

Constrained Target-ROAS Bidding for Sponsored Products via Value-Based Bids and Feedback Control

Wanchen Shao
Wayfair LLC
Mountain View, CA, USA
wshao@wayfair.com

Sergey Kolbin
Wayfair LLC
Seattle, WA, USA
skolbin@wayfair.com

Kurt Zimmer
Wayfair LLC
Boston, MA, USA
kzimmer@wayfair.com

Jinjin Zhao
Wayfair LLC
Seattle, WA, USA
jzhao5@wayfair.com

Jean-David Ruvini
Wayfair LLC
Mountain View, CA, USA
jruvini@wayfair.com

Abstract

Target return-on-ad-spend (tROAS) autobidding has become a dominant paradigm in large-scale advertising systems because it replaces manual bid tuning with a single, business-aligned objective: to deliver conversion value efficiently while meeting a return-on-investment target [1].

In the paper, we present an end-to-end tROAS system for Wayfair Sponsored Products in a multi-slot cost-per-click (CPC) auction, where an auction-time value-based bid is scaled by a campaign-level control parameter that is updated from delayed and noisy realized ROAS. We specify the interfaces among bid generation, score-based ranking, generalized second-price (GSP)-like charging, and campaign-level feedback control; introduce deployment guardrails for sparse-conversion campaigns, including cold-start handling, adaptive learning rates, and bounded multiplier updates; and quantify the intrinsic achievability of ROAS targets through an offline analysis. This analysis establishes a marketplace-specific noise floor for ROAS attainment and yields campaign-eligibility criteria based on data volume. Finally, we report findings from a large-scale online deployment and subsequent post-launch analysis, showing that 74% of enrolled campaigns with sufficient order volume achieved ROAS within $\pm 20\%$ of their target after the initial learning period, with attainment approaching the intrinsic achievability estimated offline. An observational study comparing tROAS to Wayfair’s manual bidding baseline finds tROAS delivers a 11-15% efficiency gain. In general, the paper contributes to a deployable, interpretable, and scalable tROAS system, together with the operational lessons required to make value-based bidding work in a sparse and high-variance marketplace.

CCS Concepts

• Information systems → Online advertising.

Keywords

Sponsored products advertising; autobidding; target ROAS; value-based bidding; feedback control

1 Introduction

On modern e-commerce marketplaces, sponsored-product auctions execute at high frequency across heterogeneous contexts—queries, placements, user segments, and time. As autobidding capabilities become standard across major advertising platforms, advertisers

expect automated bidding systems to deliver reliable goal-based optimization, making Target ROAS not only a performance lever but also a baseline requirement for platform competitiveness. They specify high-level goals such as Target CPA or Target ROAS, while automated agents compute auction-time decisions, shifting complexity from manual bid management to platform-side prediction and online control [1, 6, 10]

At Wayfair, we observed that when suppliers manage bids manually on a per-click basis, bid setting and campaign management can become time-consuming and imprecise. In practice, suppliers often rely on ROAS benchmarks as profitability guardrails: campaigns that fall below target may be paused, bid-reduced, or budget-limited, while campaigns that exceed expectations may receive additional budget or more aggressive bidding. It is therefore critical to offer a tROAS bidder to (i) simplify setup by replacing per-SKU manual bid tuning with a single campaign-level target, (ii) provide predictable performance over time despite noisy per-auction outcomes.

Meeting these technical challenges is further complicated by the characteristics of home-furnishings e-commerce platforms like Wayfair, where higher-priced, lower-frequency purchases (e.g., furniture versus everyday grocery items) lead to sparse conversion signals. This sparsity makes ROAS inherently noisier and feedback-driven optimization more difficult than in high-frequency retail categories, effectively imposing a stochastic ceiling on achievable performance—a point we quantify in Section 4.

Satisfying campaign-level Return-on-Spend (RoS) target is commonly formalized in the autobidding literature as a ROAS constraint [10]. Prior RoS theory shows that value-maximizing auto-bidders can be studied through constrained online optimization and primal–dual bid updates. However, these theoretical models typically rely on assumptions that differ from production sponsored-product CPC auctions. In particular, clean value-based bid rules often depend on truthful or critical-price-style auction assumptions, under which auction response terms can be absorbed into a common multiplier. In production, the mapping from predicted value to realized ROAS is less direct. First, conversion models are subject to bias and drift, meaning predicted values can systematically deviate from true performance over time, leading to miscalibrated bids. Second, GSP-like pricing in production is not generally truthful because CPC is determined by competitors’ bids and score-based rank thresholds, changing a bid can affect both allocation and payment. Therefore, the relationship between submitted bid, realized

CPC, and ROAS is indirect. Third, marketplace dynamics—such as shifts in demand and supply, as well as changes in auction composition—are inherently non-stationary, causing the relationship between bids, costs, and outcomes to evolve over time. Fourth, individual purchase decisions are inherently stochastic, especially in low-frequency, high-consideration home-furnishing categories, introducing additional noise into observed outcomes. Together, these factors decouple predicted value from realized performance, making a simple value-based bidding strategy unreliable without additional controls or adjustments.

