. As such, the impact of feature f is defined as

$$(C.1) \text{Impact}(f) = \frac{1}{n} * \sum_{i=1}^n \left(\frac{x_{fi}}{y_{fi}} - 1 \right),$$

where n is the number of analyzed positions, x_{fi} is the average CTR on position i if feature f is present, and y_{fi} is the average CTR on position i if feature f is not present. Features with cardinality greater than two are binarized⁷. For higher robustness of results, only the first 10 positions and results for positions with at least 50 samples for the computation of both x_{fi} and y_{fi} were considered. Statistical significance is assessed using the Wilcoxon signed-rank test

⁷ For binarization, a threshold was used. Values equal to or below threshold were set to 0 and values above to 1. [Thresholds] Features that involve count: 0; Technicality score: median technicality score; KW/title overlap: 0.2.

(Wilcoxon 1945) across the differences in average CTR for all positions, using alpha 0.05. This non-parametric measure was chosen due to the skewed distribution of avg. CTR values per position and thus a violation of assumptions for alternative parametric significance tests.

Ruotsalo et al. (2018) point out considerable differences in user behavior depending on their search intent. Based on this, in addition to the entire dataset, the analysis examines differences between the four intents. Figure C.1 gives a general overview of the impact of each feature for different intents. In the following, the main findings are summarized, put into context, and a more in-depth analysis is conducted where needed.

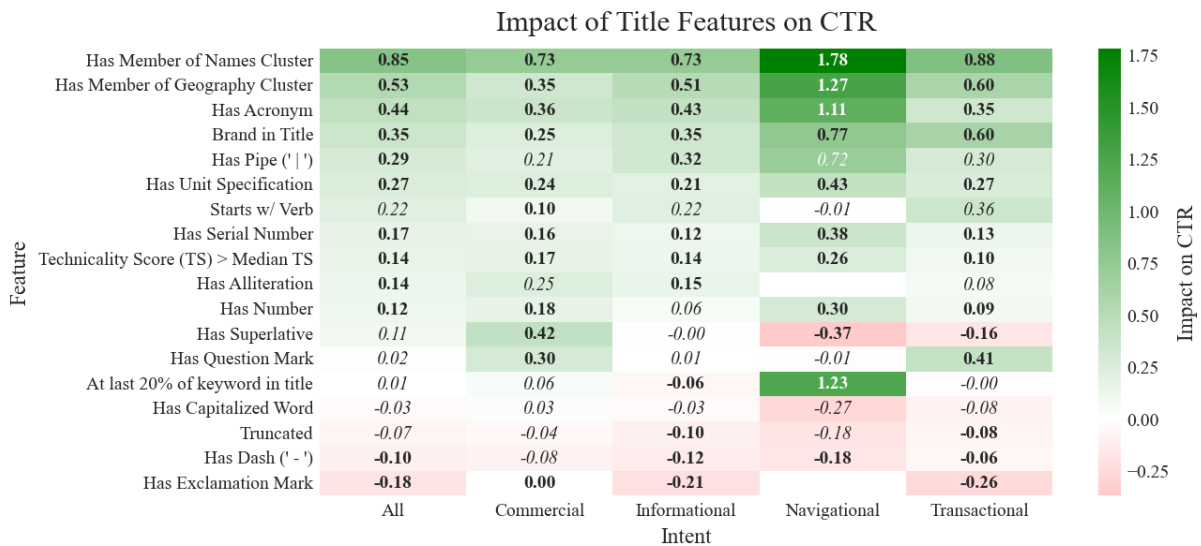


Figure C.1: Impact of the presence of title features on CTR. Values are empty, if for no position sufficient samples are available. For **bold** values, the impact is statistically significant, *italic* values are insignificant.

C.5.1 Engagement

A verb at the start of a title, superlatives, exclamation marks, and capitalized words can all be considered attempts to create more engaging titles. At the same time, users may also perceive their use as spammy and untrustworthy. Starting the title with a verb only has a statistically significant effect for commercial intents, where it slightly boosts the CTR. Superlatives work very well for searches with commercial intent while decreasing CTR in navigational and transactional searches. While the use of capitalized words has no statistically significant

impact on the CTR, the opposite holds true if exclamation marks are used. Especially for informational and transactional intent, exclamation marks decrease the CTR.

The aforementioned patterns are, however, rather extreme ways of trying to catch the attention of a user. More subtle ways of doing so can be found in the use of alliterations and question marks, which aim at creating a title that sounds good and spurs interest. Where a statistically significant impact can be observed, the use of alliterations increases CTR. The use of question marks has a statistically significant, strong positive effect for commercial and transactional intents.

As a consequence, one can conclude that the creation of more subtle engagement works much better than that through more extreme forms, such as caps or exclamation marks, which tend to do more harm than good. Furthermore, strong differences between intents are observable: While all attempts of engagement on searches with commercial intent boost CTR or have no effect, the exact opposite holds true for searches with navigational intent. This especially applies to navigational intent and is likely to be explained by the very determined nature of navigational searches, in which engaging patterns might only seem untrustworthy.

C.5.2 Content

Next to forms of engagement, the content of a title is highly relevant. The analysis found the presence of the brand name in the title to be statistically significant and extremely beneficial for the CTR, with an average increase of +35%, peaking for navigational intents at +77%. This suggests that being precise about product and brand can help to increase the CTR. The general, statistically significant increase of CTR in the presence of numbers further supports this supposition.

However, some light needs to be shed on what kinds of numbers boost CTR. If numbers occur as serial numbers, the effect on CTR is positive across all search intents and even higher than for numbers in general (+17% vs. +12%). With an average increase of +27% in CTR, the

appearance of numbers as unit specifiers has an even stronger positive impact. For both features, the positive impact is statistically significant across all intents.

The above implied positive reaction to more specific titles is further underlined by the enormous positive effect of acronyms in the title (+44%). Not only very specific content is well received, but also precise and technical language. Titles with an above-median syntactical technicality score, on average, perform better than titles with a score below the median. The impact of the presence of both acronyms and high technicality scores is statistically significant across all intents.

Together, these findings show that very precise and specific titles that display detailed information about both product and brand are beneficial. Interestingly, this effect is further magnified for navigational intent, which can be explained by the titles' ability to show the user that the result is exactly what they are looking for (while the user already knows what they expect to find).

The built clusters can offer additional insights into what logical content works well. Two clusters stand out, with both having a strong and statistically significant impact across all intents. The presence of words of a cluster consisting of first and last names leads to an average increase in CTR of +85%. In the underlying, strongly e-commerce focused dataset, first and last names are usually brand names. Thus, this finding can also be interpreted as further underlining the increase in CTR caused by the existence of the brand name in the title. Another cluster consists of city and country names. If a member of this cluster appears in a title, the CTR is on average 54% higher. This implies that specifying the geolocation of a shop helps achieve higher CTR and supports the previous finding that specific and precise titles are beneficial.

The remaining clusters deliver limited insights. They are either too broadly defined (e.g., a

range of verbs) or, when well defined (e.g., clusters of food, drinks, or drugs), have no significant impact on the CTR.

Although in this analysis, forms of engagement and content have been found to impact CTR strongly, most literature still suggests the presence of the targeted keyword as the most important influence. Interestingly, this does not hold true for commercial, transactional, and especially informational intent, according to this work's analysis. Covering more than 20% of the words in the keyword in the title has a very low or statistically insignificant impact on CTR. Also, whether only a small or a large fraction of the keyword overlaps with the title, does not have a considerable impact on CTR (see appendix C.VIII). This finding is very surprising as it is against most existing literature. A potential explanation for it is the age of these works, combined with both the continuous improvement of the Google search algorithm and changing user behavior. Nowadays, titles with synonyms or direct answers to what a user has searched for appear to work just as well as titles covering the entire searched keyword. An exception to this, again, form searches with navigational intent. For these searches, a higher resemblance of searched words in the result's title statistically significantly increases CTR. Searches where the title covers at least one-fifth of the searched words, achieve on average 123% higher CTRs. Related to previous explanations about navigational searches, a possible reason for this is that with navigational intent, the user knows exactly what they are looking for and thus becomes highly uninterested in any result not covering this.

C.5.3 Format

Lastly, a website's administrator can decide the format of titles. For the separation of contents and hierarchies in titles, pipes (' | ') work significantly better than dashes (' - '). On average, the mean CTR per position for results with pipes is 37% higher than for results with dashes. However, the use of such separators should be well considered: While the use of pipes

generally harms the CTR, the use of pipes only benefits the CTR if exactly one pipe is used. While the use of two pipes on average has no statistically significant impact on the CTR, three or more pipes harm the CTR significantly, as compared to the use of no separator (-26% average/position). However, the findings on separators are at risk of being influenced by domain effects, as one domain usually uses one separation structure. As a consequence, effects of the usage of a certain separator might only resemble the strength of all domains/brands that use it, but not the separator itself.

Keeping the title length below 60 characters to avoid truncation only has a minor impact. The only statistically significant impacts of -10% for informational searches and -8% for transactional searches are rather small compared to the impacts of other aforementioned patterns.

C.5.4 Intents

Next to the specific patterns regarding engagement, content, and format of a title, an important general pattern across different intents becomes apparent: While commercial, informational, and transactional intents behave mostly the same, searches with navigational intents inherit large deviations and special patterns. Thus, for keywords connected to navigational search intents, other rules and magnitudes of changes may apply, making it necessary to treat titles for results of such searches carefully.

C.6. Impact on model performance

To understand whether the identified patterns can also contribute to increased model performance, the modeling methodology from chapter 5 is extended by the newly generated features. As before, XGBoost is used for model performance comparison. For this comparison, all feature subset combinations from chapter 5 that do not contain the *result* feature subset yet were used again and extended by the newly generated features. This

approach compares how much the error improves when the new title features are added. Furthermore, it allows for a comparison between the predictive power of the title features and the *keywords* and *SERP features* subsets. Similarly to one-hot-encoded domains in chapter 5, word clusters are at risk of overfitting the dataset. To isolate the effect of word clusters, predictions were also compared with and without word clusters. Table C.1 summarizes the resulting scores. Appendix C.IX. contains a detailed list of all used title features.

