

SERP Interference Network and its Applications in Search Advertising

Purak Jain
Amazon
Seattle, USA
purakjn@amazon.com

Sandeep Appala
Amazon
Seattle, USA
appalar@amazon.com

ABSTRACT

Search Engine marketing teams in the e-commerce industry manage global search engine traffic to their websites with the aim to optimize long-term profitability by delivering the best possible customer experience on Search Engine Results Pages (SERPs). In order to do so, they need to run continuous and rapid Search Marketing A/B tests to continuously evolve and improve their products. However, unlike typical e-commerce A/B tests that can randomize based on customer identification, their tests face the challenge of anonymized users on search engines. On the other hand, simply randomizing on products violates Stable Unit Treatment Value Assumption for most treatments of interest. In this work, we propose leveraging censored observational data to construct bipartite (Search Query to Product Ad or Text Ad) SERP interference networks. Using a novel weighting function, we create weighted projections to form unipartite graphs which can then be used to create clusters to randomized on. We demonstrate this experimental design's application in evaluating a new bidding algorithm for Paid Search. Additionally, we provide a blueprint of a novel system architecture utilizing SageMaker which enables polyglot programming to implement each component of the experimental framework.

CCS CONCEPTS

• **Information systems** → *Computational advertising*; **Sponsored search advertising**; • **General and reference** → Experimentation.

KEYWORDS

sponsored search, experiment design, A/B testing.

ACM Reference Format:

Purak Jain and Sandeep Appala. 2024. SERP Interference Network and its Applications in Search Advertising. In . ACM, Barcelona, Spain, 6 pages. <https://doi.org/XXXXXXX.XXXXXXX>

1 INTRODUCTION

Search Engine marketing teams in the e-commerce industry manage global search engine traffic with the aim of optimizing long-term profitability by delivering the best possible customer experience

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.
AdKDD'24, August 2024, Barcelona, Spain

© 2024 Copyright held by the owner/author(s). Publication rights licensed to ACM.
ACM ISBN 978-x-xxxx-xxxx-x/YY/MM
<https://doi.org/XXXXXXX.XXXXXXX>

on the most important web pages on the internet - Search Engine Results Pages (SERPs). Figure 1 shows the prominent parts of SERP. Search Engines continue to evolve their customer experience and features due to social, technological and economic forces, including privacy concerns, and further monetization of their properties (SERPs). In anticipation of opportunities and risks that come with a shifting landscape advertisers continuously innovate with new bidding algorithms, improved paid and free search creatives, landing pages etc. Randomized experiments, or A/B tests, are the standard approach for evaluating causal effects of new features [18]. However, Search Marketing experiments are unlike conventional A/B tests in industry that can randomize on customers as advertisers don't identify their customers when they are on a search engine i.e. the ad publisher. Instead, advertisers may run A/B tests randomized by geographic locations [29] using search engine's geo-targeting capabilities but due to ad publisher's API limitations they are unable to do so without having to clone entire advertisement campaigns. The cloning of entire accounts is operationally expensive and time consuming restricting the velocity at which they can run such trials.

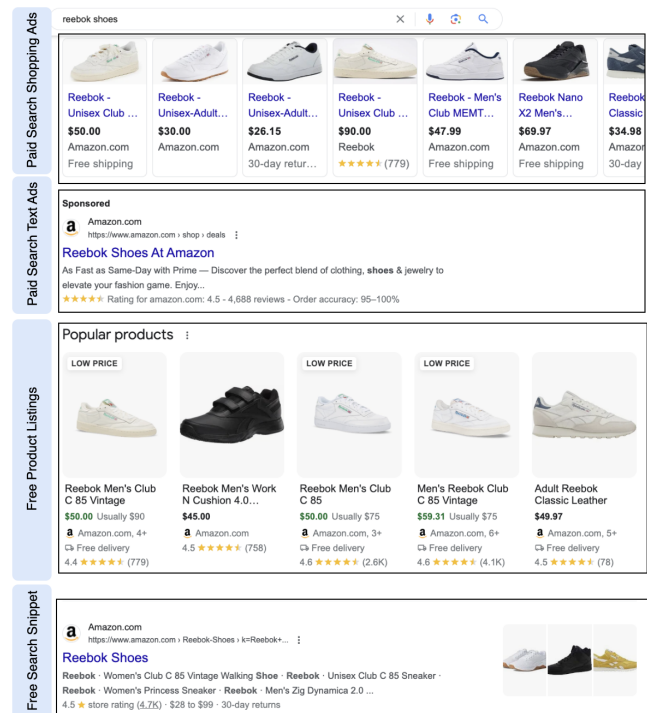


Figure 1: Components of interest on a SERP.