Contributions. This paper presents and evaluates a production tROAS autobidding system for a large-scale marketplace characterized by sparse conversion and high-average-order-value product categories, showing how a campaign-level PID controller can be made practical through auction-time value bidding, data-driven eligibility rules, and stability mechanisms for noisy conversion feedback. Our main contributions are as follows.

Production tROAS in a sparse marketplace. We present a fully integrated tROAS autobidding system for multi-slot CPC sponsored-product auctions in a large home-furnishings marketplace. A key design requirement is *bid-source neutrality*: tROAS-generated CPC bids and manual CPC bids enter the same score-based ranking and GSP-like critical-pricing mechanism, so the controller adjusts only submitted bids rather than post-hoc prices. Complementing theory of RoS-constrained autobidding and auction design [1, 10], our design focuses on the production bid/rank/charge/control interface and handles practical deployment challenges.

Feedback control with production guardrails. Our control strategy uses a single campaign multiplier applied to an auction-time value-based bid. The multiplier is updated from delayed and noisy realized ROAS signals, while adaptive learning rates, bounded updates, and cold-start handling are used to improve stability in production. We further analyze low-volume campaigns with sparse clicks and conversions and quantify a *noise floor* for achievable ROAS. This analysis provides an empirical ceiling for interpreting online attainment results and motivates the guardrails needed for sparse, high-AOV advertising environments.

Empirical validation. We support our design with large-scale data. Offline analysis of 32,000+ campaigns demonstrates how sparse signals limit ROAS accuracy, informing our thresholds. In post-deployment evaluation, within the cohort of enrolled campaigns with 50+ monthly orders, the system achieved a 74% attainment rate (ROAS within $\pm 20\%$ of target), coming close to the offline-estimated intrinsic achievability for comparable cohorts. An observational efficiency study comparing tROAS to the manual bidding baseline shows that tROAS delivers 11-15% higher efficiency, and is associated with 17% faster WSP spend growth among high-adopting suppliers, though this analysis is not causal due to advertiser self-selection. Together, these results advance the production practice of constrained autobidding, bridging the gap between prior research [1–3, 6, 8, 10, 17] and a production system for a challenging marketplace.

2 Related Work

Research on automated bidding spans constrained optimization, auction-theoretic analysis, production control systems for delivery and efficiency, and prediction models for conversion value. A

central autobidding formulation is to maximize advertiser value subject to efficiency constraints such as return-on-spend (RoS) or budget limits. Aggarwal *et al.* [1] survey the algorithmic and economic foundations of autobidding, including bidding strategies, equilibrium behavior, and auction design. Feng *et al.* [10] formalize RoS-constrained bidding as online constrained optimization with primal–dual updates and regret guarantees, while recent work extends this setting to uncertain ML-estimated values [11]. These works motivate treating tROAS as a constraint-enforcement problem rather than heuristic bid tuning. A recurring strategy in theory and practice is to scale per-auction value estimates by a campaign-level multiplier. Deng *et al.* [8] analyze RoS-constrained autobidders under value $\times$ multiplier structure and study the implications for auction design. Aggarwal *et al.* [2] connect affine constraints to truthful-auction baselines, further motivating multiplier-based bidding as a principled primitive. Our design follows this structure by computing auction-time value and applying a campaign multiplier to correct residual bias and drift.

Sponsored-search auctions are commonly modeled as multi-slot, GSP-like position auctions. Edelman *et al.* [9] and Varian [16] analyze rank-based allocation and second-price-style payments in such settings. This is directly relevant to our system because candidates are ranked by a score, not raw bid alone, and the charged CPC is the minimum bid needed to maintain the score-based position. Once ranking incorporates quality, relevance, or profit terms [14], the realized CPC is governed by the score threshold needed to beat the next competitor rather than the next highest dollar bid. A practical complication is that submitted bids and realized costs systematically differ: while many display markets have moved toward first-price auctions [15, 19], sponsored products often retain GSP-like pricing, creating a bid–price gap that a pure value-based bid may not correct on its own.

Production ad systems often use multiplicative parameters updated over time to stabilize delivery or enforce constraints. Balseiro *et al.* [7] analyze budget-management mechanisms and stability, while Balseiro *et al.* [6] compare architectures for coordinating budget pacing and RoS pacing, including decoupled, minimally coupled, and fully coupled multiplier-based designs. Apparaju *et al.* [4] frame small-budget pacing as a feedback-control problem and show how ad-hoc tuning can be unstable under noisy signals, motivating smoothing and reduced parameter volatility. Reinforcement learning has also been explored for ROI/ROAS-type constraints [18], but production systems often prioritize debuggability, stability, and guardrails, motivating simpler feedback controllers.