		RMSE of subset combinations ⁸			
		Position	Position + Keyword	Position + SERP features	Position + Keyword + SERP Features
No added features		0.082	0.073	0.075	0.068
Title features added	w/o word clusters	0.074 <i>(-0.008)</i>	0.068 <i>(-0.005)</i>	0.068 <i>(-0.007)</i>	0.065 <i>(-0.003)</i>
	w/ word clusters	0.067 <i>(-0.015)</i>	0.062 <i>(-0.011)</i>	0.063 <i>(-0.012)</i>	0.060 <i>(-0.008)</i>

Table C.1: CTR prediction error for different feature subsets, with and without title features. *Italic* values indicate change from ‘No added features’.

The new title features have strong predictive power when added to the position subset. Even without the information of word clusters, the improvement in CTR is similar to that when adding the keyword subset or the SERP features subset. When additionally adding the word cluster information, the error even decreases below that of using information from both the keyword subset and the SERP features subset. Adding the new title features to the position subset and or both of the other subsets greatly improves the error as well (up to -18%).

Overall, predicting CTR can largely benefit from the title features generated in this part of the work. Depending on whether information about word clusters is used, the predictive power is on par with, or even better than that of the keyword and SERP features subset. As a consequence, the generated title features can positively contribute to an accurate CTR prediction and its accompanying economic benefits as outlined in 1.

⁸ Scores vary from those in 5.3 since the model was also trained on the UK dataset in addition to the US.

C.7. Conclusion and limitations

This part of the work examined patterns in Google result page titles that influence the CTR of a given result. Based on a large dataset, this work has extended existing academic literature to examine the direct impact of titles on CTR as well as tested the validity of prevailing, but commonly unproven suggestions from SEO guides. It was shown that statistically significant patterns range from engagement over content to formatting of a title. Specific and technical titles are consistently related to higher CTR, while wordy or overly exaggerating titles harm CTR. However, it also became apparent that the effect of different title components strongly depends on the search intent of a user. Often, while the existence of one pattern has a positive influence on CTR for one search intent, it can harm it for another. This shows that there doesn't exist a one-size-fits-all solution to writing titles. Instead, when writing the title for a webpage, practitioners should always keep the intent of the users' search query and the composition of different patterns in the title in consideration. Furthermore, the analysis demonstrated that by incorporating title features into the prediction of CTR, the prediction error can be decreased by up to 18%.

Since the underlying dataset mostly contains data from e-commerce shops, the findings and above-mentioned implications are, however, limited to web pages and titles of this kind. Furthermore, there was remaining noise in the dataset used for this analysis, since the titles Google shows on the result page are not always the ones specified in the HTML title tag.

Future work could expand the analysis to a broader dataset, not only containing data of e-commerce websites. To achieve higher confidence, a similar analysis could be conducted using the exact titles shown by Google. Lastly, the analysis of results could be expanded from titles to descriptions, which are much richer in information and thus have immense potential for the identification of additional relevant patterns.

References

- Adames, Ingrid. 2022. "How To Write Great SEO Titles." Search Engine Journal. Accessed November 23, 2022.
<https://www.searchenginejournal.com/title-capitalization-in-the-english-language/4882/>.
- Agarwal, Abhineet, Yan S. Tan, Omer Ronen, Chandan Singh, and Bin Yu. 2022. "Hierarchical Shrinkage: Improving the Accuracy and Interpretability of Tree-Based Methods." *Proceedings of the 39th International Conference on Machine Learning* 162:111-135.
- Alwosheel, Ahmed, Sander van Cranenburgh, and Caspar G. Chorus. 2018. "Is your dataset big enough? Sample size requirements when using artificial neural networks for discrete choice analysis." *Journal of Choice Modelling* 28:167-182.
- Arik, Sercan O., and Tomas Pfister. 2021. "TabNet: Attentive Interpretable Tabular Learning." *Proceedings of the AAAI Conference on Artificial Intelligence* 38, no. 8 (May): 6679-6687.
- Arthur, David, and Sergei Vassilvitskii. 2006. "k-means++: The Advantages of Careful Seeding." *Stanford Technical Report*, (June).
- Bala, Madhu, and Deepak Verma. 2018. "A Critical Review of Digital Marketing." *International Journal of Management, IT & Engineering* 8, no. 10 (10).
- Biau, Gérard, and Erwan Scornet. 2016. "A random forest guided tour." *TEST* 25 (April): 197-227.
- Biyani, Prakhar, Kostas Tsioutsoulis, and John Blackmer. 2016. "'8 Amazing Secrets for Getting More Clicks': Detecting Clickbaits in News Streams Using Article

- Informality.” *Proceedings of the Thirtieth AAAI Conference on Artificial Intelligence* 30, no. 1 (February): 94-100.
- Bojanowski, Piotr, Edouard Grave, Armand Joulin, and Tomas Mikolov. 2017. “Enriching Word Vectors with Subword Information.” *Transactions of the Association for Computational Linguistics* 5:135-146.
- Borisov, Vadim, Tobias Leeman, Kathrin Seßler, Johannes Haug, Martin Pawelczyk, and Gjergji Kasneci. 2022. “Deep Neural Networks and Tabular Data: A Survey.” *arXiv:2110.01889v3*, (June).
- Breiman, Leo. 1996. “Bagging predictors.” *Machine Learning* 24 (August): 123-140.
- Breiman, Leo. 2001. “Random forests.” *Machine Learning* 45, no. 1 (October): 5-32.
- Chen, Junxuan, Baigui Sun, Hao Li, Hongtao Lu, and Xian-Sheng Hua. 2016. “Deep CTR Prediction in Display Advertising.” *In Proceedings of the 24th ACM international conference on Multimedia* 16:811–820. <https://doi.org/10.1145/2964284.2964325>.
- Chen, Tianqi, and Carlos Guestrin. 2016. “XGBoost: A Scalable Tree Boosting System.” *KDD '16: Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, (August), 785-794.
- Cheng, Heng-Tze, Levent Koc, Jeremiah Hermsen, Tal Shaked, Tushar Chandra, Hrishi Aradhye, Glen Anderson, et al. 2016. “Wide & Deep Learning for Recommender Systems.” *DLRS 2016: Proceedings of the 1st Workshop on Deep Learning for Recommender Systems*, (September), 7-10.
- Chevalier, Stephanie. 2022a. “Global retail e-commerce sales 2026.” Statista. Accessed October 26, 2022.
<https://www.statista.com/statistics/379046/worldwide-retail-e-commerce-sales/>.
- Chevalier, Stephanie. 2022b. “Global online e-commerce traffic by source and medium.” Statista. Accessed October 28, 2022.

<https://www.statista.com/statistics/820293/online-traffic-source-and-medium-e-commerce-sessions/>.

Chiche, Alebachew, and Betselot Yitagesu. 2022. "Part of speech tagging: a systematic review of deep learning and machine learning approaches." *Journal of Big Data* 9, no. 10 (January).

Constant Contact. 2022. "What is an SEO Title Tag?" Constant Contact. Accessed November 19, 2022. <https://www.constantcontact.com/blog/website-seo-title-tag/>.

Cui, Meng, and Songyun Hu. 2011. "Search Engine Optimization Research for Website Promotion." *2011 International Conference of Information Technology, Computer Engineering and Management Sciences*, (09). 10.1109/ICM.2011.308.

Das, Subhankar. 2021. *Search Engine Optimization and Marketing: A Recipe for Success in Digital Marketing*. N.p.: CRC Press/Taylor and Francis Group.

De Gryze, Steven, Ivan Langhans, and Martina Vandebroek. 2007. "Using the correct intervals for prediction: A tutorial on tolerance intervals for ordinary least-squares regression." *Chemometrics and Intelligent Laboratory Systems* 87:147-154.

Deng, Wei, Junwei Pan, Tian Zhou, Deguang Kong, Aaron Flores, and Guang Lin. 2021. "DeepLight: Deep Lightweight Feature Interactions for Accelerating CTR Predictions in Ad Serving." *WSDM '21: Proceedings of the 14th ACM International Conference on Web Search and Data Mining*, (March), 922-930.
<https://doi.org/10.1145/3437963.3441727>.

Filzmoser, Peter, and Klaus Nordhausen. 2020. "Robust linear regression for high-dimensional data: An overview." *WIREs Computational Statistics* 13, no. 4 (July).

Flesch, Rudolph. 1948. "A new readability yardstick." *Journal of Applied Psychology* 32 (3): 221 - 223.

- Freund, Yoav, and Robert E. Schapire. 1996. "Experiments with a New Boosting Algorithm." *Machine Learning: Proceedings of the Thirteenth International Conference*.
- Friedman, Jerome H. 2001. "Greedy function approximation: A gradient boosting machine." *Annals of Statistics* 29, no. 5 (October): 1189-1232.
- Friedman, Jerome H., and Bodgan E. Popescu. 2008. "Predictive Learning via Rule Ensembles." In *Annals of Applied Statistics*, 916-954. 2nd ed. Vol. 3. <http://www.jstor.org/stable/30245114>.
- Gharibshah, Zhabiz, Xingquan Zhu, Arthur Hainline, and Michael Conway. 2020. "Deep Learning for User Interest and Response Prediction in Online Display Advertising." *Data Science and Engineering* 5:12-26.
- Google. 2021. "An update to how we generate web page titles." Google Search Central. Accessed November 6, 2022. <https://developers.google.com/search/blog/2021/08/update-to-generating-page-titles>.
- Google. 2022a. "Click through rate definition." Clickrate (CTR) - Google Ads-Support. Accessed September 20, 2022. <https://support.google.com/google-ads/answer/2615875?hl=de>.
- Google. 2022b. "Enable Search result features for your site | Documentation." Google Developers. Accessed November 14, 2022. <https://developers.google.com/search/docs/appearance/search-result-features>.
- Google. 2022c. "Data preprocessing for ML: options and recommendations | Cloud Architecture Center." Google Cloud. Accessed November 13, 2022. <https://cloud.google.com/architecture/data-preprocessing-for-ml-with-tf-transform-pt1?hl=de>.
- Google. 2022d. "Influencing Title Links in Google Search." Google Search Central. Accessed November 27, 2022. <https://developers.google.com/search/docs/appearance/title-link>.

