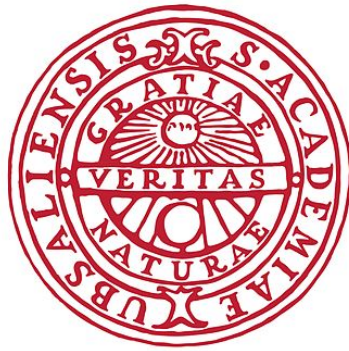




UPPSALA UNIVERSITET

Does market entry end bid-rotation?

Evidence from the Swedish market for generic substitution

Gustaf Holmberg

Master program in economics

Spring 2025

Contents

1	Introduction	3
2	Background, theory and previous literature	4
2.1	Institutional background	4
2.2	Previous literature	8
2.3	Game theory model of collusion	10
2.4	Price variation under collusion	13
3	Data	14
4	Empirical strategy	17
4.1	Overlapping permutations test	19
4.1.1	General idea	19
4.1.2	The statistical test	22
4.1.3	Parametrization and robustness checks	28
4.1.4	Choice of p-value and output of the test	29
4.2	Pre/post analysis	29
4.2.1	Fixed-effects panel regression	30
4.2.2	Event-study design	31
5	Results	33
5.1	Overlapping permutations test	34
5.2	Fixed-effects panel regression	35
5.3	Event study	39
5.4	Robustness checks	44
6	Discussion	46
7	Conclusion	50
A	Data	54
B	Empirical strategy	54
C	Robustness check – overlapping permutations test	55
D	Event-study	59
D.1	Coefficient plots	59
D.2	Tables	67

Abstract

This thesis asks a simple question: Does the entry or exit of a firm stop bid-rotation and make prices fall in Swedens market for generic drugs? To answer this question, I study fifteen years of monthly auction data from the national substitution system. I build an overlapping-permutations test that scans each sub-market for repeating winner patterns that signal bid-rotation. Next, I link these flagged periods to fixed-effects regressions and an event-study that compare prices and price spread before and after a firm enters or leaves. The test shows that bid-rotation appears almost only when two or three firms are active. When a new firm joins, the collusive pattern breaks roughly half the time and rarely returns. After a break, average prices fall by about 20–35 percent within two months, and the standard deviation of bids roughly doubles. The results are amplified when I only consider breaks in bid-rotation the coincides with market entry. These findings confirm the idea that tacit collusion is easier to keep when few firms compete and that entry tilts the balance back to price cutting. The study suggests that lowering entry barriers for generic suppliers is an effective way to protect consumers from hidden collusion.

Keywords— Bid-rotation, price competition, Bertrand paradox, repeated games

1 Introduction

Why do prices sometimes remain stubbornly high in markets that, on paper, are engineered for ferocious price wars? Few settings illustrate the puzzle more starkly than Swedens generic-drug product-of-the-month auctions – a system that was explicitly built on the textbook intuition of classical theory on price competition (Bertrand 1883). In such a setting, theory predicts that prices should be pushed to marginal cost. Yet, there are many examples of similar systems where prices have drifted far above that, echoing the Bertrand paradox that has puzzled economists since the late nineteenth century (Cuddy 2020).

A leading explanation to the paradox is found in the theory on repeated Games (Ivaldi et al. 2003) which provides an understanding of how price coordination can be preferred over price competition among firms acting in markets like the Swedish generics market (Aumann and Shapley 1994). A growing body of work points to bid-rotation as one important form of such price coordination (Athey, Bagwell, and Sanchirico 2004; Harrington Jr and Chen 2006; Skrzypacz and Hopenhayn 2004; Aoyagi 2003). In such a scheme, firms coordinate their price setting so that they take turns tendering the "cheapest" product while the others abstain from aggressive price competition. Yet, the knowledge on how fragile these rotations are when the market structure changes is still limited. This paper therefore asks:

How does market entry and exit affect price competition in markets where bid-rotation is present?

Answering this question demands both detection and measurement. I first construct an overlapping-permutations test that scans the monthly sequence for cyclically repeating pattern of winners. Suspicious patterns are flagged for bid-rotation. Inference on what patterns that should be flagged is made by a combination of a "fuzzy" scoring algorithm and the Stationary block bootstrap re-sample method (Politis and Romano 1994). Flagged sequences provide the foundation for a quasi-experimental event-study design: whenever a new firm enters or an incumbent leaves a flagged segment, that calendar month becomes a treatment time. With the event-study I trace the evolution of prices and price dispersion several months backward and forward.

Applied to fifteen years of Swedish data the combined procedure reveals three robust patterns. First, bid-rotation thrives almost exclusively when only two or three firms are present. Second, the entry of a single additional bidder breaks the rotation roughly half the time which rarely ever re-emerges. Third, when rotations collapse, average prices fall by 22-35 percent within two months, while price dispersion roughly doubles – a signature of renewed, aggressive price-cutting. The drop in prices and rise in price dispersion indicates that when bid-rotation ends competition ensues. These results are larger when only breaking of the bid-rotation by entry of a firm is considered. The results tell us that firm entry play an important role in the dissolution of tacit bid-rotation schemes and that it intensifies competition after such a scheme has ended.

The answer to the research question matters for two reasons. First, the generic market ensues when patent protection of originals runs out. As argued by Nordhaus 1969, the optimal patent policy balances the marginal dynamic benefit with the marginal static efficiency loss.¹ To assess this trade-

¹Berndt 2002 calls this conflict of perspectives "a deep and enduring one" in the market for pharmaceuticals and highlights that the severity of the issue has increased with the recent sharp rise in R&D costs within the industry.

off, one need to know how the market behaves under both patent protection and competition. I hope that my thesis will shed some light on the latter.

Second, understanding the answer is crucial for theory on competition. A central insight from repeated-game theory is that the stability of tacit collusion hinges on the number of active firms. Any empirical strategy that hopes to test that theory, therefore, must keep the firm count front and center. Linking the overlapping-permutations test to an event-study design does exactly that. The permutations screen first isolates stretches of data that look like rotation within sub-segments defined by a fixed set of firms, ensuring that a two-firm episode is never conflated with a three-firm one etc. Those precisely dated flags then feed straight into the event study: when a new entrant appears – or an incumbent disappears – the flagged segment supplies a clean before and after window in which the only structural change is the shift in number of firms. This chain of detection-then-measurement turns a qualitative theoretical prediction into a quantifiable estimate.

The general relationship between the number of firms and the price level has been studied by (Granlund and Bergman 2018). There is however (to my knowledge), no study on the direct effect of firm entry or exit in these markets as predicted by the game theory model on repeated games. This is the research gap my thesis tries to fill.

The empirical regularities found in the analysis are what repeated-game theory would predict. With more competitors, the temptation to undercut today outweighs the threat of future punishment, making price coordination hard to sustain. Entry thus chastens a cooperative equilibrium back into a competitive one. From the theory presented in section 2 each of these predictions are motivated. Stated clearly they are:

P1: Collusive behavior will be found more frequently in markets with few firms.