3 System Design

We consider a sponsored-product CPC auction platform. Each campaign i contains a set of SKUs/items. For each auction context a (query/keyword, placement, user state), the system computes bids for eligible SKUs and ranks ads according to a rank-by-score function. Figure 1 illustrates the end-to-end data flow. The system has two coupled loops. The *online serving loop* operates at auction time: a user visit triggers real-time prediction, bids are computed for eligible items, and the auction determines delivery and charged CPC. The *feedback-control loop* operates at the campaign level: realized outcomes (clicks, CPC, and attributed revenue) are aggregated over a rolling window and used to update a campaign multiplier CP_i .

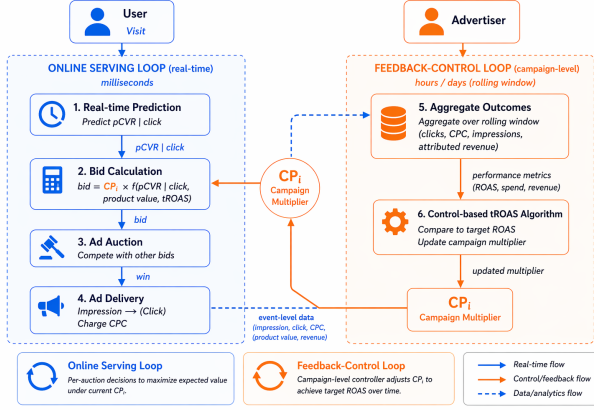


Figure 1: Architecture of Ads Auction and tROAS Bidding

that scales future bids. Separating fast per-auction decisions from slower campaign-level corrections is critical in sparse-conversion environments where ROAS observations are delayed and noisy.

Ranking and Pricing. Each auction ranks eligible campaign-SKU ad candidates from both tROAS and manual CPC campaigns in the same score-based ranking mechanism; the only difference is whether the submitted CPC bid is generated by tROAS or specified manually. The ranking score combines bid, quality, and marketplace-value signals, e.g., $\text{score}_j = f(\text{bid}_j \cdot p\text{Click}_{j,a} \cdot \hat{p}_{j|a} \cdot \text{profit}_j)$, following profit-aware ranking frameworks [14]. The charged CPC is then determined by GSP pricing: holding non-bid ranking terms fixed, the winner pays the minimum CPC required to maintain its position.

3.1 Value-Based Bid Rule

We index campaigns by i , SKUs by j , and auction contexts by a . Let $\hat{p}_{j|a}$ denote the predicted probability of conversion given a click, and let π_j denote the bid-time estimate of supplier value from a sale. The base expected value per click is

$$v_{i,j,a} = \pi_j \cdot \hat{p}_{j|a}. \quad (1)$$

For campaign i with target ROAS T_i , realized campaign-level ROAS over a measurement window W is computed from attributed revenue and spend as $\text{ROAS}_i(W) = \frac{\sum_{t \in W} \text{Revenue}_{i,t}}{\sum_{t \in W} \text{Spend}_{i,t}}$. This realized metric is used by the feedback controller after clicks, orders, and attributed revenue are observed. At bid time, however, the bidder reasons about expected value and expected cost.

A standard RoS-constrained autobidding formulation is to maximize expected conversion value subject to an expected ROAS constraint [1, 10]. Let $x_{i,j,a}(b)$ denote whether campaign-SKU pair (i, j) receives a click in auction context a under submitted bids b , and let $c_{i,j,a}(b)$ denote the corresponding expected CPC. The expected campaign-level ROAS can be written as

$$\text{ROAS}_i(b) = \frac{\sum_{a,j} x_{i,j,a}(b) v_{i,j,a}}{\sum_{a,j} x_{i,j,a}(b) c_{i,j,a}(b)}. \quad (2)$$

The objective is then to maximize expected value subject to the expected ROAS matches the campaign’s target.

$$\max_b \sum_{a,j} x_{i,j,a}(b) v_{i,j,a} \quad \text{s.t.} \quad \frac{\sum_{a,j} x_{i,j,a}(b) v_{i,j,a}}{\sum_{a,j} x_{i,j,a}(b) c_{i,j,a}(b)} = T_i. \quad (3)$$

Prior work shows that this constrained optimization can be studied through Lagrangian or primal-dual updates, and under truthful or critical-price-style auction assumptions the resulting bid takes a value-times-multiplier form [1–3, 6, 8, 10, 17]. In production, we use the same structure but learn the multiplier through campaign-level feedback. The auction-time bid for SKU j in campaign i at auction context a is therefore

$$b_{i,j,a} = \text{CP}_i \cdot \frac{v_{i,j,a}}{T_i} = \text{CP}_i \cdot \frac{\pi_j \cdot \hat{p}_{j|a}}{T_i}, \quad (4)$$

where CP_i is the campaign-level control parameter described in Section 3.3. Intuitively, $v_{i,j,a}/T_i$ converts expected value per click into an allowable CPC under the target ROAS constraint. The multiplier CP_i then adjusts this base bid to account for residual mismatch between expected and realized performance, including prediction error, auction pricing effects, attribution delay, and marketplace non-stationarity.