- Gorishniy, Yury, Ivan Rubachev, Valentin Khrulkov, and Artem Babenko. 2021. "Revisiting Deep Learning Models for Tabular Data." *Advances in Neural Information Processing Systems 34 (NeurIPS 2021)*, (May).
- Graepel, T., J. Q. Candela, T. Borchert, and R. Zerbrich. 2010. "Web-scale bayesian click-through rate prediction for sponsored search advertising in microsoft's bing search engine." *Proceedings of the 27th International Conference on Machine Learning (ICML-10)*, 13-20.
- Grinsztajn, Léo, Edouard Oyallon, and Gaël Varoquaux. 2022. "Why do tree-based models still outperform deep learning on tabular data?" *arXiv:2207.08815*, (July).
- Grips. 2022. "Grips - About Us." Grips - Transaction Intelligence for eCommerce. Accessed November 24, 2022. <https://gripsintelligence.com/about>.
- Guo, Huifeng, Ruiming Tang, Yunming Ye, Zhenguo Li, and Xiuqiang He. 2017. "DeepFM: A Factorization-Machine based Neural Network for CTR Prediction." *Proceedings of the Twenty-Sixth International Joint Conference on Artificial Intelligence*, 1725–1731.
- Gupta, Sudobh, Neha Agrawal, and Sandeep Gupta. 2016. "A Review on Search Engine Optimization: Basics." *International Journal of Hybrid Information Technology* 9 (5): 381-390.
- Hancock, J.T., and T.M. Khoshgoftaar. 2020. "CatBoost for big data: an interdisciplinary review." *Journal of Big Data* 7 (94): 1-45.
<https://doi.org/10.1186/s40537-020-00369-8>.
- Hardwick, Joshua. 2020. "How to Craft the Perfect SEO Title Tag (Our 4-Step Process)." Ahrefs. Accessed November 24, 2022. <https://ahrefs.com/blog/title-tag-seo/>.
- Heylighen, Francis, and Jean-Marc Dewaele. 1999. "Formality of Language: definition, measurement and behavioral determinants." *Internal Report, Center "Leo Apostel", Free University of Brussels*.

- Iqbal, Muhammad, Muhammad Noman Khalid, Amir Manzoor, Malik Muneeb Abid, and Nazir Ahmed Shaikh. 2022. "Search Engine Optimization (SEO): A Study of important key factors in achieving a better Search Engine Result Page (SERP) Position." In *Sukkur IBA Journal of Emerging Technologies*, 1-15. Vol 6 ed. Vol. 1. <http://journal.iba-suk.edu.pk:8089/sibajournals/index.php/sjcms/article/view/924/309>.
- Jais, Imran K., Amelia R. Ismail, and Syed Q. Nisa. 2019. "Adam Optimization Algorithm for Wide and Deep Neural Network." *Knowledge Engineering and Data Science* 2 (1).
- Joulin, Armand, Edouard Grave, Piotr Bojanowski, Mathijs Douze, H erve J egou, and Tomas Mikolov. 2016. "FastText.zip: Compressing text classification models." *arXiv:1612.03651*, (December).
- Joulin, Armand, Edouard Grave, Piotr Bojanowski, and Tomas Mikolov. 2017. "Bag of Tricks for Efficient Text Classification." *Proceedings of the 15th Conference of the European Chapter of the Association for Computational Linguistics* 2:427-431.
- Kadra, Arind, Marius Lindauer, Frank Hutter, and Josif Grabocka. 2021. "Well-tuned Simple Nets Excel on Tabular Datasets (NeurIPS 2021)." *35th Conference on Neural Information Processing Systems*.
- Kingsnorth, Simon. 2022. *Digital Marketing Strategy: An Integrated Approach to Online Marketing*. N.p.: Kogan Page.
- Kossen, Jannik, Neil Band, Clare Lyle, Aidan N. Gomez, Tom Rainforth, and Yarin Gal. 2021. "SelfAttention Between Datapoints: Going Beyond Individual Input-Output Pairs in Deep Learning." *arXiv:2106.02584*, (June).
- Krrabaj, Samedin, Fesal Baxhaku, and Dukagjin Sadrijaj. 2017. "Investigating search engine optimization techniques for effective ranking: A case study of an educational site." *2017 6th Mediterranean Conference on Embedded Computing (MECO)*, (June).

- Le, Claire. 2022. "Google's Title and Meta Descriptions Length (Updated 2022)." SEOPressor. Accessed November 25, 2022. <https://seopressor.com/blog/google-title-meta-descriptions-length/>.
- Ledford, Jerri L. 2009. *Search Engine Optimization Bible*. N.p.: Wiley.
- Lewandowski, Dirk, Sebastian Sünkler, and Nurce Yagci. 2020. "The influence of search engine optimization on Google's results: A multi-dimensional approach for detecting SEO." *13th ACM Web Science Conference 2021* 21 (06): 12-20.
- Luh, Cheng-Jye, Sheng-An Yang, and Ting-Li D. Huang. 2016. "Estimating Google's search engine ranking function from a search engine optimization perspective." *Online Information Review*, (April).
- Lyons, Kelly. 2022. "What Is a Title Tag & How to Optimize Title Tags for Google." Semrush. Accessed November 27, 2022. <https://www.semrush.com/blog/title-tag/>.
- Ma, Chao, Yuze Liao, Yuan Wang, and Zhen Xiao. 2016. "F2M: Scalable Field-Aware Factorization Machines." *30th Conference on Neural Information Processing Systems*.
- MacQueen, J. 1967. "Some methods for classification and analysis of multivariate observations." *Berkeley Symposium on Mathematical Statistics and Probability* 5 (1): 281-297.
- Marutho, Dhendra, Sunarna H. Handaka, Ekaprana Wijaya, and Muljono. 2018. "The Determination of Cluster Number at k-Mean Using Elbow Method and Purity Evaluation on Headline News." *2018 International Seminar on Application for Technology of Information and Communication*, (September).
- Matošević, Goran, Jasminka Dobša, and Dunja Mladenović. 2021. "Using Machine Learning for Web Page Classification in Search Engine Optimization." *Future Internet* 13, no. 1 (January).

- Messenger, Robert, and Lewis Mandell. 1971. "A Modal Search Technique for Predictive Nominal Scale Multivariate Analysis." *Journal of the American Statistical Association* 67, no. 340 (November): 768-772.
- Mifsud, Justin. 2019. "15 Title Tag Optimization Guidelines For Usability and SEO." Usability Geek. Accessed November 25, 2022.
<https://usabilitygeek.com/15-title-tag-optimization-guidelines-for-usability-and-seo/>.
- Montti, Roger. 2020. "Three Ways a Pipe or Dash in Title Tag Makes a Difference." Search Engine Journal. Accessed November 22, 2022.
<https://www.searchenginejournal.com/pipe-or-dash-in-title-tag/378099/>.
- Morgan, James N., and John A. Sonquist. 1963. "Problems in the Analysis of Survey Data, and a Proposal." *Journal of the American Statistical Association* 58, no. 302 (June): 415-434.
- Morningscore. 2022. "What are Title Tags, and do they help with SEO? [+13 writing tips]." Morningscore. Accessed November 26, 2022.
<https://morningscore.io/what-are-title-tags-in-seo/>.
- Morris, Corey. 2021. "Title Tag Optimization: A Complete How-to Guide." Search Engine Journal. Accessed November 25, 2022.
<https://www.searchenginejournal.com/on-page-seo/title-tag-optimization/#close>.
- Moseley, Brian. 2021. "Case Study: Should You Add Pipes or Dashes to Your Title Tag?" Semrush. Accessed November 22, 2022.
<https://www.semrush.com/blog/case-study-should-you-add-pipes-or-dashes-to-your-title-tag/>.
- Moz. 2022. "What Is A SERP Feature? Common Types And How To Win Them." Moz. Accessed November 12, 2022. <https://moz.com/learn/seo/serp-features>.

- Olson, Eric M., Kai M. Olson, Andrew J. Czaplewski, and Thomas Martin Key. 2021. "Business strategy and the management of digital marketing." *Business Horizons* 64, no. 2 (04): 285-293.
- Omer, Sagi, and Lior Rokach. 2018. "Ensemble learning: A survey." *WIREs Data Mining Knowl Discov* 8 (4).
- O'Neill, Aaron. 2022. "Japan GDP 1987-2027." Statista. Accessed November 27, 2022. <https://www.statista.com/statistics/263578/gross-domestic-product-gdp-of-japan/>.
- Pan, Junwei, Jian Xu, Alfonso Lobos Ruiz, Wenliang Zhao, Shengjun Pan, Yu Sun, and Quan Lu. 2018. "Field-weighted Factorization Machines for Click-Through Rate Prediction in Display Advertising." *In Proceedings of the 2018 World Wide Web Conference* 18.
- Pernice, Kara, Kathryn Whitenton, and Jakob Nielsen. 2020. *How People Read on the Web: The Eyetracking Evidence*. 2nd ed. N.p.: Nielsen Norman Group.
- Porter, Martin F. 1980. "An algorithm for suffix stripping." *Program: electronic library and information systems* 14, no. 3 (March): 130-137.
- Prokhorenkova, Liudmila, Gleb Gusev, Aleksandr Vorobev, Anna V. Dorogush, and Andrey Gulin. 2018. "CatBoost: unbiased boosting with categorical features." *Advances in Neural Information Processing Systems* 31.
- Quinlan, John R. 1986. "Induction of Decision Trees." *Machine Learning* 1 (March): 81-106.
- Ribeiro, Marco Tulio, Sameer Singh, and Carlos Guestrin. 2016. "Model-agnostic interpretability of machine learning." *ICML Workshop on Human Interpretability in Machine Learning*. <https://arxiv.org/abs/1606.05386>.
- Richardson, M., E. Dominowska, and R. Ragno. 2007. "Predicting clicks: estimating the click-through rate for new ads." *Proceedings of the 16th international conference on World Wide Web*, 521-530.