P2: The breakdown of collusive behavior will decrease the price in the market.

P3: The breakdown of collusive behavior will increase the variation in prices in the market.

These predictions are setup as stepping stones in answering the research question.

The remainder of the thesis proceeds as follows. Section 2 reviews the institutional setting, prior empirical work, and the theoretical foundation for the three predictions. Section 3 describes the data. Section 4 details the empirical strategy, including the new screening algorithm. Section 5 reports the results, Section 6 discusses their implications for policy and theory, and Section 7 concludes.

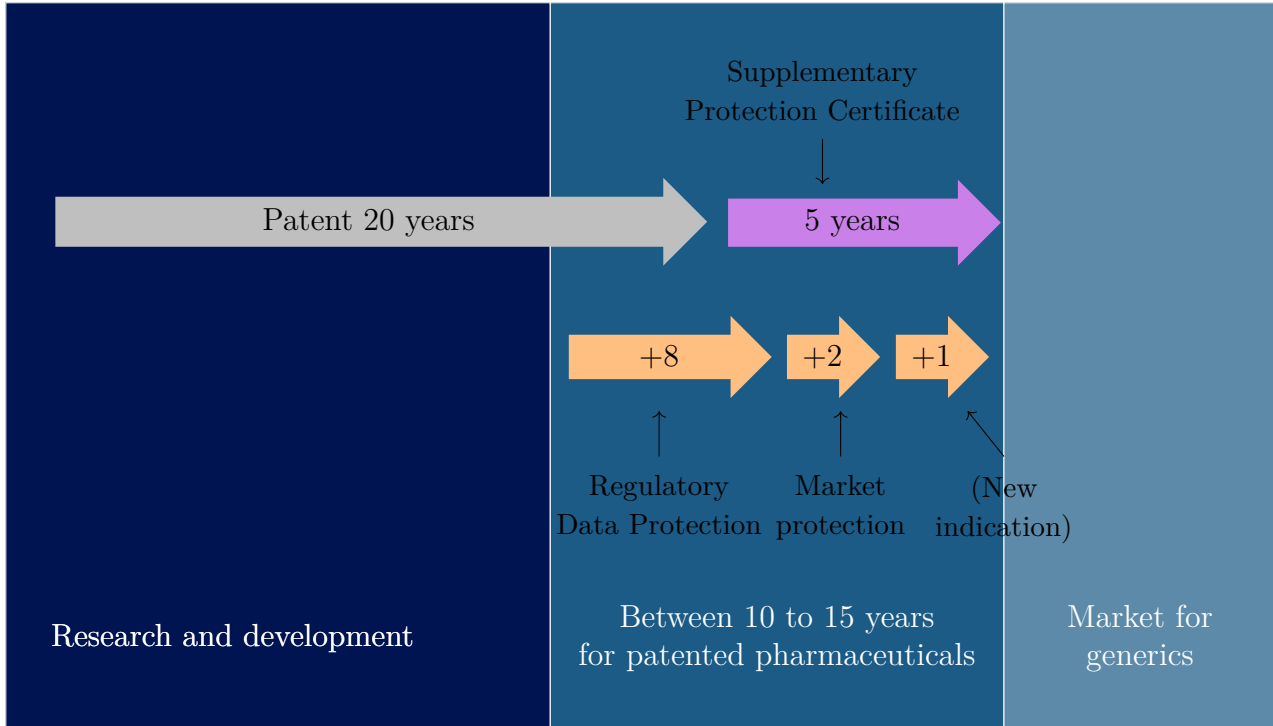
2 Background, theory and previous literature

2.1 Institutional background

The market for pharmaceuticals has two distinct features: high fixed costs and low marginal costs. The effort to develop a new treatment is costly (Wouters, McKee, and Luyten 2020). Once a new drug has been invented, the cost of producing one additional unit is often low (Lakdawalla 2018) and economies of scale are not considered important in the production of already invented drugs (Caves et al. 1991). Without patent protection, competitors can free-ride on the R&D made by other firms

once a new drug formula is developed (Oronsky et al. 2023).² This is a classic case of market failure due to the non-excludability of knowledge (Grossman and Helpman 1993).

In Sweden, EU-regulations give companies the possibility to apply for patents extending 20 years from the date of application.³ Once the patent protections have expired, other firms can enter the market to produce and sell generics.⁴ The situation is depicted in figure 1.



Based on Rondahl 2022

Figure 1: Timing of patents

A national market for generics in Sweden was introduced in 2002, together with a system of generic substitution. The system has the following features. Products on the generics market are narrowly defined within "substitutions groups". A substitution group is defined by a substance $\times$ strength $\times$ form $\times$ package combination⁵ – a setup homogenizes goods to a great extent. As an example of a substitution group we can consider the active substance *ibuprofen*. A product with this substance is called "Ipren" and is commonly sold in Sweden with the strength of 400mg in packages of 30 tablets. Here we see that "Ipren" is one product within the substitution group *ibuprofen* $\times$ 400mg $\times$ tablets $\times$ 30. Thus we can think of a substitution group as a sub-market within a larger market defined solely by the active substance. I will use the words "substitution group" and "sub-market" interchangeably for the rest of the thesis.

²Within the World Trade Organization there is an international legal agreement under the acronym TRIPS (Agreement on Trade-Related Aspects of Intellectual Property Rights) which has been in place since the 1990 (Grossman and Lai 2004) which serves as a system of patent protection for new innovations.

³Once the product comes to market, other legal protections such as Regulatory Data Protection (RDP) and Supplementary Protection Certificate (SPC) can provide the company exclusive right to produce and sell a pharmaceutical for at most 15 years. The most common range that a pharmaceutical is patent protected when it has entered the market is between 10 to 15 years (Rondahl 2022).

⁴A generic is a copy of the original pharmaceutical with the same active ingredient which makes it bio-equivalent to the original.

⁵The decision of what constitutes a suitable substitution group is made by the Swedish medical product agency (Janssen 2022).

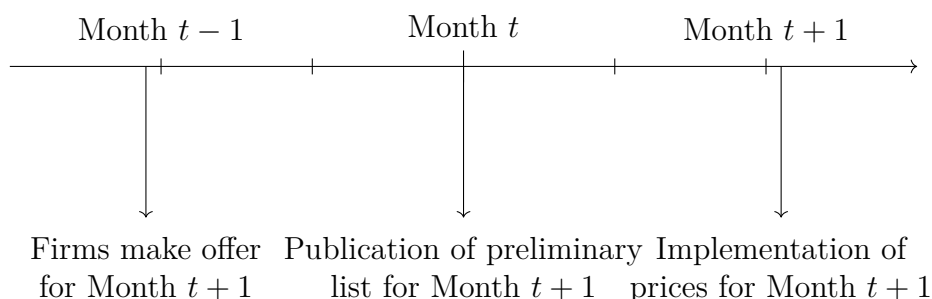
For a company to sell a product within a substitution group, it must partake in a price auction held every month by the Dental and Pharmaceutical Benefits Agency (DPBA). Price bids are uniform across Sweden and include transport to the pharmacies. The lowest price bid wins the auction and gets a special status: "product-of-the-month". This status makes it so that for that month, this particular product is recommended to customers by all pharmacists in Sweden (TLV 2025). All other products are sold at their proposed prices.