3.2 Sources of bid-to-ROAS mismatch.

In production, several factors cause auction-time value estimates to differ from realized campaign-level ROAS. First, like any learned predictor used in production, the post-click conversion-rate model is imperfect and, therefore, contributes estimation error to the system. Second, the post-click conversion-rate model is learned from historical attribution labels and is therefore subject to attribution-window mismatch. In particular, the attribution definition used to train pCVR may differ from the window used in the campaign-level ROAS feedback loop and from the measurement suppliers observe in reporting. Third, the conversion value available at bidding time is only an estimate of the sale value for suppliers. The realized order value can differ from the bid-time value proxy because the final purchased SKU option, quantity, or transaction value is only observed after purchase. Fourth, in GSP-like score-based auctions, the submitted bid, realized clearing CPC, and win probability are connected indirectly through competition and ranking mechanics. The resulting bid-to-CPC relationship can vary across auctions and over time. Finally, purchase outcomes are inherently stochastic, especially in low-frequency, high-consideration home-furnishing categories, which adds substantial noise to observed ROAS.

3.3 Feedback Controller

The control parameter CP_i serves as a multiplicative feedback loop that adjusts bids to drive the observed ROAS toward the target. The controller is inspired by PID (Proportional-Integral-Derivative) control, the most widely adopted feedback-control paradigm and one that delivers robust performance without a detailed model of the underlying process [5].

Update rule. Let $\overline{\text{ROAS}}_i$ denote the smoothed observed ROAS for the campaign i calculated over a rolling window of N -days. The normalized error is as follows:

$$e_i^{(t)} = \frac{\overline{\text{ROAS}}_i - T_i}{T_i} \quad (5)$$

The control parameter is updated multiplicatively:

$$\text{CP}_i^{(t)} = \text{CP}_i^{(t-1)} \cdot (1 + \alpha_i^{(t)} \cdot e_i^{(t)} + k_d \cdot (e_i^{(t)} - e_i^{(t-1)})) \quad (6)$$

where $\alpha_i^{(t)}$ is an adaptive learning rate (defined below) and k_d is the derivative gain. Intuitively, if the observed ROAS exceeds the target ($e_i^{(t)} > 0$), bids should increase to spend more and bring the ROAS down; if the observed ROAS is below the target, bids should decrease to spend less and bring the ROAS up. The multiplicative structure means that cumulative updates embed an integral effect:

$$\log CP_i^{(t)} \approx \log CP_i^{(0)} + \sum_{\tau=1}^t \alpha_i^{(\tau)} e_i^{(\tau)} + k_d (e_i^{(t)} - e_i^{(0)}) \quad (7)$$

so past errors accumulate in log-space while the k_d term dampens oscillation. To avoid abrupt oscillations, the per-update relative change is capped via $\Delta \in (0, 1)$:

$$CP_i^{(t)} \in [CP_i^{(t-1)}(1 - \Delta), CP_i^{(t-1)}(1 + \Delta)]. \quad (8)$$

CP updates are also subject to floor and ceiling guardrails to avoid extreme overbidding or underbidding: $CP_i \in [CP_{\min}, CP_{\max}]$.

Adaptive learning rate. The learning rate $\alpha_i^{(t)}$ decays over time from the start of the tROAS campaign (or the last reset event) to balance responsiveness against oscillation — large early to correct the cold-start bias, small later for stability.

$$\alpha_i^{(t)} = \max(\alpha_{\min}, \alpha_0 \cdot e^{-\lambda(t-t_{\text{start}})}) \quad (9)$$

where t_{start} is reset on trigger events such as a major change in the campaign's tROAS target, budget, or SKU composition. The decay reaches $\alpha_{\min}$ after approximately 14 days, stabilizing the controller once sufficient data accumulate.

Cold-start handling. The cold-start problem is particularly important in Wayfair's low-order, high-order value environment, where sparse conversions lead to poor and noisy feedback. For new campaigns with limited history, even a single order can substantially move the ROAS error signal. We address this with three mechanisms: (a) *CP freeze*: the control parameter is held at its initial value until the campaign records its first attributed order, preventing a downward spiral from zero-ROAS readings; (b) *historical warm-start*: for campaigns duplicated from an existing campaign, the historical ROAS of the parent is seeded into the rolling window; (c) *for new non-duplicated campaigns without history*, smoothed ROAS over a rolling window W with a Laplace-smoothing constant factor κ_C to dampen extreme ratios when spend or revenue is near zero or low:

$$\overline{\text{ROAS}}_i = \frac{\sum_{t \in W} \text{Revenue}_{i,t} + \kappa_C}{\sum_{t \in W} \text{Spend}_{i,t} + \kappa_C} \quad (10)$$

Pseudocode. Algorithm 1 outlines the iterative update process used in the tROAS feedback controller. The adjustable parameters it references are configured based on a combination of business guardrails, historical bid and CP distributions, attribution-delay profiles, and validated using offline simulation. The disclosure of absolute parameter used in our production system is restricted in accordance with corporate policy.