- Ruotsalo, Tuuka, Jaako Peltonen, Manuel J. Eugster, Dorota Głowacka, Patrik Floréen, Petri Myllymäki, Giulio Jacucci, and Samuel Kaski. 2018. "Interactive Intent Modeling for Exploratory Search." *ACM Transactions on Information Systems* 36, no. 4 (October): 1-46.
- Sam-Martin, Kristina. 2020. *Google's featured snippets in the Context of Strategic Content Marketing*. N.p.: E-Book - Kristina SamMartin.
- Semrush. 2021. "Basic docs." Semrush Developers. Accessed September 22, 2022. <https://developer.semrush.com/api/v3/analytics/basic-docs/#serp-features>.
- Semrush. 2022a. "A Beginner's Guide to Keyword Search Volume." Semrush. Accessed November 3, 2022. <https://www.semrush.com/blog/keyword-search-volume/>.
- Semrush. 2022b. "What Are SERP Features? An In-Depth Guide." SEMrush. Accessed December 1, 2022. <https://www.semrush.com/blog/serp-features-guide/#how-to-tell-if-your-competitors%E2%80%99-sites-have-serp-features>.
- Shih, Bih-Yaw, Chen-Yuan Chen, and Zih-Siang Chen. 2013. "Retracted: An Empirical Study of an Internet Marketing Strategy for Search Engine Optimization." *Human Factors and Ergonomics in Manufacturing and Service Industries* 23, no. 6 (11).
- Snoek, Jasper, Hugo Larochelle, and Ryan P. Adams. 2012. "Practical Bayesian Optimization of Machine Learning Algorithms." *Advances in Neural Information Processing Systems* 25 (NIPS 2012).
- Stevenson, Angus. 2010. *Oxford Dictionary of English*. N.p.: OUP Oxford, 2010.
- Su, Xiaogang, Xin Yan, and Chih-Ling Tsai. 2012. "Linear regression." *WIREs Computational Statistics* 4, no. 3 (April): 275-294.
- Taher, Sara. 2022. "What should the title tag length be in 2023?" Search Engine Land. Accessed November 26, 2022. <https://searchengineland.com/title-tag-length-388468>.

- Tenet Partners. 2020. "2020 Top 100 Most Powerful Brands," The essential brand rises - adapting to core consumer needs. Ranking the brands. Accessed November 28, 2022. <https://www.rankingthebrands.com/PDF/Top%20100%20Most%20Powerful%20Brands%202020,%20Tenet.pdf>.
- Tober, Marcus. 2022. "Zero-clicks Study." Semrush. Accessed December 6, 2022. <https://www.semrush.com/blog/zero-clicks-study/>.
- Turner, Ryan, David Eriksson, Michael McCourt, Juha Kiili, Eero Laaksonen, Zhen Xu, and Isabelle Guyon. 2021. "Bayesian Optimization is Superior to Random Search for Machine Learning Hyperparameter Tuning: Analysis of the Black-Box Optimization Challenge 2020." *arXiv:2104.10201*, (April).
- Uyanık, Gül den K., and Neşe Güler. 2013. "A study on mutiple linear regression analysis." *Procedia - Social and Behavioral Sciences* 106 (December): 234-240.
- von der Osten, Barbara. 2022. "SEO Case Studies: Learn From These 7 Success Stories." Rock Content. Accessed December 2, 2022. <https://rockcontent.com/blog/seo-case-studies/>.
- Wang, Fang, Warawut Suphamitmongkol, and Bo Wang. 2013. "Advertisement Click-Through Rate Prediction using Multiple Criteria Linear Programming Regression Model." *Procedia Computer Science* 17:803-811.
- Wang, Lih-Werm, Michael J. Miller, Michael R. Schmitt, and Frances K. When. 2013. "Assessing readability formula differences with written health information materials: Application, results, and recommendations." *Research in Social and Administrative Pharmacy* 9 (5): 503-516. <https://doi.org/10.1016/j.sapharm.2012.05.009>.
- Wang, Zhiqiang, Qingyun She, and Junlin Zhang. 2021. "MaskNet: Introducing Feature-Wise Multiplication to CTR Ranking Models by Instance-Guided Mask." *Proceedings of DLP-KDD 2021*.

- Wheelhouse. 2022. “How SERP Features Can Negatively Impact Your Business (Updated 2022).” Wheelhouse DMG. Accessed December 9, 2022.
<https://www.wheelhousedmg.com/blog/important-google-serp-features-and-how-to-optimize-for-them/>.
- Wilcoxon, Frank. 1945. “Individual Comparisons by Ranking Methods.” *Biometrics Bulletin* 1, no. 6 (December): 80-83.
- Yan, Ling, Wu-Jun Li, Gui-Rong Xue, and Dingyi Han. 2014. “Coupled Group Lasso for Web-Scale CTR Prediction in Display Advertising.” *Proceedings of the 31st International Conference on Machine Learning* 32 (2): 802-810.
- Yang, Yanwu, and Panyu Zhai. 2022. “Click-through rate prediction in online advertising: A literature review.” *Information Processing & Management* 59, no. 2 (03).
- Yuchin, Juan, Yong Zhuang, Wei-Sheng Chin, and Chih-Jen Lin. 2016. “Field-aware Factorization Machines for CTR Prediction.” *Proceedings of the 10th ACM Conference on Recommender Systems* 16 (09): 43-50.
- Yujian, Li, and Liu Bo. 2007. “A Normalized Levenshtein Distance Metric.” *IEEE Transactions on Pattern Analysis and Machine Intelligence* 29, no. 6 (June): 1091-1095.
- Zhang, Weinan, Jiarui Qin, Wei Guo, Ruiming Tang, and Xiuqiang He. 2021. “Deep learning for click-through rate estimation.” *arXiv preprint arXiv:2104.10584*.
- Zhou, Guorui, Xiaoqiang Zhu, Chenru Song, Ying Fan, Han Zhu, Xiao Ma, Yanghui Yan, Junqi Jin, Han Li, and Kun Gai. 2018. “Deep Interest Network for Click-Through Rate Prediction.” *Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining* 18 (07): 1059-1068.

Ziakis, Christos, Maro Vlachopoulou, Theodosios Kyrkoudis, and Makrina Karagkiozidou.

2019. "Important Factors for Improving Google Search Rank." *Future Internet* 11, no. 32 (01).

List of Figures

Figure 1: Mean CTR over keyword complexity score for branded and unbranded keywords. Mean CTR is calculated as a centered moving average with a total window size of 24. Keywords are more complex for low values and simpler for high values.

Figure 2: Mean CTR per position if a SERP feature is present (1) compared to when it is not present (0). Left plot shows pattern 'Lower', middle plot pattern 'Higher', and right plot pattern 'Lower → Higher'.

Figure 3: Loss in CTR per position. Values are avg. CTR for each position in% of avg. CTR on position 1. Only combinations with > 200 occurrences are considered (n=59). Dashed lines represent the top three most frequent SERP feature combinations. Shaded areas represent the range in which 50% and 90% of all values fall.

Figure 4: The Kernel Density Estimation (KDE) plot for the top six domains by occurrence, normalized to unified area under curve and x-axis limited to [0.0; 1.0] CTR-interval.

Figure 5: Friedman's H-statistic for two-way feature interactions. Highly correlated features and features with small number of occurrences were dropped, as they are likely to bias the H-statistic.

Figure 6: Mean CTR for branded and unbranded result pages over the count of page-level SERP features.

Figure C.1: Impact of the presence of title features on CTR. Values are empty, if for no position sufficient samples are available. For **bold** values, the impact is statistically significant, *italic* values are insignificant.

List of Tables

Table 1: Impact of the presence of a SERP feature on the average CTR, split between page and positional SERP features. ‘Lower $\rightarrow$ Higher’ refers to cases where the CTR within the first three positions is lowered and on subsequent positions higher. SERP features that also appear in Figure X, are underlined. SERP features without a clear pattern or $n < 1000$ are not listed.

Table 2: Benchmark results on different feature subsets. The top results for each feature subset are **bold**. We also underline the second-best result

Table C.1: CTR prediction error for different feature subsets, with and without title features. Italic values indicate change from ‘No added features’.

Appendix

I. SERP Features Illustration



Illustration of SERP features. For a detailed description and examples of each SERP feature we also recommend visiting the Semrush SERP feature guide: <https://www.semrush.com/blog/serp-features-guide>. Image source: Semrush.

II. Data Dictionary

The texts in the description fields with data source “Semrush” are taken from Semrush (2021).