The next obvious choice for unit of randomization is usually products or search queries. However for any A/B test, splits of the unit of randomization should satisfy the assumptions of the Neyman - Rubin causal framework [25], that is, no interference, unconfoundness, overlap and no hidden treatment variations. For example, if we simply randomly select products into control and treatment groups, it should hold the unconfoundness and the overlap assumption given a large sample size but we still need to check if the "no interference" assumption holds. We observe that for most treatments of interest (e.g. new bidding algorithms for paid search programs or improved title headlines for free search snippets), the SERP page leads to interference between treatment and control units causing the Stable Unit Treatment Value Assumption (SUTVA) to fail, and consequently induces bias in the standard estimators used to evaluate the value generated by the treatment. A standard answer [7, 10] to this problem is to replace the "product-split" experiment design with a "time-split" (or "switchback") design, where the entire market switches repeatedly between treatment and control. In practice, such designs turn out to be equally time consuming as geo-based splits since we need to account for long lengths of adjustment period between switches due to the presence of an intermediary i.e. search engine that applies the treatment and takes its own time which advertisers cannot control.

Another approach to dealing with spillovers or interference is given by clustered experiments [3, 24] or in social network settings, by network bucketing testing [4] where nodes that are relatively clustered together are given the same assignment of treatment or control [28]. Our work is inspired by similar clustered experiments methods that have been applied to estimate and reduce bias in marketplace experiments [13]. Specifically, one such example involves experimentation in internet ad auctions, where each auction consists of a keyword along with a set of advertisers who submit competing bids in order for their ads to be displayed when the keyword is queried by a user. There is cross-unit interference because the same advertiser or keyword may appear in multiple auctions. Basse et al. [6] and Ostrovsky and Schwarz [23] make the observation that the auction type used for one keyword does not meaningfully affect how advertisers bid for other keywords. They then consider experiments that group auctions into clusters by their keywords and randomize auction formats across these keyword clusters, rather than across advertisers, as a means to avoid problems with interference. More broadly, in our context of Search Marketing, this idea of cluster-level randomization corresponds to identifying product or search query clusters that are relatively isolated from each other and randomizing the interventions across product-clusters rather than across products. Our primary contribution lies in leveraging observational data to build bipartite (Search Query - Product) and tripartite (Search Query - Paid Search Product - Free Search URL) SERP interference networks. We introduce an innovative weight function to generate weighted projections, transforming these networks into unipartite graphs. These graphs facilitate the clustering of products that co-appear on SERPs through Paid Search Shopping Ads, Text Ads, or Free Product Listings. The resultant clusters can then be randomized during A/B tests to generate insights.

Note that more recently, Johari et al. [15] and Bajari et al. [5] have proposed newer experiment designs where both search query and product units are randomized simultaneously. While having a

similar flavor, neither framework applies easily to our problem of interest. To begin with, we cannot control "search-query" assignment as that is determined by the search engine i.e. the ad publisher. Johari et al. [15] use a choice model to capture spillovers, which captures a different kind of market than the one we consider, where interference is mediated by a matching algorithm. Bajari et al. [5] imposes a local interaction assumption, which does not hold in our setting. However, when the graph is a bi-partite graph it holds some similarity which we plan to explore in future work for measuring the magnitude of spillovers.

The rest of this paper proceeds as follows. In Section 2, we use the two-population search query - product case as a motivation to build SERP interference network to test out new bidding algorithms. In Section 3, we describe in greater detail our experiment design. In Section 4, we further share details on testing a new bidding model using this experimentation design. Section 5 provides an overview of a system architecture blueprint for deploying such experimentation frameworks. Finally, we discuss our findings and future extensions in Section 6.