4 Offline ROAS Noise-Floor Benchmark

Before deployment, we quantify the finite-window ROAS variation that remains even when bids are correct in expectation. In a CPC auction, bidders buy clicks, while ROAS is realized only through

Algorithm 1: Constrained tROAS Bidding with Campaign-Level Feedback Control

Input: Target ROAS T_i ; campaign multiplier CP_i ; eligible SKUs j_a ; predicted purchase probability $\hat{p}_{j|a}$; value proxy π_j ; parameters $\alpha_i^{(t)}$, k_d , λ , Δ , ϵ , $CP_{\min}$, $CP_{\max}$, κ_C ,

Output: Bids $\{b_{i,j,a}\}$ and updated multiplier CP_i

Online Serving (per auction a):

foreach tROAS candidate $(i, j) \in C_a^{\text{tROAS}}$ **do**

$v_{j,a} \leftarrow \pi_{i,j} \cdot \hat{p}_{j|a}$;
 $b_{i,j,a} \leftarrow CP_i \cdot v_{j,a}/T_i$;

Apply bid floor and bid cap guardrails;

Merge tROAS-generated bids with manual CPC bids from C_a^{manual} ;

Submit all candidates to the shared score-based auction;

Periodic Control Update (over window W):

$V \leftarrow \sum_{(i,t) \in W} \text{Revenue}_{i,t}$;

$C \leftarrow \sum_{(i,t) \in W} \text{Spend}_{i,t}$;

$\bar{R} \leftarrow (V + \kappa_C)/(C + \kappa_C)$;

$e_i^{(t)} \leftarrow \bar{R}/T_i - 1$;

if $|e_i^{(t)}| < \epsilon$ **then**

return CP_i [within freeze zone; no update];

$CP' \leftarrow CP_i \cdot (1 + \alpha_i^{(t)} \cdot e_i^{(t)} + k_d \cdot (e_i^{(t)} - e_i^{(t-1)}))$;

$CP' \leftarrow \min(CP', CP_i(1 + \Delta))$;

$CP' \leftarrow \max(CP', CP_i(1 - \Delta))$;

$CP_i \leftarrow \text{clip}(CP', CP_{\min}, CP_{\max})$;

return $\{bid_{i,j,a}\}, CP_i$;

stochastic post-click purchases and order values. This benchmark sets a realistic reference for post-launch tROAS achievement by separating controller error from irreducible purchase randomness.

Split-sample AA estimator. We estimate this noise floor using $N = 32,140$ campaigns with positive spend during a two-week window. For each campaign, we randomly partition customers into two equally sized groups, A and A' . Let $R_{ig} = V_{ig}/S_{ig}$ denote split-level ROAS for $g \in \{A, A'\}$, where V is attributed order value and S is CPC spend, and let $R_i = (V_{iA} + V_{iA'})/(S_{iA} + S_{iA'})$ denote full-window campaign ROAS.

Because both splits come from the same campaign, auction, bidding policy, and measurement window, differences between R_{iA} and $R_{iA'}$ reflect finite-sample variation in post-click purchases and order values rather than bidding or controller behavior. We define the rescaled AA deviation $D_i = |R_{iA} - R_{iA'}|/(2R_i)$ and count campaign i as measurable within a $\pm 20\%$ target band when $D_i \leq 0.20$. The factor of 2 maps the difference between two independent half-sample estimates to the error scale of the full-window estimate; this relationship is exact for means and a standard delta-method approximation for ROAS ratios. Campaigns with zero attributed orders are reported separately and excluded from the target-band calculation. We bucket campaigns by full-window attributed orders divided by two.

Results. Table 1 reports target-band achievability by weekly order volume. The results show a steep sparsity gradient: campaigns with 51+ weekly orders fall within the rescaled $\pm 20\%$ band 82.4% of

Table 1: ROAS target achievability by weekly order band (best possible based on AA test).

Orders/Week	% Within $\pm 20\%$
0	—
1–10	15.9
11–30	51.8
31–50	65.2
51+	82.4

the time, while campaigns with 1–10 weekly orders do so only 15.9% of the time. This sparsity is material in our market. For example, a campaign with 500 clicks and a 2% post-click conversion rate is expected to generate only 10 orders per week; even with correctly calibrated bids, random variation in order counts and order values can produce large ROAS swings. The challenge is even more pronounced for high-consideration, higher-priced home-furnishing categories with lower purchase frequency.