Variable	Description	Data Source	Example
<i>Ads</i>	Any type of ad on the page. This can include ads on top of the search, at bottom, or shopping ads. Describes whether the SERP feature appears at least once on the page.	Self-engineered	1
<i>AdWords Bottom</i>	A series of ads (up to 4) that appear at the bottom of the first search results page. Describes whether the SERP feature appears at least once on the page.	Semrush	1
<i>AdWords Top</i>	A series of ads (up to 4) that appear at the top of the first search results page. Describes whether the SERP feature appears at least once on the page.	Semrush	1
<i>AMP</i>	AMP pages (i.e., results marked with the word "AMP" and a gray lightning bolt) shown in search results on mobile devices.	Semrush	1
<i>Avg. Monthly Volume</i>	Search volume averaged over months	Keywordplanner	1312.65
<i>Branded</i>	If the keyword includes a brand name.	Grips	1
<i>Carousel</i>	A row of horizontally scrollable images displayed at the top of search results. Describes whether the SERP feature appears at least once on the page.	Semrush	1
<i>Clicks</i>	Number of clicks on targeted search result	Grips	25
<i>Competition</i>	Measure of competition on search term with 0.0 being the lowest and 1.0 being the highest.	Keywordplanner	0.22
<i>Count Page SERP Feat.</i>	Count of positive (=1) page SERP features for search result page.	Self-engineered	2
<i>Count Positional SERP Feat.</i>	Count of positive (=1) positional SERP features for search result.	Self-engineered	5

Variable	Description	Data Source	Example
<i>CPC</i>	Cost per click of ad entry.	Keywordplanner	2.34
<i>CTR</i>	Click-through rate (clicks divided by impressions).	Grips	0.25
<i>Domain</i>	Domain of the URL of each search result.	Semrush	'jomashop.com'
<i>FAQ</i>	A list of questions related to a particular search that shows up for a particular organic search result. When clicked on, each of the questions reveals the answer. Describes whether the SERP feature appears at least once on the page.	Semrush	1
<i>FAQ (Positional)</i>	A list of questions related to a particular search that shows up for a particular organic search result. When clicked on, each of the questions reveals the answer. Describes whether the result is featured in the SERP feature.	Semrush	1
<i>Featured Images</i>	A collection of images is usually displayed at the top of the SERP if Google considers visual results to be more relevant than text results. Only for mobile devices. Describes whether the SERP feature appears at least once on the page.	Semrush	1
<i>Featured Snippet</i>	A short answer to a user's search query with a link to the third-party website it is taken from that appears at the top of all organic search results. Describes whether the SERP feature appears at least once on the page.	Semrush	1
<i>Featured Snippet (Positional)</i>	A short answer to a user's search query with a link to the third-party website it is taken from that appears at the top of all organic search results. Describes whether the result is featured in the SERP feature.	Semrush	1
<i>Featured Video</i>	A video result to a search query that is displayed at the top of all organic search results. Describes whether the SERP feature appears at least once on the page.	Semrush	1
<i>Flights</i>	A block that displays flights related to a search query. Flight results include information on flight dates, duration, the number of transfers and prices. Data is taken from Google Flights. Describes whether the SERP feature appears at least once on the page.	Semrush	1

Variable	Description	Data Source	Example
<i>Hotels Pack</i>	A block that displays hotels related to a search query. Hotel results include information on prices and rating, and allows users to check availability for certain dates. Describes whether the SERP feature appears at least once on the page.	Semrush	1
<i>Image</i>	An image result with a thumbnail displayed along with other organic search results. Describes whether the SERP feature appears at least once on the page.	Semrush	1
<i>Image (Positional)</i>	An image result with a thumbnail displayed along with other organic search results. Describes whether the result is featured in the SERP feature.	Semrush	1
<i>Image Pack</i>	A collection of images related to a search query that is usually displayed between organic search results. Describes whether the SERP feature appears at least once on the page.	Semrush	1
<i>Image Pack (Positional)</i>	A collection of images related to a search query that is usually displayed between organic search results. Describes whether the result is featured in the SERP feature.	Semrush	1
<i>Impressions</i>	Impressions of the result, including estimates for results on the second search page..	Grips	100
<i>Instant Answer</i>	A direct answer to a user's search query that is usually displayed at the top of organic search results in the form of a gray-bordered box. Describes whether the SERP feature appears at least once on the page.	Semrush	1
<i>Intent: Commercial</i>	Trying to learn more before making a purchase decision (e.g. "Subaru vs. Nissan")	Semrush	1
<i>Intent: Informational</i>	Trying to learn more about something (e.g., "What's a good car?")	Semrush	1
<i>Intent: Navigational</i>	Trying to find something (e.g., "Subaru website")	Semrush	1
<i>Intent: Transactional</i>	Trying to complete a specific action (e.g., "buy Subaru Forester")	Semrush	1

Variable	Description	Data Source	Example
<i>Jobs Search</i>	A number of job listings related to a search query that appear at the top of the search results page. Job listings include the job title, the company offering the job, a site where the listing was posted, etc. Describes whether SERP feature appears at least once on the page.	Semrush	1
<i>Keyword</i>	The searched keyword.	Semrush	'jomashop burberry scarf'
<i>Keyword Complexity Score</i>	Flesch reading ease score (Flesch 1948) of the keyword, with higher values indicating easier-to-read keywords.	Self-engineered	112
<i>Keyword Count</i>	Number of words the search query (keyword) consists of.	Grips	3
<i>Keyword Difficulty</i>	Difficulty to rank for given keyword expressed from 0 to 100	Semrush	87
<i>Keyword Included in Domain</i>	Keyword Included in Domain.	Self-engineered	0
<i>Keyword Included in URL</i>	Keyword Included in URL.	Self-engineered	1
<i>Keyword Length</i>	Number of characters in keyword	Grips	23
<i>Knowledge Panel</i>	Panel on right side of results that often includes images, facts, social media links, and other relevant information to the search query. Describes whether SERP feature appears at least once on the page.	Semrush	1
<i>Knowledge Panel (Positional)</i>	Panel on right side of results that often includes images, facts, social media links, and other relevant information to the search query. Describes whether the result is featured in the SERP feature.	Semrush	1
<i>Local Pack</i>	Embedded Google Maps frame on top of search results. Describes whether SERP feature appears at least once on the page.	Semrush	1

Variable	Description	Data Source	Example
<i>Local Pack (Positional)</i>	Embedded Google Maps frame on top of search results. Describes whether the result is featured in the SERP feature.	Semrush	1
<i>Number of Results</i>	Total number of results for search of keyword	Semrush	456789
<i>People Also Ask</i>	A series of questions that may relate to a search query that appears in an expandable grid box labeled "People also ask" between search results. Describes whether SERP feature appears at least once on the page.	Semrush	1
<i>People Also Ask (Positional)</i>	A series of questions that may relate to a search query that appears in an expandable grid box labeled "People also ask" between search results. Describes whether the result is featured in the SERP feature.	Semrush	1
<i>Position</i>	Position of the result among other search results	Semrush	3
<i>Position (monthly)</i>	Average position in the last month.	Grips	2.45
<i>Position Difference</i>	Position subtracted from previous position, i.e. positive values mean a decrease in position	Semrush	-1
<i>Previous Position</i>	Position last month	Semrush	2
<i>Reviews</i>	Organic search results marked with star ratings and including the number of reviews the star rating is based on. Describes whether SERP feature appears at least once on the page.	Semrush	1
<i>Reviews (Positional)</i>	Organic search results marked with star ratings and including the number of reviews the star rating is based on. Describes whether the result is featured in the SERP feature.	Semrush	1
<i>Search Volume</i>	The search volume of a keyword, cleaned from the original (inflated) value.	Grips	75643

Variable	Description	Data Source	Example
<i>SERP</i>			
<i>Features by Keyword</i>	SERP features listed by ID, before one-hot-encoded.	Semrush	1
<i>Shopping Ads</i>	A row of horizontally scrollable paid shopping results that appear at the top of a search results page for a brand or product search query, and include the website's name, pricing, and product image. Describes whether SERP feature appears at least once on the page.	Semrush	1
<i>Site Links</i>	A set of links to other pages of a website that is displayed under the main organic search result and for brand-related search queries. Describes whether SERP feature appears at least once on the page.	Semrush	1
<i>Site Links (Positional)</i>	A set of links to other pages of a website that is displayed under the main organic search result and for brand-related search queries. Describes whether the result is featured in the SERP feature.	Semrush	1
<i>Top Stories</i>	A card-style snippet presenting up to three news-related results relevant to user's search query, which is usually displayed between organic search results. Describes whether SERP feature appears at least once on the page.	Semrush	1
<i>Top Stories (Positional)</i>	A card-style snippet presenting up to three news-related results relevant to user's search query, which is usually displayed between organic search results. Describes whether the result is featured in the SERP feature.	Semrush	1
<i>Trends</i>	How much interest web searchers have shown in a given keyword in the last 12 months.	Semrush	1
<i>Tweet</i>	A card-style snippet displaying the most recent tweets related to a search query. Describes whether SERP feature appears at least once on the page.	Semrush	1
<i>URL</i>	The url of a google search result	Semrush	'https://www.jomas-hop.com/burberry-4031051.html'

Variable	Description	Data Source	Example
<i>Video</i>	Video results with a thumbnail displayed along with other organic search results. Describes whether SERP feature appears at least once on the page.	Semrush	1
<i>Video (Positional)</i>	Video results with a thumbnail displayed along with other organic search results. Describes whether the result is featured in the SERP feature.	Semrush	1
<i>Video Carousel</i>	A row of horizontally scrollable videos displayed among search results. Describes whether SERP feature appears at least once on the page.	Semrush	1
<i>Video Carousel (Positional)</i>	A row of horizontally scrollable videos displayed among search results. Describes whether the result is featured in the SERP feature.	Semrush	1

III. Feature Subsets

Subset	Features		
Position	Position	Position Difference	Position (monthly)
Keyword	Number of Results	Intent: Transactional	Keyword Length
	Keyword Difficulty	Competition	Keyword Count
	Intent: Commercial	CPC	Branded
	Intent: Informational	Search Volume	Complexity Score
	Intent: Navigational	Avg. Monthly Search Vol.	
SERP Feat.	Instant Answer	Featured Snippet	Site Links (Positional)
	Knowledge Panel	Image	Reviews (Positional)
	Carousel	Jobs Search	Video (Positional)
	Local Pack	Video Carousel	Featured Snippet (Positional)
	Top Stories	People Also Ask	Image (Positional)
	Image Pack	FAQ	Video Carousel (Positional)
	Site Links	Flights	People Also Ask (Positional)
	Reviews	Knowledge Panel (Positional)	FAQ (Positional)
	Tweet	Local Pack (Positional)	Ads
	Video	Top Stories (Positional)	Count Page SERP Feat.
	Featured Video	Image Pack (Positional)	Count Positional SERP Feat.
Result	Keyword in Domain	Keyword in URL	Domain (<i>One-Hot-Encoded</i>)