2 SETTING AND MOTIVATION

Shopping Ads is one of the ad formats supported on SERPs. To place ads within the shopping ad carousel advertisers need to participate in an auction competing with other advertisers. The format of the auction is considered close to second price Vickrey-Clarke-Groves (VCG) [9, 12, 30], although the exact ad publisher implementation is a blackbox for us. As such to maximize long term profitability, it is important to constantly develop, test and launch new bidding algorithms responsible for valuating products worldwide. Let's say, to test out a new bidding algorithm we simply split on products. The Stable Unit Treatment Value Assumption (SUTVA) presumes that the valuation assigned to one product by the new algorithm does not influence the profitability of other products. However, in the context of shopping advertisements, this assumption may be violated due to potential between-product interference. This interference occurs when both a product with a treatment bid from the new model and another with a control bid from the current model participate in the same auction triggered by a search query, deemed relevant by the ad publisher for both products. See figure 2 for an example. Such scenarios clearly breach SUTVA, challenging the validity of our evaluation method. If one product happens to be assigned to treatment group and the other one to control, then the difference in financial performance between the two products will be resulted from the combined effect of treatment and between-product spillover effect, thus making the treatment effect indistinguishable from the product spillover effect.

3 PRODUCT-CLUSTER RANDOMIZED CONTROL TRIAL DESIGN

To address the challenge of interference in experimental designs, we propose a preemptive modeling strategy that incorporates interference networks during the design phase. This approach allows us to shift the unit of randomization from individual products to clusters of products, as illustrated in Figure 3. Importantly, traditional constrained randomization methods [20], such as segmenting by

Table 1: Mock row in a Search Engine's Query Report shared with the advertiser.

Search Query	Impressions	Product/Keyword	Clicks	Metric Day	Ad Campaign
chopper axe	6	B00EOVRX06	1	2023-09-01	12345

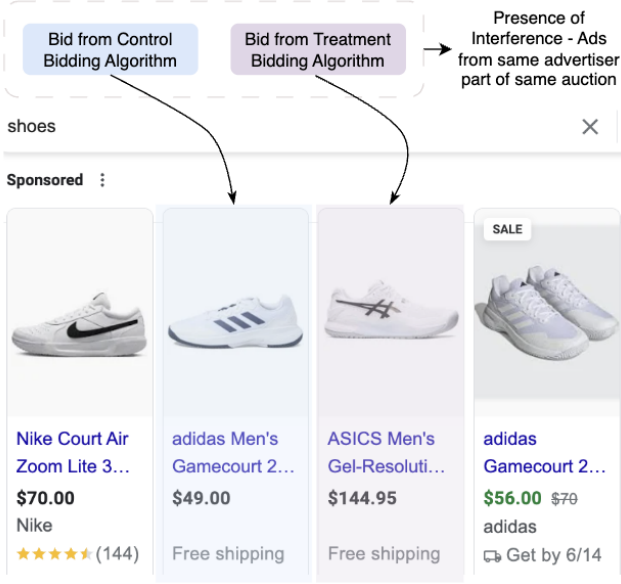


Figure 2: Note the two Shopping Ad advertisements shown in the same carousel. If these two products have bids from different bidders (i.e. control and treatment) then SUTVA is violated.

product categories, prove ineffective. This is because search engines can associate broad upper funnel search queries (e.g., "Harry Potter") with a diverse range of products across multiple categories (e.g., a book, toy, or blanket related to Harry Potter). By leveraging interference networks, our method ensures more robust and accurate experimental outcomes.

Modeling Network Interference The notion of interference in the network [14] we construct has to be aligned with the notion of interference we are trying to estimate. Since the relevance of products to a user search query is determined by the search engine and ranking algorithm, an advertiser cannot use its internal datasets that provide product to keyword mapping e.g. e-commerce website's own search to product results. Instead, we use daily reports provided by the ad publisher itself. These reports have information on which actual user search query on the search engine was mapped to which shopping ads product by the ad publisher. A sample mock row from such a report is shown in Table 1. We use these search query reports to construct an undirected bipartite search query - product graph using the number of impressions as edge weight.