This noise floor has two operational implications. First, it informs the *eligibility threshold*: campaigns need sufficient order volume before target ROAS can be evaluated reliably. Second, it motivates the controller-stability mechanisms in Section 3.3; when ROAS feedback is noisy, aggressive learning rates and cold-start updates can amplify random fluctuations rather than improve convergence.

5 Online Experiments

5.1 Rollout & Burst Test Results

The system was deployed via a phased burst test, gradually scaling from 20 campaigns to 590 over five months before a broad launch for all advertisers. All experiments use CPC billing with GSP-like pricing on our large-scale sponsored-product platform.

Burst Test Evaluation. We report results from the phased burst test covering 590 campaigns with 50+ monthly orders and were assessed after a 14-day learning window. In total, 95% of campaigns achieved at least 80% of their tROAS goal (meet-or-beat), and 64% performed within $\pm 20\%$ of the target (attainment). Since above-target ROAS can indicate conservative bidding and under-delivery, we report ROAS attainment together with budget utilization. As shown in Figure 2(a), 60% of campaigns were both within target and high-utilization, while 12% were above target but low-utilization, suggesting unmet spend opportunity.

Post-Deployment Evaluation. Within 90 days of broad deployment, approximately 20% of all Wayfair campaigns were enrolled in tROAS. Because enrollment is advertiser-driven, this post-deployment analysis is descriptive rather than causal: it measures steady-state target attainment among enrolled campaigns, not lift relative to a randomized holdout.

To compare online performance against the offline achievability baseline in Table 1, we focus on the cohort of enrolled campaigns having 50+ monthly orders, ensuring both reliable ROAS measurement and a fair comparison. Within this evaluable cohort, the system achieved a **90%** meet-or-beat rate ($\text{ROAS} \geq 80\%$ of target) and a **74%** attainment rate (ROAS within $\pm 20\%$ of target) after a 14-day learning period under steady-state conditions. The attainment rate approaches the intrinsic achievability bound estimated from the offline analysis, suggesting that post-deployment performance is

(a) ROAS attainment by budget utilization

ROAS band	Low util	High util	Total
Under: $< 80\%$ target :	2%	3%	5%
Meet: $80\text{--}120\%$ target	4%	60%	64%
Over: $> 120\%$ target	12%	19%	31%

Low util. = $< 80\%$ budget used;
High util. = $\geq 80\%$ budget used.

(b) Covariate-adjusted efficiency comparison

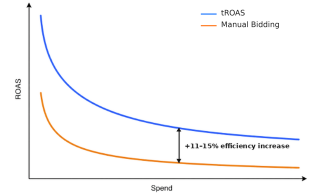


Figure 2: Post-deployment evaluation summary. (a) ROAS attainment by budget utilization after the 14-day burst test ($N = 590$). (b) Covariate-adjusted bidding efficiency comparison of tROAS and manual bidding campaigns

close to the level expected given conversion sparsity and stochastic ROAS variation.

To evaluate whether tROAS delivers greater efficiency than manual bidding, we ran an observational post-deployment comparison using covariate-adjusted quantile regression, since randomized holdout was not feasible due to voluntary supplier adoption. Figure 2(b) shows that, controlling for product category, average-order-value bucket, and campaign week, tROAS achieved 11–15% higher efficiency than manual bidding. A separate descriptive analysis finds that high-tROAS adopters increased ad spend 17% faster and active-SKU participation 8% faster than non-adopters. Although observational rather than causal, these results align with improved efficiency and growing supplier confidence in automated bidding.

5.2 Iterative Improvements via A/B Tests

A/B testing methodology. After initial deployment, controller changes were evaluated with campaign-level randomized A/B tests among enrolled tROAS campaigns. Each campaign was assigned to control or treatment before serving, and all traffic for that campaign inherited the same assignment. This preserves auction-time consistency and avoids within-campaign interference. Enrollment into tROAS is advertiser-driven, so the experimental population reflects self-selection; conditional on enrollment, however, treatment assignment is randomized, so the A/B comparison is internally valid for the enrolled set.

In the iterative-improvement test, Control A used the previously deployed controller, while Treatment B added three stability mechanisms from Section 3.3: smoothed ROAS error, a decaying update schedule, and caps on per-update CP changes. These changes were designed to reduce overshooting under sparse, noisy, and delayed conversion feedback. The analysis removed the top 1% of campaigns by volume to reduce sensitivity to outliers, yielding comparable group sizes ($n(A) = 2,320$, $n(B) = 2,408$). Randomization quality was validated with bootstrap-based balance checks and sample-ratio-mismatch diagnostics; ratio metrics such as ROAS were analyzed using delta-method approximations [12, 13].