Winning the status "product-of-the-month" entails the following things: If an individual has a prescription for a drug and walks into *any* pharmacy in Sweden, the pharmacist who serves that individual is required *by law* to inform the customer of the "product-of-the-month". The information contains two important facts. First, the product of the month is bio-equivalent to the original and other brands.⁶ And second, the product of the month is the cheapest option. This sharing of information nullifies potential search costs for getting the cheapest price.⁷

Firms present their price bids for products within a substitution group to the auction for month t in month $t - 2$. On the fifth workday of month $t - 1$, DPBA announces all purchase prices and the retail pharmacy prices. This makes it so that the bidding in month $t - 2$ happens when prices in month $t - 1$ have already been announced (Granlund and Bergman 2018). With monthly auctions, the history of bids is updated regularly along with the number of active firms in the market. The information on bidding history is central to the strategic behavior of firms (Skrzypacz and Hopenhayn 2004) and the number of active firms determine the incentives for firms to collude (as shown in section 2.3). The timing of the bidding is depicted in figure 2.

The system is set up so as to incentivize low price bids. To be included in the government-funded benefits scheme (*högkostnadsskyddet*), firms need to have their proposed price bids accepted by the DPBA. Three-fourths of the cost of prescription drugs for Swedish patients between the years 2006 to 2012 were covered by the benefits scheme (Wallström et al. 2012). Any price bid lower than the highest price in the previous month gets accepted automatically to the auction.

There are also structures within the system to limit supply shortages of the drug with the "product-of-the-month-status". Pharmaceutical companies have to confirm to the regulatory agency (DBPA) that they can service the whole market before their proposed price is included in the national reimbursement program (Janssen 2022). Pharmacies are required by law to provide a medication that a consumer demands within 24 hours (Granlund and Bergman 2018).



Based on Janssen 2022

Figure 2: Timing of the generics market

⁶Meaning that it has the exact same effect as any other drug within that substitution group.

⁷In Granlund and Rudholm 2012 the authors find that 83% of patients agreed to substitution when they had an option.

The price bidding system, the legal requirements put on pharmacists and the division of pharmaceutical into substitution groups sets the stage for competition almost solely on prices. With homogenized goods, there is very little to compete on except prices.

Under price competition, described by the Bertrand model (Bertrand 1883), the market price will converge to marginal costs in markets with more than one firm. The reason is that for whatever price one firm sets, the other firm(s) can undercut it and capture the whole market. Since this applies to all firms, everyone has incentives to undercut each other. This is true until prices are equal to marginal costs.⁸

Two features of the system strengthens the mechanism whereby a firm with the lowest price captures the whole market. Since pharmacists have to inform of the "product-of-the-month", it is impossible for a customer to *not* find the cheapest price. In economic terms, we would say that search costs are nullified.⁹ The requirement placed on firms and pharmacies to supply the "product-of-the-moth" limits the capacity constraints – a feature that dampens price competition.¹⁰

In sum, the Swedish market for generic substitution is meticulously constructed so as to foster price competition. The reason is obvious. Classical theory predicts that such markets should have low prices.¹¹ This prediction is sometimes in stark contrast to the real world. As an example, the United States have a system much similar to that of Sweden's generics market, which for the most part has been deemed successful.¹² However, amidst the praise given to the American system, worrying signs started appearing in 2013 with unexpected price increases for hundreds of generic drugs (Berndt, Conti, and Murphy 2018). The discrepancy between predictions from the Bertrand theory on price competition, and empirical observations is in the Economics literature referred to as the Bertrand paradox. Given the institutional setting in Sweden, and the similarities with the American system, the Swedish generics market is an ample "laboratory" to study this paradox.

A common way to explain the Bertrand paradox is that firms are engaged in collusion. Collusion cases in the generic pharmaceuticals market have recently been brought in Great Britain, Italy, and South Korea. In the United States, several manufacturers of generic drugs have been indicted on antitrust charges for what is regarded as the largest corporate cartel in the countrys history (Cuddy 2020). Collusion is prohibited by law, and thus under scrutiny from the competition authorities. Even so, there is still possibilities for firms to coordinate their price setting without formal communication, that is, they can collude *tacitly*.¹³ It is this explanation to the Bertrand paradox that this thesis concerns itself with.

⁸Prices lower than marginal costs would mean that firms make negative profits any time they sell a tier goods.

⁹Search costs can be though of the time spent to find out what the lowest price is. If information on prices are hard to come by, customers must spend resources (time) just to find out what the best price is. With large search costs, the actual price of buying a good is higher than what is displayed on the price tag. Thus, with large search costs present we would expect imperfect price competition.

¹⁰Capacity constraints means that no firm can service the whole market. Under such circumstances, when the cheapest product is out of stock, more demand will be directed towards the next cheapest product. In such a case, the cheapest product won't captures the whole market and price competition is undermined.

¹¹A rapport from the Swedish Competition Authority concludes that the system has been successful in keeping prices low (Habib 2017). On a general note, Starc and Wollmann 2025 writes that "Generic drugs are a competition success story."

¹²Morton and Boller 2017 writes that "The U.S. generic market is one of the most dynamic and cost-effective in the world..." and Cuddy 2020 calls it "one of the rare bargains in the US healthcare system".

¹³Tacit collusion is not illegal, although it poses the same problem of hampering competition.