Unipartite Projection To apply one mode projection of the bipartite graph onto the product nodes in order to model the between-network interference, we needed a scoring function to attribute

weights to the resulting graph edges. Since, we start with large number of products (200M+), we could not directly use the edge weighting functions proposed by Stram et al. [27] due to the computational complexity. Instead we propose the following edge weight function that requires significantly less computation:

$$W_{\text{uni}}(a, b) = \sum_{i=1}^n \frac{1}{\log_e(f_{sq_i})} \frac{\min(W_{bi}(sq_i, a), W_{bi}(sq_i, b))}{\max(W_{bi}(sq_i, a), W_{bi}(sq_i, b))} I[sq_i, a, b] \quad (1)$$

where:

- (1) $W_{\text{uni}}(a, b)$ is the edge weight in the unipartite graph (one-mode projection) between product a and b .
- (2) $W_{bi}(sq_i, a)$ is edge weight between search query i and a in the original bi-partite graph.
- (3) f_{sq_i} is the number of distinct products that a particular search query drives impressions to. Since the distribution is right-skewed i.e. few upper funnel queries drive impressions to only a few distinct products, weighing down by the log of frequency of search query helped us to weigh down edge weight contributions between two products from very generic queries.
- (4) $I[sq_i, a, b]$ is 1 if search query sq_i trigger an impression for both product a and b as represented by the presence of an edge in the original bipartite graph, otherwise 0.

Using the above approach to take a weighted one-mode projection of the bipartite graph leads to an increase in the number of edges, since if a search query links to n products, we need to consider $\binom{n}{2}$ pairs of edges.

Graph Partitioning Methodology Constructing a product graph following the above approach then allows us to use network dismantling algorithms [8] as oppose to naive connected components approach to creating product clusters. Historically, the community identification problem is a well studied problem in computer science literature [11, 21, 22]. However, a lot of the proposed methods wouldn't scale up since we are dealing with graphs that are as large as 200M+ nodes and 400M+ edges. Thus, anything that runs in $O(|V|^2)$ or $O(|E|^2)$ is not practical. Moreover during the graph partitioning phase, we needed to find a balance between two objectives:

- (1) Maximize the number of product clusters as they translate to randomization units. More clusters equal more power for our test.
- (2) Minimize the between clusters edge weights.

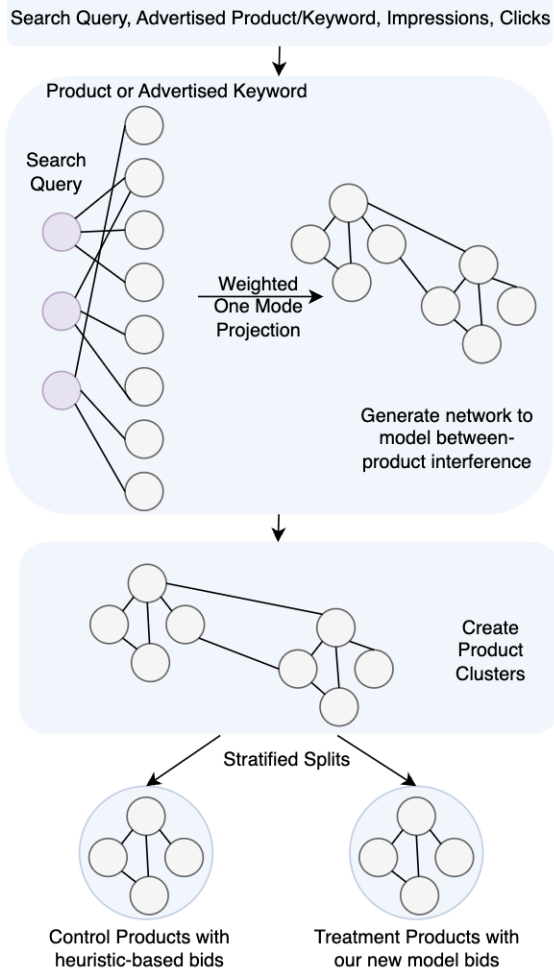


Figure 3: Birds-eye view of our experiment design. First, we fetch the search query reports from the ad publisher and create a search-query, product bipartite graph. Then we take a weighted one mode projection and finally use a clustering algorithm to cluster products. We then do a stratified random split of product clusters.