Results. Table 2 summarizes the main results. Treatment B significantly improved lower-bound attainment, reduced CPC, and cut CP volatility nearly in half. Spend, exposure, CTR, CVR, ROAS, and order revenue per campaign showed no statistically significant changes, although ROAS (+5.79%) and order revenue per campaign

Table 2: A/B test of controller improvements. Relative lifts are Treatment B vs. Control; bold rows indicate $p < 0.05$.

Metric	Lift	p -value
Lower-bound attainment (ROAS $\geq 0.8T$)	+9.32%	0.0021
ROAS	+5.79%	0.1240
CPC	-9.28%	0.0080
CP volatility (std.)	-49.74%	<0.0001
Spend / campaign	+0.81%	0.8992
Order revenue / campaign	+6.65%	0.3709
Ads exposure / campaign	+4.35%	0.5626
CVR	+0.33%	0.9410
CTR	+6.49%	0.1440

(+6.65%) moved directionally higher. Thus, the primary measurable gains were improved target attainment and control stability, without evidence of reduced delivery or supplier order value.

Diagnostic analysis confirmed the mechanism: smoother and bounded CP updates reduced transient bid spikes caused by short-term noisy ROAS feedback, especially in low-data phases. Lower bid spikes reduced realized CPC while preserving aggregate delivery, consistent with the observed stability and cost improvements.

6 Conclusion and Future Work

We presented a constrained target ROAS bidder for CPC sponsored products that separates the system into (i) an auction-time value-based bid proportional to expected conversion value per click divided by target ROAS, and (ii) a campaign-level multiplicative feedback controller that corrects systematic bias from prediction error, auction pricing discounts, and non-stationarity. This design aligns with the uniform bid-scaling structure in autobidding theory[8] and with stability-aware control approaches for pacing[4, 7].

An offline analysis of 32,000+ campaigns established that conversion sparsity alone limits ROAS achievability to ~16% for low-volume campaigns and ~82% for high-volume ones—setting a realistic performance ceiling for any autobidding system in low-frequency purchase categories. Online experiments demonstrated that, within the evaluable cohort of enrolled campaigns with sufficient order volume, the system achieved a 90% meet-or-beat rate (ROAS $\geq 80\%$ of target) and a 74% attainment rate (ROAS within $\pm 20\%$ of target) after a 14-day learning period under steady-state conditions. The latter approaches the intrinsic achievability bound estimated from the offline analysis, suggesting that observed target attainment is close to the level expected given conversion sparsity and stochastic ROAS variation. The comparison against manual bidding provides descriptive evidence that tROAS improves efficiency, while the randomized controller A/B supports the narrower causal claim that stability mechanisms improve lower-bound attainment, reduce CPC, and reduce multiplier volatility.

Future directions include: 1) **Exploration of under-observed SKUs.** New and low-volume SKUs suffer from a cold-start feedback loop: uncertain $\hat{p}_{j|a}$ leads to low bids, limited impressions, and no data to improve predictions. Introducing a multi-armed bandit layer with Upper Confidence Bound (UCB) adjustments on $\hat{p}_{j|a}$ could break this loop by trading short-term ROAS precision for long-term SKU-level data acquisition. 2) **Seasonality and promotion forecasting.** Major promotions produce ROAS spikes of up to 40%, triggering large CP oscillations. Replacing reactive, manual promo multipliers with a proactive, class-level CVR forecasting model would smooth controller behavior during demand shocks. 3) **Advanced constrained learning.** Primal–dual RoS-constrained

online learning [10] and RL variants [18] offer principled alternatives to PID-style control, potentially improving convergence under non-stationarity while balancing production requirements for simplicity and debuggability.