2.2 Previous literature

Any colluding practice, tacit or overt, leaves traces. The nature of these traces depend on how the firms coordinate their behavior in the market. In the literature of cartel detection, Harrington Jr and Chen 2006 have characterized several collusive pricing patterns. One central finding is that cartel prices are less sensitive to cost shocks than competitive prices. Thus low price variance can be considered as a marker for collusive behavior. Bolotova, Connor, and Miller 2008 find in their analysis of the citric acid conspiracy (1991-1995) and the lysine conspiracy (1992-1995) higher prices and lower price variance under the period of collusion. When the cartels were dissolved, prices decreased and price variance increase. Similarly, Abrantes-Metz et al. 2006 find an increase in the price variance and a decrease in the mean price in the retail gasoline industry in Louisville in 1996-2002 after the dissolution of a local cartel. By employing a model of an infinitely repeated Bertrand game where costs are private information but i.i.d over time and across firms both Athey, Bagwell, and Sanchirico 2004 and Harrington Jr and Chen 2006 find similar results.¹⁴

In repeated auction settings (like the "product-of-the-month" auction held by DPBA), Skrzypacz and Hopenhayn 2004 have demonstrated that tacit collusion is possible in the following way: If bidders publicly observe past bids or winners, they can coordinate on a rotating-winner pattern without ever explicitly communicating. The history of play (e.g. whose turn it was in the last auction) serves as a signal for whose turn is next, allowing collusion to be sustained through implicit understanding.¹⁵

Aoyagi 2003 formally demonstrated that a simple bid-rotation scheme can yield higher joint profits than competitive bidding, and that this can be an equilibrium of an infinitely repeated auction game. In such a rotation equilibrium, each firm wins the auction only periodically, but when it is its turn, the firm faces little competition (as others bid non-aggressively or not at all) and thus can secure the contract at an elevated price. The collusive profit is then shared over time as each conspirator gets designated wins. It is the bid-rotation mode of tacit collusion that I will study in this paper.¹⁶

Bid-rotation has previously been studied in the Swedish context. The first contribution to this literature was made by Cletus 2016. Starting from the assumption that each active bidder has an identical, independent chance of winning every month, the author develops an Overlapping-Permutations Test that flags suspicious winning patterns. Using roughly 1 900 off-patent products auctioned 61 times between March 2010 and March 2015, the screen identifies about 12% of them as showing clear signs of bid-rotation.¹⁷ The results show that average prices on flagged markets are fully five times those on comparable competitive markets. Further, price dispersion is found to collapse under coordination: the coefficient of variation of winning bids falls to about 6% on collusive markets versus 20% on matched controls.

This paper contributes to the collusion-detection literature by providing a winner-sequence screen

¹⁴Although they find them for different reasons. In (Athey, Bagwell, and Sanchirico 2004) results are explained by difficulties in sharing cost information between the cartel members and in (Harrington Jr and Chen 2006) the lower variance in price results from the disincentive of the cartel to pass thru cost shifters as this could increase the chance of detection.

¹⁵A question not pursued in the thesis is that of *how* collusive agreements are setup. A justified question in the context of tacit collusion is: How do firms develop a mutual understanding of coordinated strategies in absence of explicit communication? One answer to this question is provided by Byrne and De Roos 2019.

¹⁶Previous studies on the Swedish market for generics has found this mode of price coordination to be the most common by far. Granlund and Rudholm 2023 find it in over 95% of sub-markets flagged for signs of collusion. Other modes of coordinating prices under repeated auctions can be considered, such as parallel bidding which is studied by Cletus 2016.

¹⁷Collusion is observed almost exclusively in segments with only two or three active bidders.

that avoids the specification risks of structural bidding models. Its chief limitations lie in its simplifying assumptions on the probability of winning the auction in the absence of tacit collusion. The paper also uses a narrow definition of collusive signatures—episodes that do not generate neat alternation go undetected. Even so, the pioneering work of Cletus 2016 remains a useful benchmark for empirical work on the Swedish generics market.

Inspired by the work done by Cletus 2016, Granlund and Rudholm 2023 develop a Bayesian screening framework with similar goals. Whereas Cletus 2016 derived a simple model based on stringent assumptions, Granlund and Rudholm 2023 propose a more comprehensive framework that obviates them. The authors make use of Bayes theorem to obtain the probability of collusion, given a certain pattern of winners. Thus the paper moves the detection literature from binary red-flag alerts toward a continuous measure.¹⁸

Methodologically the two papers sit at opposite ends of a spectrum. Cletus 2016 supplies an indicator that requires nothing more than winner identities, but at the cost of assuming i.i.d. victories and treating every suspect episode alike. Granlund and Rudholm 2023, by contrast, embed a behavioral model that tolerates asymmetric success probabilities and outputs a continuous probability attached to every single auction. The effect on prices that follow from tacit collusion is in (Cletus 2016) derived from cross market comparisons. In (Granlund and Rudholm 2023) the comparison is made across time, when the probability of collusion shifts from certain competition to certain collusion.

The model developed by Granlund and Rudholm 2023 is a formidable attempt at dealing with collusion detection in a complex environment. Their more comprehensive model is attractive because it doesn't rely on stringent assumptions, yet its very breadth makes the communicability of the evidence somewhat lesser than that in (Cletus 2016). The logical chain from assumptions to conclusions is harder to trace in a sprawling model. For my own model, I will try to strike a balance between doing away with the stringent assumption made by (Cletus 2016), while at the same time maintaining a moderate level of model complexity.

My thesis will contribute three things to the collusion-detection literature. **(1)** I will link my results directly to the change in number of firms active in a market. I do this by setting up an overlapping permutations test that flags time-segments in the market *defined by the number of active firms*. This links the shift in number of firms to the change in flag. Using this shift I then set up a fixed panel regression and an event-study to do a pre/post comparison. This exercise will add to the understanding of how the number of firms affect the presence of collusive behavior. It will also provide suggestive evidence on what effects firm entry and exit have on prices in markets where collusion is present. The event-study will let us trace that effect out over time from when collusion ends.

(2) I will further the collusion detection literature by extending the work done by Cletus 2016 in two important ways. The overlapping permutations test that I develop will not rely on strong assumptions on the probability of individual firms winning the "product-of-the-month" auctions. Rather, it will utilize resampling methods for its hypothesis testing. **(3)** The test will also be able to handle noise in the pattern of winners by introducing a "fuzzy" scoring algorithm.

Why firm entry/exit is important to study is motivated by Game Theory. In the next section I present the game theoretical underpinnings of price competition and tacit collusion. This, together with a simple supply and demand diagram will provide the theoretical basis for why we expect the

¹⁸This continuous measure is useful for enforcement agencies when deciding on which markets to prioritize for investigation and intervention.

number of firms in a market to affect collusion; why prices would be higher; and why price variation would be lower.

2.3 Game theory model of collusion

The Bertrand theory of price competition predicts marginal pricing on the generics market, which is in contrast to empirical findings. One explanation for this discrepancy is that firms may be engaging in tacit collusion. To illustrate some key features of tacit collusion I will present a simple model from Ivaldi et al. 2003.