Since, the more clusters we create, the less isolated they are, these two objectives are conflicting. For example, if we want to have zero connections across clusters, then the obvious solution is to have one cluster only. This, of course, would not lend itself to an A/B test. To balance the above two objectives, first we look at the percentage edge weight across k clusters ($C_1, \dots, C_k$) i.e. leakage as L :

$$L = \frac{\sum_{i=1}^k \sum_{j \in C_i, k \notin C_i} W_{\text{uni}}(j, k)}{\sum_{j,k} W_{\text{uni}}(j, k)}$$

Secondly, we use a clustering algorithm that is designed for balanced clustering, that is, all clusters should have roughly equal

size. We evaluated naive connected components, power iteration clustering (PIC) [19], and METIS [17]. Connected components minimizes leakage, but suffers from extreme imbalance. PIC improves cluster imbalance, but suffers from high leakage. We choose to use METIS partitioning algorithm which is an extremely efficient and fast implementation of graph partitioning algorithm for undirected weighted graph. METIS adopts an objective function to minimize the number of weighted edges whose vertices belong to different partitions. The METIS graph partitioning consists of three phases: (i) In the graph coarsening phase, a series of successively smaller graphs is derived from the input graph. This process continues until the size of the graph has been reduced to just a few hundred vertices, (ii) In the initial partitioning phase, a partitioning of the coarsest and hence, smallest, graph is computed and finally (iii) in the un-coarsening phase, the partitioning of the smallest graph is projected to the successively larger graphs by assigning the pairs of vertices that were collapsed together to the same partition. After each projection step, the partitioning is refined using heuristics to iteratively move vertices between partitions as long as such moves improve the quality of the partitioning. The advantages of this methodology are threefold:

- (1) It runs in $O(|E|)$ time, which is extremely efficient for large graphs.
- (2) It is the only algorithm that allows precise control of both the number partitions and the balances of the overall split.
- (3) It is the only algorithm that is specifically trying to minimize the edgecut (defined as weighted sum of edges that straddle between different clusters).

Optimal number of clusters: We want to be able to identify as many nearly independent clusters as possible with leakage controlled within the tolerance. We plot leakage against various choice of k (number of partitions), and identify a k that is as large as possible where the leakage is as small as possible (i.e. identifying the elbow point). For the new shopping ad bidder experiment, we ended up having 10,000 clusters and 36% edge weight across clusters. Note, the above measure (L) overstates the spillover effects as they consider spillover between clusters that may end up being in the same group (C or T).

Magnitude of Spillover: The search query - product bipartite graph we construct usually has a clustering coefficient [20] of around ~ 0.6 for most marketplaces which indicates tightly knit groups in the network suggesting high spillover. However, to empirically provide a lower bound on the magnitude of bias due to interference we need to conduct a meta-experiment that randomizes over two experiment designs: one Bernoulli randomized, one cluster randomized. We can then check for a statistically significant difference between the total average treatment effect estimates obtained with the two designs [26]. In the absence of business approval to run such a meta-experiment, the next best directional data point we have is from our previous attempt to run simple product-split A/B test. The impact measured from that experiment had been largely overstated ($\sim 44\%$ lift) when compared to the actual lift ($\sim 24\%$ lift) observed.

4 APPLICATION

In this section, we discuss the use-case motivated in Section 2 to show an application of the product-cluster randomized control trial design. We had developed a new machine learning based product valuation model for our shopping ads program to improve over the current in production heuristic bidding algorithm and we wanted to run an online experiment to understand the impact on the long term profit. We used the methodology described in Section 3 to create product clusters based on search query reports from the ad publisher for the past one year.

Constrained Randomization Once we had the clusters, we created strata of clusters with similar characteristics instead of randomizing them in a simple bernoulli fashion. We measure the net impressions, clicks, cost and profit of each of the product cluster and stratify clusters on those axis.