References

- [1] Gagan Aggarwal, Ashwinkumar Badanidiyuru, Santiago R. Balseiro, Kshipra Bhawalkar, Yuan Deng, Zhe Feng, Gagan Goel, Christopher Liaw, Haihao Lu, Mohammad Mahdian, Jieming Mao, Aranyak Mehta, Vahab Mirrokni, Renato Paes Leme, Andrés Perlroth, Georgios Piliouras, Jon Schneider, Ariel Schwartzman, Balasubramanian Sivan, Kelly Spendlove, Yifeng Teng, Di Wang, Hanrui Zhang, Mingfei Zhao, Wennan Zhu, and Song Zuo. 2024. Auto-bidding and Auctions in Online Advertising: A Survey. *ACM SIGecom Exchanges* 22, 1 (2024), 159–183. arXiv:2408.07685
- [2] Gagan Aggarwal, Ashwinkumar Badanidiyuru, and Aranyak Mehta. 2019. Auto-bidding with Constraints. In *Web and Internet Economics – WINE 2019 (Lecture Notes in Computer Science, Vol. 11920)*. Springer, 17–30. doi:10.1007/978-3-030-35389-6_2
- [3] Gagan Aggarwal, Giannis Fikioris, and Mingfei Zhao. 2024. No-Regret Algorithms in Non-Truthful Auctions with Budget and ROI Constraints. *arXiv preprint arXiv:2404.09832* (2024). arXiv:2404.09832
- [4] Sreeja Apparaju, Yichuan Niu, and Xixi Qi. 2025. Feedback Control for Small Budget Pacing. *arXiv preprint arXiv:2509.25429* (2025). arXiv:2509.25429
- [5] Karl Johan Åström and Tore Hägglund. 1995. *PID Controllers: Theory, Design, and Tuning* (2 ed.). Instrument Society of America (ISA), Research Triangle Park, NC.
- [6] Santiago R. Balseiro, Kshipra Bhawalkar, Zhe Feng, Haihao Lu, Vahab Mirrokni, Balasubramanian Sivan, and Di Wang. 2024. A Field Guide for Pacing Budget and ROS Constraints. In *Proceedings of the 41st International Conference on Machine Learning (Proceedings of Machine Learning Research, Vol. 235)*. PMLR, 2607–2638.
- [7] Santiago R. Balseiro, Anthony Kim, Mohammad Mahdian, and Vahab S. Mirrokni. 2021. Budget-Management Strategies in Repeated Auctions. *Operations Research* 69, 3 (2021), 859–876.
- [8] Yuan Deng, Jieming Mao, Vahab S. Mirrokni, and Song Zuo. 2021. Towards Efficient Auctions in an Auto-bidding World. In *Proceedings of The Web Conference 2021*. 3965–3973. arXiv:2103.13356
- [9] Benjamin Edelman, Michael Ostrovsky, and Michael Schwarz. 2007. Internet Advertising and the Generalized Second-Price Auction: Selling Billions of Dollars Worth of Keywords. *American Economic Review* 97, 1 (2007), 242–259.
- [10] Zhe Feng, Swati Padmanabhan, and Di Wang. 2023. Online Bidding Algorithms for Return-on-Spend Constrained Advertisers. In *Proceedings of the ACM Web Conference 2023*. ACM, 3550–3560. doi:10.1145/3543507.3583491
- [11] Jiale Han, Chun Gan, Chengcheng Zhang, Jie He, Zhangang Lin, Ching Law, and Xiaowu Dai. 2026. Auto-Bidding under Return-on-Spend Constraints with Uncertainty Quantification. In *Proceedings of the ACM Web Conference 2026*. ACM, arXiv:2509.16324 doi:10.1145/3774904.3792109
- [12] Ron Kohavi, Roger Longbotham, Dan Sommerfield, and Randal M. Henne. 2009. Controlled Experiments on the Web: Survey and Practical Guide. *Data Mining and Knowledge Discovery* 18, 1 (2009), 140–181.
- [13] Ron Kohavi, Diane Tang, and Ya Xu. 2020. *Trustworthy Online Controlled Experiments: A Practical Guide to A/B Testing*. Cambridge University Press.
- [14] Manavender Malgireddy, Dongyue Xie, Sergey Kolbin, Kurt Zimmer, Masoum Mosmer, and Patrick Phelps. 2025. Profit Aware Ad Ranking with Relevance Constraint. In *Proceedings of the AdKDD 2025 Workshop*. Springer. doi:10.1007/978-3-032-16358-5_6
- [15] Renato Paes Leme, Balasubramanian Sivan, and Yifeng Teng. 2020. Why Competitive Markets Converge to First-Price Auctions. In *Proceedings of The Web Conference 2020*. ACM, 596–605. https://research.google/pubs/why-competitive-markets-converge-to-first-price-auctions/
- [16] Hal R. Varian. 2007. Position Auctions. *International Journal of Industrial Organization* 25, 6 (2007), 1163–1178.
- [17] Sushant Vijayan, Zhe Feng, Swati Padmanabhan, Karthikeyan Shanmugam, Arun Suggala, and Di Wang. 2025. Online Bidding under RoS Constraints without Knowing the Value. *arXiv preprint arXiv:2503.03195* (2025). arXiv:2503.03195
- [18] Haozhe Wang, Chao Du, Panyan Fang, Shuo Yuan, Xuming He, Liang Wang, and Bo Zheng. 2022. ROI-Constrained Bidding via Curriculum-Guided Bayesian Reinforcement Learning. In *Proceedings of the 28th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*. ACM. doi:10.1145/3534678.3539211
- [19] Tian Zhou, Hao He, Shengjun Pan, Niklas Karlsson, Bharatbhushan Shetty, Brendan Kitts, Djordje Gligorijevic, San Gultekin, Tingyu Mao, Junwei Pan, Jianlong Zhang, and Aaron Flores. 2021. An Efficient Deep Distribution Network for Bid Shading in First-Price Auctions. In *Proceedings of the 27th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*. ACM, 3996–4004. arXiv:2107.06650 doi:10.1145/3447548.3467167