We start with a market with two price competitors, Firm A and B. Each firm set its price at p_A and p_B respectively. In the market there exists D consumers who each want to buy one unit. The willingness to pay for a unit is at most V . Marginal costs c are constant and we assume that $V > c$. We further assume homogeneous goods and preferences which means that consumers only care about the price of a good when deciding to buy it or not. If one firm offers a lower price, all demand is directed to that product. (Here, we see the clear link to the institutional setting in the Swedish generics market.) If the firms set equal prices demand is split evenly between the two firms. Thus the demand faced by each firm i is given by

$$d_i(p_i, p_{-i}) = \begin{cases} D & \text{if } p_i < p_{-i} \text{ and } p_i < V \\ \frac{1}{2}D & \text{if } p_i = p_{-i} \\ 0 & \text{otherwise} \end{cases} \quad (1)$$

The range in where the firms can set their prices is $[c, V]$ since $p = c$ is the lowest price a firm can offer without making losses and $p = V$ is the highest price consumers on the market will accept (it is also the price a monopolist would set). If the price is set only once, meaning that we model the market as a one shot game, the equilibrium outcome is that both firms set their price equal to marginal costs. This is commonly referred to as the Bertrand Trap. It is the Bertrand model of price competition embedded in Game Theory. The profit for each firm in this case would then be $\pi_i = \frac{1}{2}D(p - c) = 0$.

Now if firms could in some way coordinate and set the price to $p = V$ the profit would be $\pi_i = \frac{1}{2}D(V - c) > 0$. The problem is that if such an agreement were in place it would still be rational to deviate from the agreed price by some small amount $\varepsilon < \frac{1}{2}(V - c)$ and reap the profits $\pi_i = D(V - \varepsilon - c) > \frac{1}{2}D(V - c)$. By denoting the profits under cooperation $\pi_i^C = \frac{1}{2} \times D(V - c)$ π_i , and profits under defection $\pi_i^D = \frac{1}{2} \times D(V - \varepsilon - c)$ we can depict the situation as a normal form game presented in figure 3.

		Firm A	
		Cooperate	Defect
Firm B	Cooperate	π_B^C, π_A^C	$0, \pi_A^D$
	Defect	$\pi_B^D, 0$	$0, 0$

Figure 3: One stage competition with two firms

From figure 3 we can see that the only Nash Equilibrium is $\{0, 0\}$. This is in essence the classical

prisoners dilemma.¹⁹ While both firms would be better off *jointly* if they chose to cooperate, they will *individually* be better off if they deviate while the other firm chooses to cooperate. Since this individual incentive is symmetric, the only stable outcome is that both firms deviate. This is however only true if the game is played only once. If we extend the game and play indefinitely, coordination can be rational and thus possible. This is yet another classical result from Game Theory called the Folk Theorem.²⁰

Extending the game like this can be thought of as a repeated strategic interaction. New rounds of the same game are played in perpetuity. To this, we add a common discount factor $\delta \in (0, 1)$. This can be thought of either as the value players place on future outcomes of the game, or as a measure of their patience in waiting for reward to materialize.²¹ A higher discount factor means that future events are valued similarly to present ones, and vice versa. Lastly, we consider a strategy called the Grim Trigger. The strategy consists of two rules:

1. In the first round, charge $p = V$
2. Continue to charge $p = V$ in subsequent round as long as the other firm charged $p = V$ in the previous round. If the other firm charged a price $p < V$, charge $p = c$ in all (infinitely many) coming rounds.

Now the present value of cooperation becomes

$$PV_C = \pi^C \sum_{t=0}^{\infty} \delta^t = \frac{1}{2} D(V - c) \frac{1}{1 - \delta} \quad (2)$$

and the present value of deviation becomes

$$PV_D = \pi^D + \sum_{t=1}^{\infty} \delta^t \times 0 = D(V - \varepsilon - c) \quad (3)$$

Cooperation will be a best response if $PV_C \geq PV_D$, that is if

$$\frac{1}{2} D(V - c) \frac{1}{1 - \delta} \geq D(V - \varepsilon - c) \quad (4)$$

Now this expression need to hold for an arbitrarily small $\varepsilon < 0.5(V - c)$ which means that we can consider $\varepsilon = 0$. Substituting this into equation 4 we get

$$\begin{aligned} \frac{1}{2} \frac{D(V - c)}{1 - \delta} &\geq D(V - c) \\ \frac{1}{2} &\geq 1 - \delta \\ \delta &\geq \frac{1}{2} \equiv \delta^* \end{aligned} \quad (5)$$

¹⁹The prisoners dilemma has been studied by economists and political scientists ever since it was first formally formulated in (Flood 1958).

²⁰The name "Folk theorem" was mentioned first by Aumann and Shapley 1994, although the result dates back to Friedman 1971

²¹An equally valid interpretation of the discount factor is to view it as the probability of the game ending in the next round.

which means that the inequality in (4) holds if $\delta \geq \frac{1}{2} \equiv \delta^*$. This means that for two firms, their common discount factor has to be at least $\delta = \frac{1}{2}$ for collusion to be stable and we denote this threshold δ^* .

Extending the argument to a case with N firms we have a similar situation, however, the market is now split among more firms which means that the present value of cooperation becomes

$$PV_C = \frac{1}{N}D(V - c) \quad (6)$$

and the condition $PV_C \geq PV_D$ can be written like this

$$\frac{1}{N}D(V - c)\frac{1}{1 - \delta} \geq D(V - \varepsilon - c) \quad (7)$$

Again we can consider $\varepsilon = 0$ and get

$$\begin{aligned} \frac{1}{N} \frac{D(V - c)}{1 - \delta} &\geq D(V - c) \\ \frac{1}{N} &\geq 1 - \delta \\ \delta &\geq 1 - \frac{1}{N} \equiv \delta^*(N) \end{aligned} \quad (8)$$

The inequality in equation (8) means that as the number of firms in a market grows, the higher discount factor is required for collusion to be stable.²² As a numerical example, if the common discount factor is $\delta = \frac{3}{5}$ then two firms would sustain collusion since $\delta^*(N = 2) = 1 - \frac{1}{2} = \frac{1}{2} < \frac{3}{5} = \delta$. However, three firms would sustain collusion since $\delta^*(N = 3) = 1 - \frac{1}{3} = \frac{2}{3} > \frac{3}{5} = \delta$. If the discount factor reflects the interest rate, meaning that $\delta = \frac{1}{1+R}$ then one could also express (8) as $R \leq R^*(N) = \frac{1}{N-1}$. In this case the interest rate threshold would drop from 1 to 0.5 when the number of competitors rose from 2 to 3.