Experiment Setup The goal of this experiment was verifying the null hypothesis that the new bidding strategy is better than the current bidding strategy in term of bidding efficiency i.e. increase of net long term profitability while maintaining the total ad spend. We matched the spend between control and treatment groups to control for elasticity as well as to comply with spend constraints at account level. Finally, we run a simulation based power analysis for cluster randomized designs using difference-in-differences (DID) estimation [1].

Measurement To measure the impact of the proposed valuation method, a DID analysis for cluster randomized designs is performed for two weeks of periods where spends are closely matched. The results from the DID analysis showed a lift in click-through-rate for the treatment group which was consistent with the lift observed post roll out of the new bidding model. Note that since model errors can be correlated within cluster, failure to control for within-cluster error correlation can lead to misleading small standard error and consequently low p-values. Although we do not control for within-cluster error correlation in the model, post-estimation we obtain cluster-robust standard errors as proposed by White [31].

5 SYSTEM ARCHITECTURE

Our product-cluster randomized control methodology as detailed in Section 3 asks for a highly scalable and flexible infrastructure with very different compute requirements and library support for each step. To address these challenges we propose the "Search Marketing Lab" using AWS SageMaker [16] pipelines which allows to define a series of interconnected processing steps where each step (i) can be provided its own docker image that has our code in preferred language and (ii) can have its own compute environment. This allows for polyglot programming. Here, we briefly focus on the split generation component. In particular, we break the approach into 3 modules:

- (1) **Graph Generation** which requires parsing >1 year of daily search query report data and Keyword data to build a graph

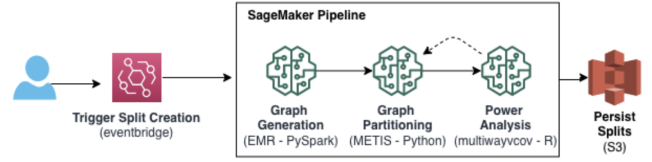


Figure 4: Birds-eye view of Product-cluster split generation system.

edge list - thus requiring spark's distributed compute. We use a cluster of memory-optimized instances for this step.

- (2) **Graph Partitioning** In this step we use the METIS graph partitioning algorithm. METIS is written entirely in ANSI C (no distributed implementation) but there is a python wrapper for the METIS library [17] that we use. We use a single compute-optimized instance for this step.
- (3) **Power Analysis** In this module we obtain cluster-robust standard errors after fitting a linear mixed effect model, using R's cluster.vcov [2] implementation to return a multi-way cluster-robust variance-covariance matrix and perform inference for estimated coefficients using R's coeftest.

SageMaker Pipeline executions can be scheduled using Amazon EventBridge passing run-time parameters. This allows to define a single pipeline with multiple executions (e.g. one per marketplace) based on input parameters. The serves as a blueprint for a large scale production system combining multiple languages (R, Python on Spark) utilizing each to their respective strengths (R for statistical analysis modules, python on Spark for ETL) triggering SageMaker processing jobs orchestrated via SageMaker Pipelines.

6 CONCLUSION AND FUTURE WORK

In this paper, we present a cluster-based randomized control test design which enables search marketing e-commerce teams to do fast online experiment launch while minimizing interference between experimental groups. Our key idea is to use observational data to construct bipartite (Search Query - Product) SERP interference networks and use a novel weight function to take weighted projections to form unipartite graphs which can be use to create clusters of products appearing together on SERP (via Paid Search shopping ads, text ads or Free Search listings), and then using those clusters to randomize on. Online A/B testing results for the treatment group are consistent with the lift observed post roll out of a new bidding model thereby showing that the A/B test design gives a good estimate of the actual lift. In our previous attempts to run simple product-split A/B test the impact measured from experiments had been largely overstated because of spillover effects. Lastly, we present a novel simplified system architecture using SageMaker which allows scientist to do polyglot programming using compute and language suitable for each scientific module.