IV. Tested Hyperparameters

Model	Parameter	Values considered for hyperparameter tuning
Random Forest	Max depth	Categorical: [None, 2, 3, 4], p=[0.7, 0.1, 0.1, 0.1]
	Number of estimators	Integer Log Uniform: 10 $\rightarrow$ 3000
	Criterion	Categorical: ['squared_error', 'absolute_error']
	Max features	Categorical: ['sqrt', 'log2', None, 0.1, 0.2, 0.3, 0.4, 0.5, 0.6, 0.7, 0.8, 0.9]
	Min samples split	Categorical: [2, 3], p=[0.95, 0.05]
	Min samples leaf	Integer Log Uniform: 1 $\rightarrow$ 50
	Bootstrap	Categorical: [True, False]
	Min impurity decrease	Categorical: [0.0, 0.01, 0.02, 0.05], p=[0.85, 0.05, 0.05, 0.05]
GBDT	Loss	Categorical: ['squared_error', 'absolute_error', 'huber']
	Learning rate	Real Log Uniform: 0.01 $\rightarrow$ 10.0
	Subsample	Real Uniform: 0.5 $\rightarrow$ 1.0
	Number of estimators	Integer Log Uniform: 10 $\rightarrow$ 1000
	Criterion	Categorical: ['friedman_mse', 'squared_error']
	Max depth	Categorical: [None, 2, 3, 4, 5], p=[0.1, 0.1, 0.6, 0.1, 0.1]
	Min samples split	Categorical: [2, 3], p=[0.95, 0.05]
	Min samples leaf	Integer Log Uniform: 1 $\rightarrow$ 50
	Min impurity decrease	Categorical: [0.0, 0.01, 0.02, 0.05], p=[0.85, 0.05, 0.05, 0.05]
	Max leaf nodes	Categorical: [None, 5, 10, 15], p=[0.85, 0.05, 0.05, 0.05]
XGBoost	Max depth	Integer Uniform: 1 $\rightarrow$ 11
	Number of estimators	Integer Uniform: 100 $\rightarrow$ 1000
	Min child weight	Integer Log Uniform: 1 $\rightarrow$ 100
	Subsample	Real Uniform: 0.5 $\rightarrow$ 1.0
	Learning rate	Real Log Uniform: 1e-5 $\rightarrow$ 0.7
	Col sample by level	Real Uniform: 0.5 $\rightarrow$ 1.0
	Col sample by tree	Real Uniform: 0.5 $\rightarrow$ 1.0
	Gamma	Real Log Uniform: 1e-8 $\rightarrow$ 7.0
	Lambda	Real Log Uniform: 1.0 $\rightarrow$ 4.0
	Alpha	Real Log Uniform: 1e-8 $\rightarrow$ 100.0

Model	Parameter	Values considered for hyperparameter tuning
CatBoost	Max depth	Integer Uniform: 3 $\rightarrow$ 10
	Learning rate	Real Log Uniform: 1e-5 $\rightarrow$ 1.0
	Bagging temperature	Real Uniform: 0.0 $\rightarrow$ 1.0
	L2 leaf regression	Real Log Uniform: 1.0 $\rightarrow$ 10.0
	Leaf estimation iterations	Integer Uniform: 1 $\rightarrow$ 10
TabNet	Number decision steps	Categorical: [3, 5, 10]
	Layer size	Categorical: [8, 16, 64]
	Learning rate	Categorical: [0.01, 0.02]
Wide&Deep	Layer 1 size	Categorical: [64, 128, 256]
	Layer 2 size	Categorical: [64, 128, 256]
	Layer 3 size	Categorical: [None, 64, 128, 256]
	Dropout ratio	Categorical: [0.0, 0.01, 0.05]
	Embedding dimension	Categorical: [4, 16, 32]
DeepFM	Layer 1 size	Categorical: [64, 128, 256]
	Layer 2 size	Categorical: [64, 128, 256]
	Layer 3 size	Categorical: [None, 64, 128, 256]
	Use batch normalization	Categorical: [True, False]
	Dropout ratio	Categorical: [0.0, 0.01, 0.05]
	Embedding dimension	Categorical: [4, 16, 32]

V. Used Hyperparameters

Model	Parameter	Default	Used for Subset					
			Position	Position + Keyword	Position + SERP Feat	Position + Result	Position + Keyword + SERP Feat	Position + Keyword + SERP Feat + Result
Random Forest	Max depth	None	None	None	None	None	None	None
	Number of estimators	100	897	823	3000	3000	175	516
	Criterion	squared_error	squared_error	squared_error	squared_error	squared_error	squared_error	squared_error
	Max features	1.0	0.1	sqrt	0.3	0.4	0.3	0.6
	Min samples split	2	3	2	2	3	2	2
	Min samples leaf	1	38	5	8	9	1	2
	Bootstrap	True	True	False	False	True	False	True
	Min impurity decrease	0.0	0.0	0.0	0.0	0.0	0.0	0.0
GBDT	Loss	squared_error	squared_error	squared_error	squared_error	squared_error	squared_error	huber
	Learning rate	0.10	0.067	0.033	0.030	0.351	0.126	0.767
	Subsample	1.0	0.517	1.0	0.711	0.833	0.634	0.819
	Number of estimators	100	128	497	220	196	487	268
	Criterion	friedman_mse	friedman_mse	friedman_mse	friedman_mse	friedman_mse	friedman_mse	friedman_mse
	Max depth	3	2	None	None	4	5	2
	Min samples split	2	3	2	2	3	3	2
	Min samples leaf	1	24	50	1	10	24	9
	Min impurity decrease	0.0	0.010	0.050	0.050	0.010	0.010	0.010
	Max leaf nodes	None	15	None	15	5	5	15
XGBoost	Max depth	6	11	11	8	11	11	11
	Number of estimators	100	406	689	603	786	460	714
	Min child weight	1	100	1	1	1	1	1
	Subsample	1.0	1.0	0.993	1.0	1.0	1.0	0.668
	Learning rate	0.30	0.015	0.030	0.035	0.016	0.039	0.031
	Col sample by level	1.0	0.50	0.50	0.50	0.50	0.917	0.50
	Col sample by tree	1.0	1.0	1.0	0.50	0.50	0.708	1.0

			Used for Subset					
Model	Parameter	Default	Position	Position + Keyword	Position + SERP Feat	Position + Result	Position + Keyword + SERP Feat	Position + Keyword + SERP Feat + Result
	Gamma	0.0	1.0e-08	7.7e-05	0.009	1.0e-08	1.0e-08	0.001
	Lambda	1.0	4.0	4.0	2.859	1.0	4.0	1.0
	Alpha	0.0	1.0e-08	0.080	0.113	1.0e-08	0.014	1.0e-08
CatBoost	Max depth	6	5	10	9	9	10	9
	Learning rate	0.030	0.013	0.041	0.033	0.026	0.056	0.054
	Bagging temperature	1.0	0.885	1.0	1.0	0.0	1.0	0.0
	L2 leaf regression	3.0	6.601	10.0	10.0	10.0	10.0	10.0
	Leaf estimation iterations	<i>Dynamic</i>	10	10	1	6	10	10
TabNet	Num decision steps	3	3	3	3	3	3	3
	Layer size	8	8	8	8	8	8	64
	Learning rate	0.02	0.02	0.02	0.02	0.02	0.02	0.025
Wide & Deep	Layer 1 size	256	256	256	256	256	256	256
	Layer 2 size	128	128	128	128	128	128	128
	Layer 3 size	64	64	64	64	64	64	64
	Dropout ratio	0.0	0.0	0.0	0.0	0.0	0.0	0.01
	Embedding dimension	4	4	4	4	4	4	16
DeepFM	Layer 1 size	256	256	256	256	256	256	256
	Layer 2 size	128	128	128	128	128	128	128
	Layer 3 size	64	64	64	64	64	64	64
	Use batch normalization	False	False	False	False	False	False	False
	Dropout ratio	0.0	0.0	0.0	0.0	0.0	0.0	0.01
	Embedding dimension	4	4	4	4	4	4	16

VI. Improvement Through Hyperparameter Tuning

		Improvement over default hyperparameters					
		Position	Position + Keyword	Position + SERP Feat	Position + Result	Position + Keyword + SERP Feat	Position + Keyword + SERP Feat + Result
Linear Regression	OLS	-	-	-	-	-	-
	Poly2	-	-	-	-	-	-
	Ridge	-	-	-	-	-	-
Tree-Based	Random Forest	0.009	0.003	0.007	0.008	0.003	0.01
	GBDT	0.000	0.009	0.002	0.003	0.007	0.007
	XGBoost	0.001	0.005	0.001	0.001	0.004	0.003
	CatBoost	0.000	0.003	0.000	0.000	0.002	0.001
Neural Network	TabNet	-	-	-	-	-	0.033
	Wide&Deep	-	-	-	-	-	0.010
	DeepFM	-	-	-	-	-	0.013

C.I. Data Acquisition

Since scraping the Google SERP is both tedious and restricted by several hurdles imposed by Google, it was impractical to get the title for 900k data points this way. Instead, the actual URLs were scraped and the text of the `<title>` meta tag was extracted. In order to get results as close to what the users actually saw, the page requests were conducted using a VPN, emulating an IP address for either the US or UK. In order to overcome as many anti-scraping measures imposed by websites as possible, two types of request headers were used: One imitating the request of a human user in a desktop web browser and the other imitating a request by the Google crawler. Responses with a status code other than “200” or text patterns

such as “404” or “are you a bot?” in the title, were considered denied. As a result, for 81.7% of all unique URLs, a successful response with a valid title could be obtained.