We can analyze the expression in 8 further by considering the limit as $N \rightarrow \infty$. We get

$$\lim_{N \rightarrow \infty} 1 - \frac{1}{N} = 1 < \delta \quad (9)$$

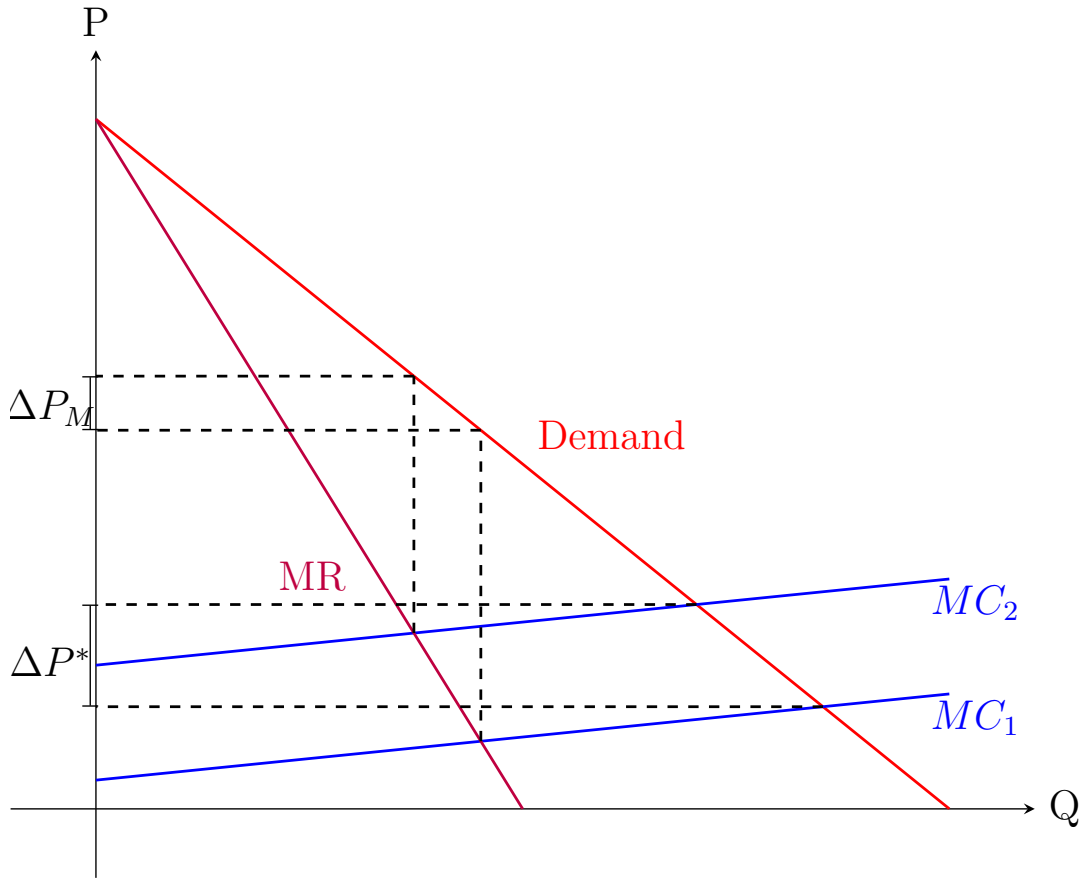
Here we see that as the number of firms in a market grows larger, collusion eventually becomes impossible. Remember that $\delta \in (0, 1)$ which means that we can never have a common discount factor equal to 1. Such a discount factor is ruled out by the model since it implies that firms value outcomes infinitely far into the future as much as they value those in the present.

It is from this theory that the two first predictions presented in section 1 are derived. The first of, **P1** states that: *Collusive behavior will be found more frequently in markets with few firms*. The parameter δ models the value placed on future outcomes. A higher δ thus implies a higher willingness to wait for future profits to materialize. While patience might be high among firms they will still prefer getting paid today than tomorrow. Thus it is reasonable to assume that higher values of δ will be shared by less firms. Therefore, when the number of firms increase, the required $\delta^*(N)$ for collusion to be sustain will be shared by fewer firms and thus collusive behavior will be found less frequently. Or put in another way: "Collusive behavior will be found more frequently in markets with few firms".

The second prediction **P2** states that: *The breakdown of collusive behavior will decrease the price*

²²One thing that is noteworthy of the expression in equation (8) is that it is independent of the market size.

Figure 4: Illustration of variation in price



in the market. This is motivated by the game theory model presented above. The goal of colluders in our context is to avoid the Bertrand trap, meaning that they coordinate rather than compete on prices. This is sustained as long as the smaller, but continuous profits from collusion are valued higher than the greater, but short term, profits from deviation. When collusion is no longer sustained the equilibrium will be that of price competition and thus prices will be lower.

The third prediction will be derived from the theory presented next.

2.4 Price variation under collusion

The third prediction presented in section 1, **P3** states that: *The breakdown of collusive behavior will increase the variation in prices in the market.* By simply assuming that a cartel acts as a monopoly we can employ a micro-economic model to show that shifts in marginal costs cause greater shifts in prices under competition than under monopoly/cartels. The reasoning is depicted in figure 4. Since the marginal revenue curve is twice as steep as the demand curve, a shift in the marginal cost will cause the equilibrium price to change less if the market is run by a cartel (that acts as a monopoly). This can be seen by comparing ΔP^M which denotes the price change when there is monopoly power, and ΔP^* which denotes changes in the equilibrium price under competition.

The intuition is this, since colluding firms act as a monopoly they have monopoly market power. With this they can adapt to a cost shock by both raising prices and supply. Therefore a shock to costs doesn't transmit fully to their offered price. Under market competition, firms can't influence the price

by changing the supply and thus the full cost-shock is carried over to the market price.

Harrington Jr and Chen 2006 provide a different model to explain the difference in price variation under collusion compared to competition. Below I present a simplified version of their model. We assume an auction game where bidding repeats infinitely (as in the Game theory model): the time horizon is $t = (1, 2, \dots)$, and firms share a common discount factor $\delta \in (0, 1)$. Each period firms in the market face a cost shock $c_t \sim \text{i.i.d.}(\mu, \sigma^2)$. We can think of this as a shock to some ingredient in the active substance for a given substitution group. Under pure competition, firms set prices to $P_t = c_t$. We posit that the market can be in either two states, competition and collusion. We can then express the price as

$$P_t = \alpha + \varphi c_t \begin{cases} \varphi = 1, \alpha = 0 & \text{(competition)} \\ \varphi = \varphi^* \in [0, 1), \alpha > 0 & \text{(collusion)} \end{cases} \quad (10)$$

where α is the cartel mark-up and φ^* is a parameter which dampens the transmission of cost into prices. What (Harrington Jr and Chen 2006) show is that a cartel that fears detection (either from buyer or competition authorities) will tone down its price reactions to cost shocks, meaning that $\varphi^* < 1$. When firms compete, no such fear is present and the cost fully "bleed-through" into prices ($\varphi = 1$).

Under competition, $P = c_t$ and $\text{Var}(P_t) = \varphi^2 \sigma^2 = \sigma^2$ since $\varphi = 1$ but when firms collude, $P_t = \alpha + \varphi c_t$ and $\text{Var}(P_t) = \varphi^{*2} \sigma^2 < \sigma^2$ since $\varphi^* < 1$. The moment the agreement breaks – through internal defection or external prosecution – φ snaps back to one, and the observed variance jumps up to its competitive benchmark.²³

A third reason for price variance being lower under bid-rotation is the structure of the bidding scheme. To maintain high profits, the colluding firms take turns on placing the lowest bid and thereby capturing the whole market. For a firm to maximize profits when it is their turn to win, they would want to bid as close to the next lowest bid. Conversely, when markets are characterized by competition, we can expect more aggressive bidding from firms since they don't have a guaranteed win in a particular month (as firms have in a bid-rotation setup).