One downside of inferring interference network from search query report data is that such observational data is censored, that

is, we only have data when we win the auction. In future, we are investigating using SERP page data from platforms like seoClarity to get better visibility into SERP interference networks and allow us to incorporate not just shopping ad products but also Text Ads keyword and Free Search URLs to build comprehensive ad units spanning across all Search channels - Text Ads, Shopping Ads and Free Search. More recently, this also includes large language models powered results like shown in Appendix. We can then use these ad units to design cross-channel substitution experiments. We are also working on investigating further into the stability of these clusters over time and that they can be updated in real time as more data flow in from search engines. Finally, we are exploring recent proposed experiment designs by Bajari et al. [5] to measure the actual magnitude of spillovers.

REFERENCES

- [1] Joshua D Angrist and Jörn-Steffen Pischke. 2009. *Mostly harmless econometrics: An empiricist's companion*. Princeton university press.
- [2] Manuel Arellano. 1987. Computing robust standard errors for within-groups estimators. *Oxford Bulletin of Economics & Statistics* 49, 4 (1987).
- [3] Peter M Aronow and Cyrus Samii. 2017. Estimating average causal effects under general interference, with application to a social network experiment. (2017).
- [4] Lars Backstrom and Jon Kleinberg. 2011. Network bucket testing. In *Proceedings of the 20th international conference on World wide web*. 615–624.
- [5] Patrick Bajari, Brian Burdick, Guido W Imbens, Lorenzo Masoero, James McQueen, Thomas Richardson, and Ido M Rosen. 2021. Multiple randomization designs. *arXiv preprint arXiv:2112.13495* (2021).
- [6] Guillaume W Basse, Hossein Azari Soufiani, and Diane Lambert. 2016. Randomization and the pernicious effects of limited budgets on auction experiments. In *Artificial Intelligence and Statistics*. PMLR, 1412–1420.
- [7] Iavor Bojinov, David Simchi-Levi, and Jinglong Zhao. 2023. Design and analysis of switchback experiments. *Management Science* 69, 7 (2023), 3759–3777.
- [8] Alfredo Braunstein, Luca Dall'Asta, Guilhem Semerjian, and Lenka Zdeborová. 2016. Network dismantling. *Proceedings of the National Academy of Sciences* 113, 44 (2016), 12368–12373.
- [9] Edward H Clarke. 1971. Multipart pricing of public goods. *Public choice* (1971), 17–33.
- [10] Thomas D Cook and David L DeMets. 2007. *Introduction to statistical methods for clinical trials*. Chapman and Hall/CRC.
- [11] Michelle Girvan and Mark EJ Newman. 2002. Community structure in social and biological networks. *Proceedings of the national academy of sciences* 99, 12 (2002), 7821–7826.
- [12] Theodore Groves. 1973. Incentives in teams. *Econometrica: Journal of the Econometric Society* (1973), 617–631.
- [13] David Holtz, Ruben Lobel, Inessa Liskovich, and Sinan Aral. 2020. Reducing interference bias in online marketplace pricing experiments. *arXiv preprint arXiv:2004.12489* (2020).
- [14] Michael G Hudgens and M Elizabeth Halloran. 2008. Toward causal inference with interference. *J. Amer. Statist. Assoc.* 103, 482 (2008), 832–842.
- [15] Ramesh Johari, Hannah Li, Inessa Liskovich, and Gabriel Y Weintraub. 2022. Experimental design in two-sided platforms: An analysis of bias. *Management Science* 68, 10 (2022), 7069–7089.
- [16] Ameet V Joshi and Ameet V Joshi. 2020. Amazon's machine learning toolkit: Sagemaker. *Machine learning and artificial intelligence* (2020), 233–243.
- [17] George Karypis and Vipin Kumar. 1997. METIS: A software package for partitioning unstructured graphs, partitioning meshes, and computing fill-reducing orderings of sparse matrices. (1997).
- [18] Ron Kohavi, Alex Deng, Brian Frasca, Toby Walker, Ya Xu, and Nils Pohlmann. 2013. Online controlled experiments at large scale. In *Proceedings of the 19th ACM SIGKDD international conference on Knowledge discovery and data mining*. 1168–1176.
- [19] Frank Lin and William W Cohen. 2010. Power iteration clustering. (2010).
- [20] Lawrence H Moulton. 2004. Covariate-based constrained randomization of group-randomized trials. *Clinical trials* 1, 3 (2004), 297–305.
- [21] Mark EJ Newman. 2004. Detecting community structure in networks. *The European physical journal B* 38 (2004), 321–330.
- [22] Mark EJ Newman. 2004. Fast algorithm for detecting community structure in networks. *Physical review E* 69, 6 (2004), 066133.
- [23] Michael Ostrovsky and Michael Schwarz. 2011. Reserve prices in internet advertising auctions: A field experiment. In *Proceedings of the 12th ACM conference on Electronic commerce*. 59–60.
- [24] Chris Roberts and Stephen A Roberts. 2005. Design and analysis of clinical trials with clustering effects due to treatment. *Clinical trials* 2, 2 (2005), 152–162.
- [25] Donald Rubin. 2005. Causal Inference Using Potential Outcomes. *J. Amer. Statist. Assoc.* 100, 469 (2005), 322–331. <https://doi.org/10.1198/016214504000001880>
- [26] Martin Saveski, Jean Pouget-Abadie, Guillaume Saint-Jacques, Weitao Duan, Souvik Ghosh, Ya Xu, and Edoardo M Airolidi. 2017. Detecting network effects: Randomizing over randomized experiments. In *Proceedings of the 23rd ACM SIGKDD international conference on knowledge discovery and data mining*. 1027–1035.
- [27] Rotem Stram, Pascal Reuss, and Klaus-Dieter Althoff. 2017. Weighted one mode projection of a bipartite graph as a local similarity measure. In *Case-Based Reasoning Research and Development: 25th International Conference, ICCBR 2017, Trondheim, Norway, June 26–28, 2017, Proceedings 25*. Springer, 375–389.
- [28] Johan Ugander, Brian Karrer, Lars Backstrom, and Jon Kleinberg. 2013. Graph cluster randomization: Network exposure to multiple universes. In *Proceedings of the 19th ACM SIGKDD international conference on Knowledge discovery and data mining*. 329–337.
- [29] Jon Vaver and Jim Koehler. 2011. Measuring ad effectiveness using geo experiments. *Google Inc* (2011).
- [30] William Vickrey. 1961. Counterspeculation, auctions, and competitive sealed tenders. *The Journal of finance* 16, 1 (1961), 8–37.
- [31] Halbert White. 2014. *Asymptotic theory for econometricians*. Academic press.