C.II. Comparison of HTML vs. displayed titles & descriptions

In order to validate that the HTML title tag also resembles the title shown in the Google SERP, 2,000 sample Google SERPs for keyword / URL combinations were collected. The Google SERP title was compared to the HTML title using normalized Levenshtein similarity (Yujian and Bo 2007), a string comparison metric that measures how many characters of two strings need to be changed to make them equal. Scores range from zero to one with one being perfectly equal strings. For both titles, only the first 50 characters were compared, to only consider the minimum shown part of the title in the SERP. Two titles with Levenshtein similarity > 0.5 were considered sufficiently equal. As a result, 77.6% of all titles were found to be equal. In addition, varying patterns between domains could be observed: while for some domains a large number of titles had a similarity close to one, other domains had solely highly unequal titles. Based on this, 9 domains with more than 30 comparison samples where more than 50% of all titles were unequal, could be identified and removed. The impacted domains are:

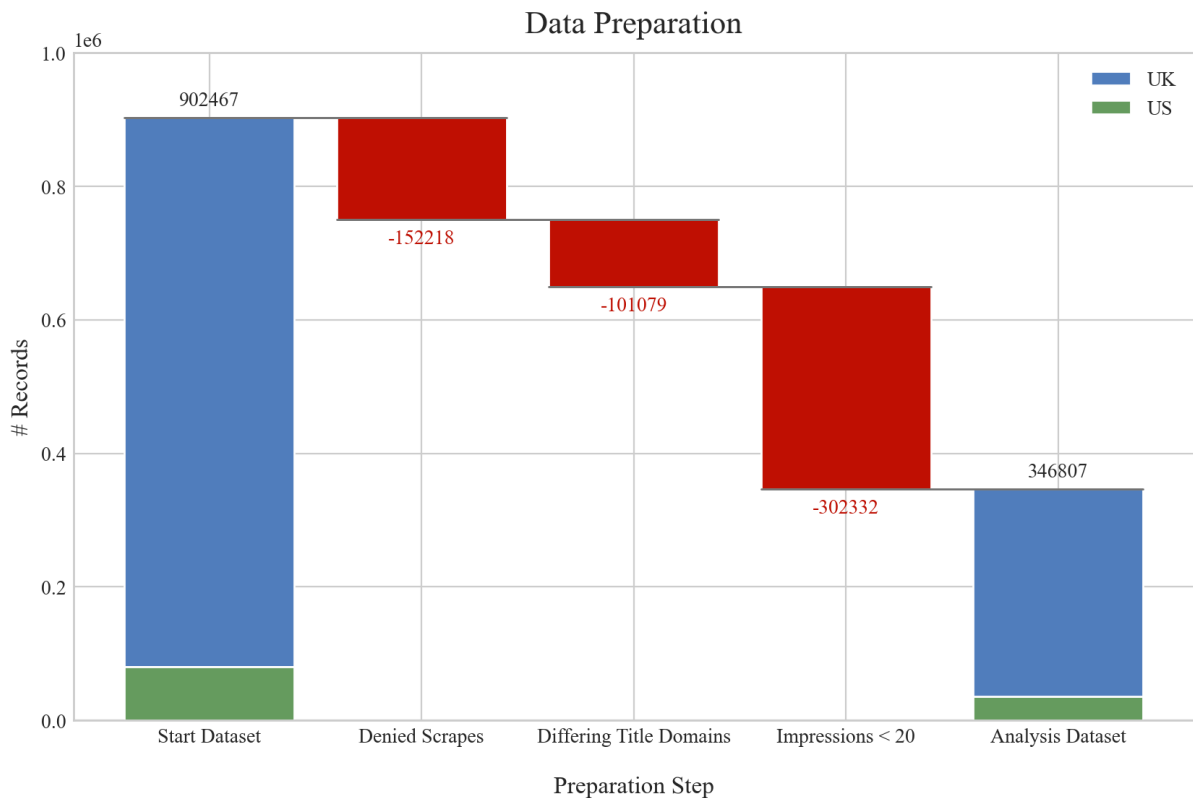
- 'aquarterof.co.uk'
- 'bissell.com'
- 'champneys.com'
- 'hotelchocolat.com'
- 'katespade.com'
- 'sainsburys.co.uk'
- 'spirithalloween.com'
- 'studio.co.uk'

- 'vizio.com'

By removing these domains, the expected share of similar titles could be increased by 11.8% to 89.4%, which was considered a sufficiently large share to confidently assume that the titles in the HTML tags were the one shown to the user during data generation.

The same methodology was applied to descriptions. The first 150 description characters in the HTML `<meta>` tag were compared to the first 150 characters on the Google SERP description. In only 31.6% of all cases, both descriptions were sufficiently similar. Furthermore, no frequently differing domains whose removal could increase the expected equality, such as for titles, could be identified.

C.III. Dataset preparation data loss



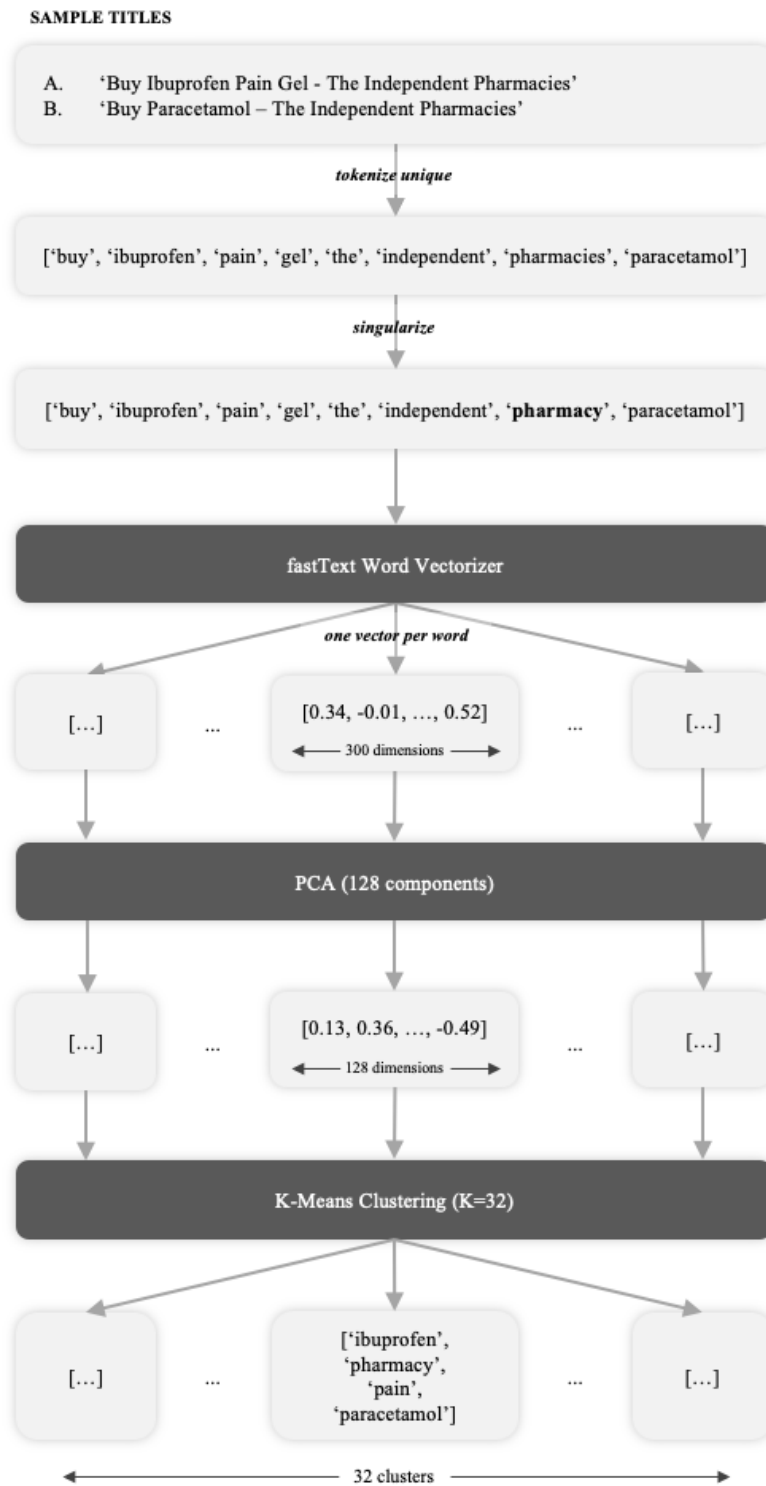
C.IV. Title Features Data Dictionary

Feature	Description	Example	
		Input	Resulting Value
Brand in title	Whether the brand name appears in the title. Checks whether the domain name (except the top-level domain, i.e. '.com', '.co.uk', etc.) is featured in the title.	Title = "Shoes Amazing Shop Inc."; Domain = "amazingshop.co.uk"	1
Count acronyms	Count of acronyms in title. Similar to Biyani, Tsioutsoulouklis, and Blackmer (2016), an acronym is considered as a series of two to four capital letters.	Title = "This UK food is AMAZING"	1
Count capitalized	Number of capitalized words in title. In line with Biyani, Tsioutsoulouklis, and Blackmer (2016), capitalized words are words with a minimum of 5 all-capitalized characters.	Title = "This UK food is AMAZING and CHEAP"	2
Count dashes	Number of appearances of '-' in the title. Does not consider hyphens in-between words.	Title = "3-dimensional glasses - The Glass Shop"	1
Count pipes	Number of appearances of ' ' in the title.	Title = "iPhone 14 Phones Apple"	2
Count superlatives*	Number of superlatives in the title. A word is considered a superlative, if its POS-Tag lies within ['RBS', and 'JJS'].	Title = "The largest and best Online Front Bicycle Basket Shop!"	2
Has alliteration	Count of alliterations in the title. Considers two subsequent words that start with the same letter as alliteration.	Title = "Four free tips that help increasing interaction"	1
Has exclamation mark	Whether an exclamation mark (!) is present in the title.	Title = "The largest Online Rear Bicycle Basket Shop!"	1
Has number	Whether the title contains a number. Number can be infinitely large, but must be numeric, i.e. written as a series of '0' to '9'.	Title = "5 Amazing tips to lose weight"	1
Has question mark	Whether a question mark (?) is present in the title.	Title = "How can I get more clicks on Google?"	1
Has serial number	Whether the title contains a serial number. Serial numbers are considered as combinations of letters and numbers, eventually connected through a hyphen, as well as numbers with at least 5 digits.	Title = "27" FHD Curved Monitor (C27FG73) Gaming Monitor"	1
Has unit specification	Whether the title contains a unit specification. Unit specifications are extracted with the python package quantulum3, where accepted unit dimensions are currency, volume, length/area, mass, power, angle, and data storage.	Title = "iPhone 14 256GB"	1
Keyword/title overlap	The share of words in the keyword that overlap with the title. Words in both title and keyword are stemmed using the Porter Stemming algorithm (Porter 1980) and overlaps of bi-grams with unigrams (and vice versa) are accepted.	Keyword = "i phone 14 chargingcable"; Title = "iPhone Charging Cable"	0.75