Based on the above discussion we can form the prediction **P3** that *The breakdown of collusive behavior will increase the variation in prices in the market.*

3 Data

The data that I have collected comes from the Dental and Pharmaceutical Benefits Agency (DPBA). The panel data spans over fifteen years and has data on prices and sales volumes for all unique medical products defined by the NPL system within the government-funded benefits scheme (*högkostnadsskyddet*).²⁴

The initial raw data was compiled into a data set containing roughly one million observations of prices over time for 3030 products. These were sold within 2957 different substitution groups and produced by 191 unique firms. As stated in section 2.1 a substitution group represents a sub-market within a market for an active substance.

²³This logic has been used by (Abrantes-Metz et al. 2006) to develop a collusion screen.

²⁴The National Register for Medicinal Products

A restriction was put on the data so that it only includes sub-markets (substitution groups) which had more than 5% of sales of the total market (substance level).²⁵ Small sub-markets are excluded for two reasons. Firstly, small markets have a small impact on overall medical costs which makes them of limited interest. Secondly, products in small sub-markets are likely produced for a class of consumers we certain, uncommon needs, and thus highly specialized. Therefore there is reason to believe that the markets for these products have some features that differentiate them from markets for products sold to a more general group of patients.

The number of sub-markets and their market share in sales volume are presented figure 5 for this restricted sample. The corresponding histogram for the whole sample is found in figure 15 in appendix A. Since both Cletus 2016 and Granlund and Rudholm 2023 investigate the same market as I do, a comparison of samples seems appropriate. The sample used by Cletus 2016 spans the years 2010 – 2015 and contains around 1 000 sub-markets. The data used by Granlund and Rudholm 2023 spans the years 2010 – 2020 and contains around 1 500 sub-markets. Thus my sample spans a longer time horizon and includes more sub-markets.

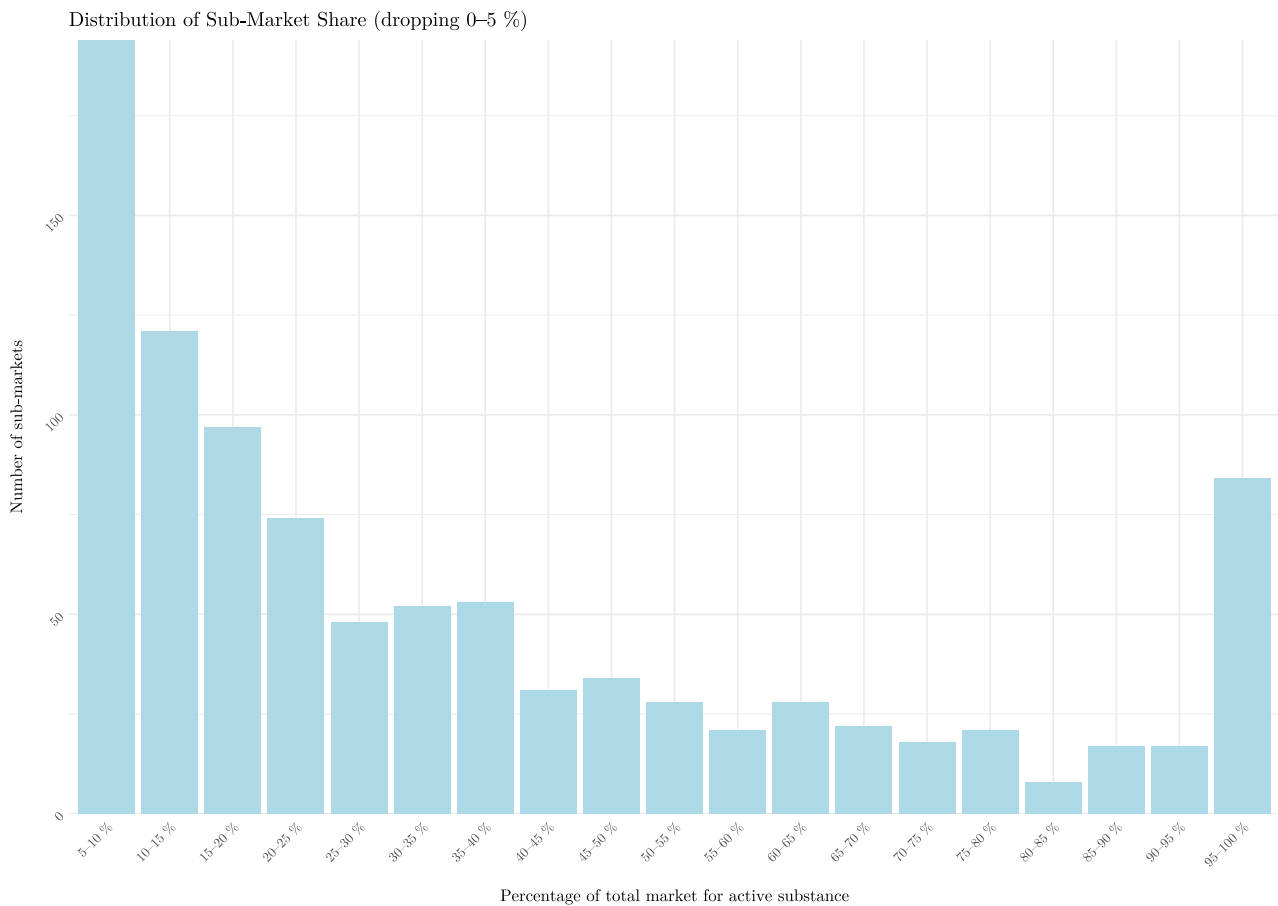


Figure 5: Sub-market size – restricted sample

Summary statistics for the restricted sample is presented in table 1. The number of observations has dropped by a fifth compared to the original data. The data set now contains 2843 products which were sold on 1852 different sub-markets and produced by 180 unique firms. Along with summary statistics for the whole sample, the table reports on the characteristics of sub-markets within each

²⁵The total market is defined by the active substance and a sub-market by the combination substance $\times$ strength $\times$ package $\times$ size.

substance and characteristics for the products within those sub-markets.