A COMPONENTS OF INTEREST ON RECENT LLM POWERED SERPS

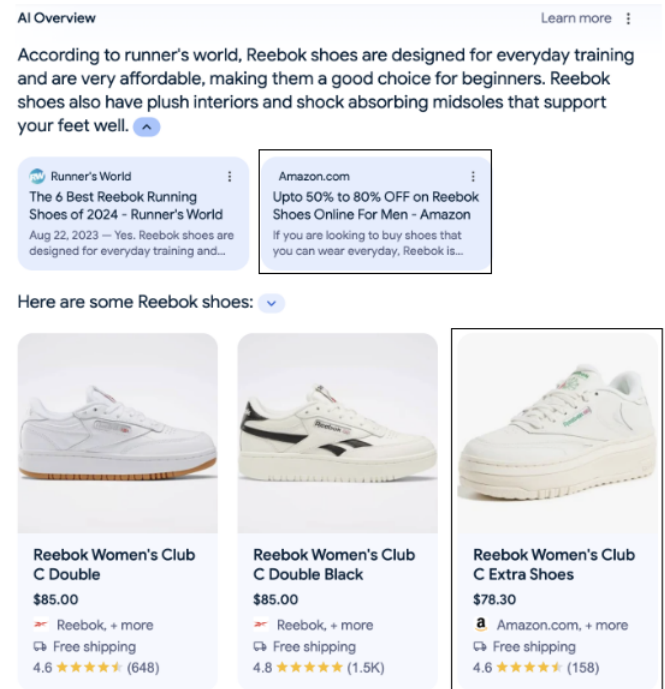


Figure 5: Evolving Large Language Model powered SERPs with components of interest highlighted in rectangle boxes.