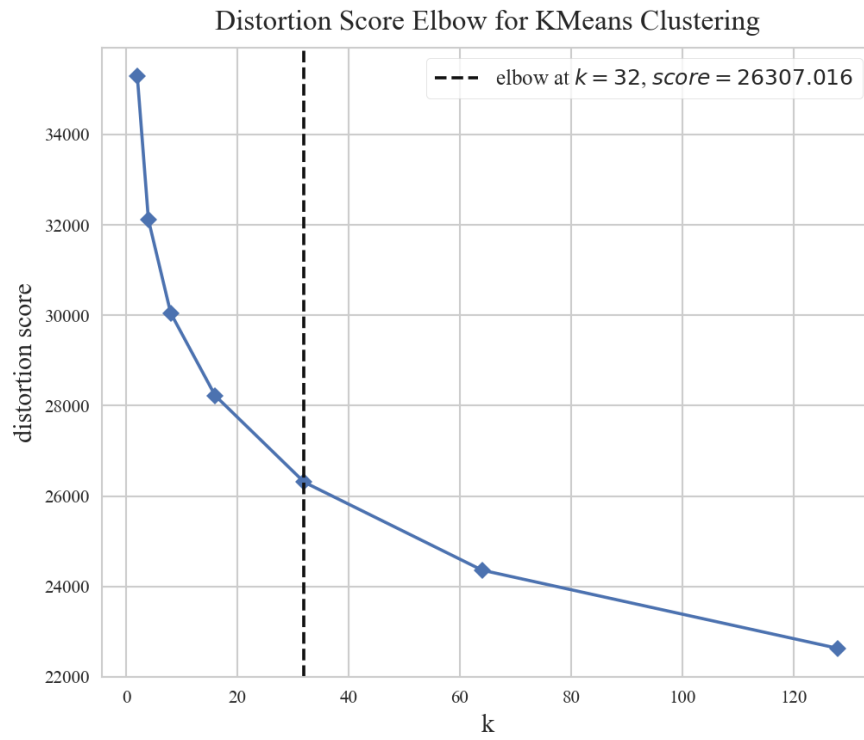
Example			
Feature	Description	Input	Resulting Value
Technicality score	Measure of syntactical technicality. Computed as (noun freq + adjective freq + number freq - preposition freq - article freq - pronoun freq - verb freq - adverb freq - interjection freq + 100) / 2.	Title = “Office Trolleys AJ Products”	52.0
Truncated	Whether a title exceeds 60 characters.	Title = “10 cool tips to create short titles”	0
Verb at start*	Whether the title starts with a verb in present tense. A word is considered in present tense, if its POS-Tag lies within ['VB', 'VBP', and 'VBZ'].	Title = “Buy the best wines in town - Winery UK”	1

* For the computation of the variable, part-of-speech (POS) tags are used. A POS tag is a grammatical classification that commonly includes verbs, adjectives, adverbs, nouns, or similar (Chiche and Yitagesu 2022). In this work, POS tags are generated using the implementation of the NLTK Python library.

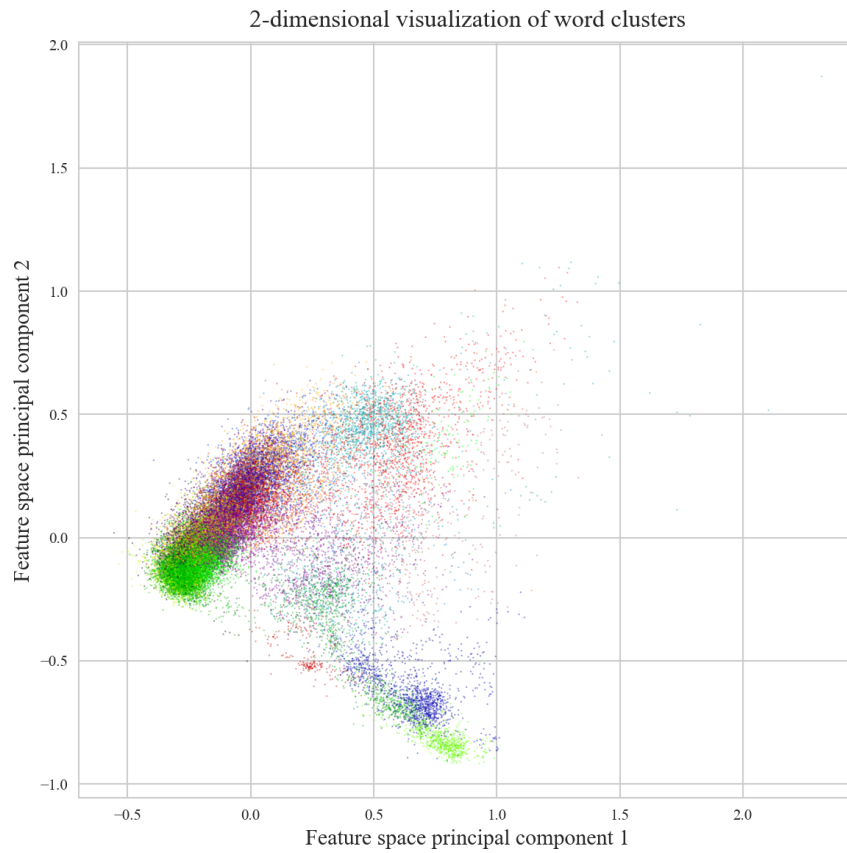
C.V. Clustering Process Visualization



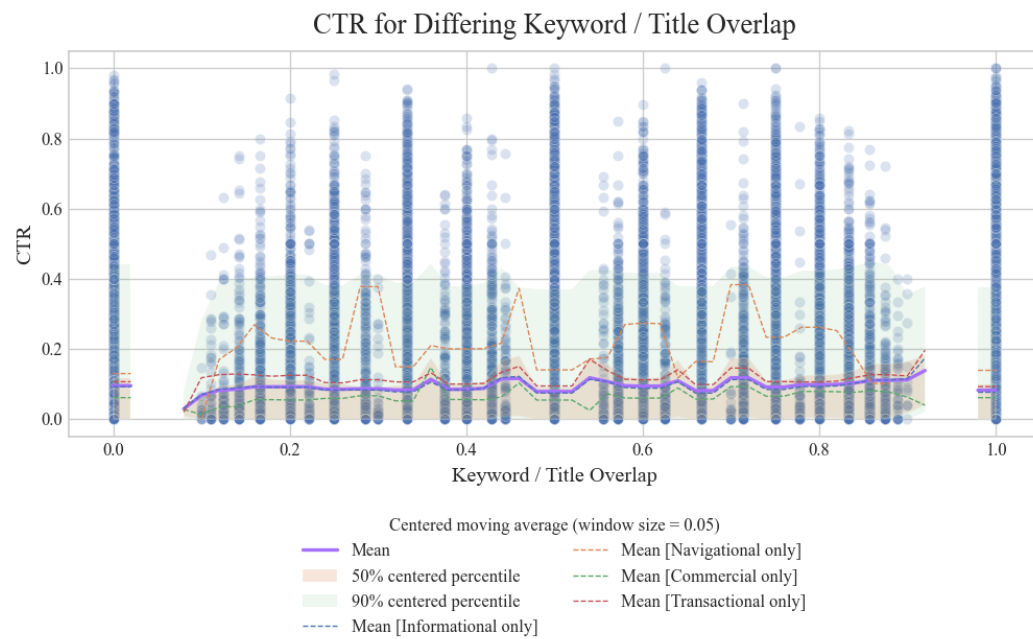
C.VI. Word Clustering Elbow Determination



C.VII. Clustering Result Visualization



C.VIII. CTR for differing CTR keyword/title overlap, split by intents



C.IX. Features used for modeling

Feature Subset	Features
w/o word clusters	• Has alliteration • Count superlatives • Technicality score • Verb at start • Has exclamation mark • Has question mark • Truncated • Count acronyms • Count capitalized • Has number • Has unit specification • Has serial number • Count pipes (‘ ‘) • Count dashes (‘ - ‘) • Keyword/title overlap • Brand in title
Title features added	• Has alliteration • Count superlatives • Technicality score • Verb at start • Has exclamation mark • Has question mark • Truncated • Count acronyms • Count capitalized words • Has number • Has unit specification • Has serial number • Count pipes (‘ ‘) • Count dashes (‘ - ‘) • Keyword/title overlap • Brand in title • <i>Count of words for each of the 32 clusters</i>