Table 1: Summary Statistics of Data

Time period	2010 (Jan 1st)	2025 (Jan 2nd)	
N. Obs		846 679	
<i>Across Panel</i>			
Number of Substitution Groups (sub-markets)		1 852	
Number of Firms		180	
Number of products		2 843	
Number of sub-market segments		11 843	
<i>Across Substances</i>			
<i>Sub-markets</i>		<i>Products within sub-markets</i>	
Median N. of sub-markets	3	Median N. of products	3
Mode N. of sub-markets	2	Mode N. of products	2
Mean N. of sub-markets	4.68	Mean N. of products	3.95
SD N. of sub-markets	4.53	SD N. of products	2.91
Min N. of sub-markets	1	Min N. of products	1
Max N. of sub-markets	38	Max N. of products	25

For each substance level in the sample, there are different substitution groups/sub-markets. Returning to the example in section 2.1 of the substance *ibuprofen*. For this substance there is one substitution group defined as $ibuprofen \times 400\text{mg} \times tablets \times 30$. Another common substitution group (sub market) for this substance is defined similarly but with a lower strength of 200mg instead of 400mg. A third could be defined by products that have a larger package size of 50 tablets. If these were the only substitution groups for the substance *ibuprofen* then there would be 3 sub-markets within this substance level.

The number of sub-markets for each substance is quite a disperse. We see this by looking at the bottom left panel in table 1. The maximum number of sub-markets within one active substance is 38 while the minimum number is 1. The three measures of centrality are located between 2 and 4.7 and the standard deviation is 4.53. This means that for most of substances in the sample, the number of sub-markets is between 1 and 10. This is clearly seen in figure 6.

Within sub-markets there are different products. Again consider the case with the substitution group (sub-market) defined by $ibuprofen \times 400\text{mg} \times tablets \times 30$. Within this group there are a number of different products being sold under brands such as "Ipre", "Ibuprofen Orifarm" and "Ibumetin". If these were the only brands sold within this sub-market, the number of products would be 3.

Within sub-markets, the dispersion of the number of products is also quite large, ranging from 1 to 25. The measures of centrality are located between 2 and 3.95 with a standard deviation of 2.91.

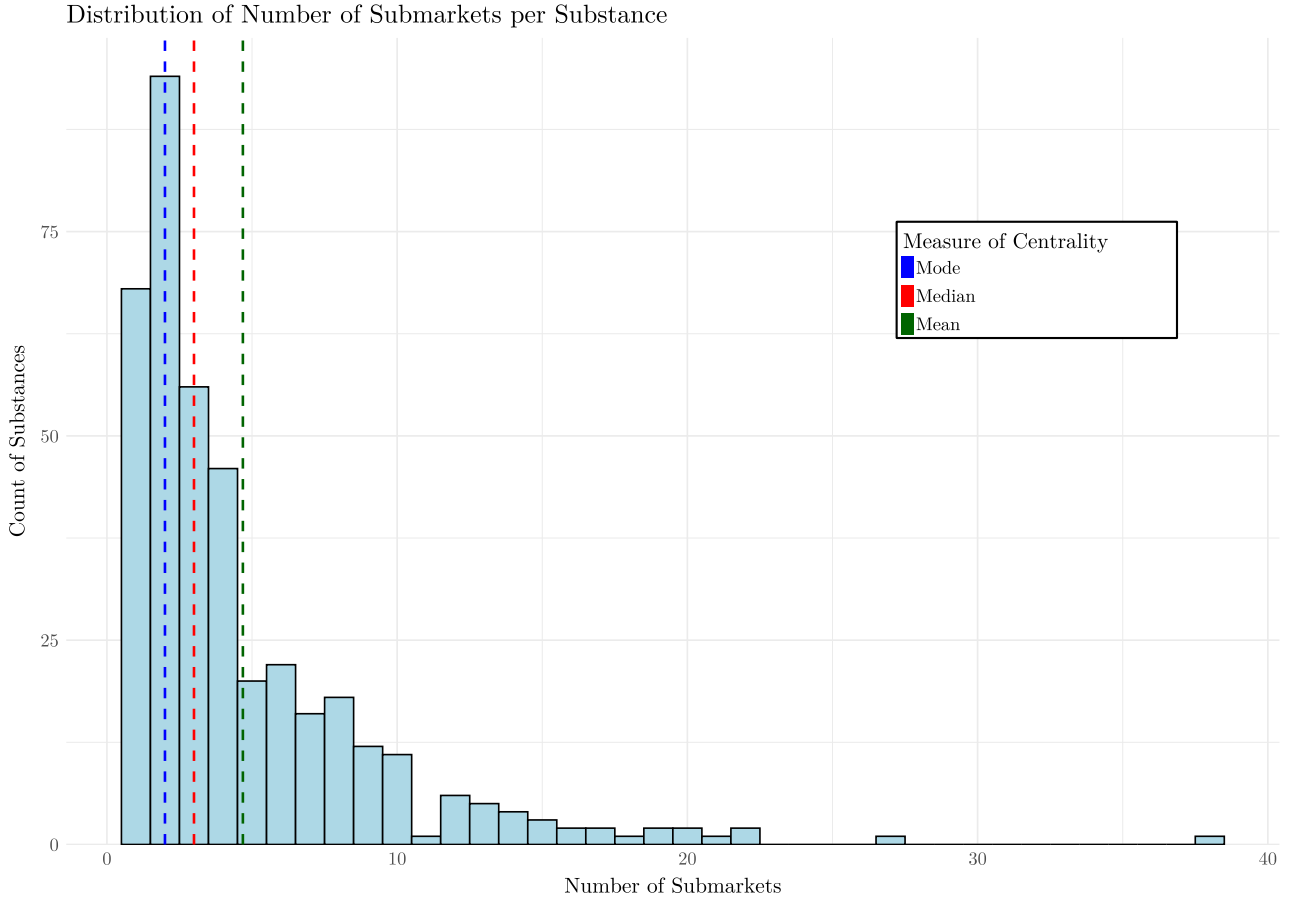


Figure 6: Distribution of number of submarkets

This means that within sub-markets, the common number of products sold is between 1 and 7.

Figure 7 depicts a sub-market in my sample. When I set up my overlapping permutations test I will denote them by M . There are 1852 of these in total in the data set (see table 1). Each period between two dashed lines I call a "sub-market segment" and denote it M_T . The period of each of these sub-market segments has varying time length, measured in moths, which I denoted by T . The number of active firms with each sub-market segment will be denoted by F .

By looking at the bottom panel of figure 7 we can see how the number of active firms within this sub-market changes over time. In the top panel, each change in the number of firms is marked by a dashed line. The time period between two dashed lines is thus a time period where the number of firms doesn't change.

Within the particular sub-market M that is depicted in figure 7, there are 8 different sub-market segments M_T with varying lengths T . For the sub-market segment M_T that starts around the year 2019 and ends somewhere in the year 2022, T is equal to $12 \times 3 = 36$, assuming the period lasted for exactly three years. The number of active firms within this sub-market segment F is equal to 4 which can be read off the y-axis in the bottom panel.

4 Empirical strategy

The empirical strategy used in this thesis consists of two steps. In the first step I conduct an overlapping permutations test to identify sub-markets time segments (M_T) where there is evidence of

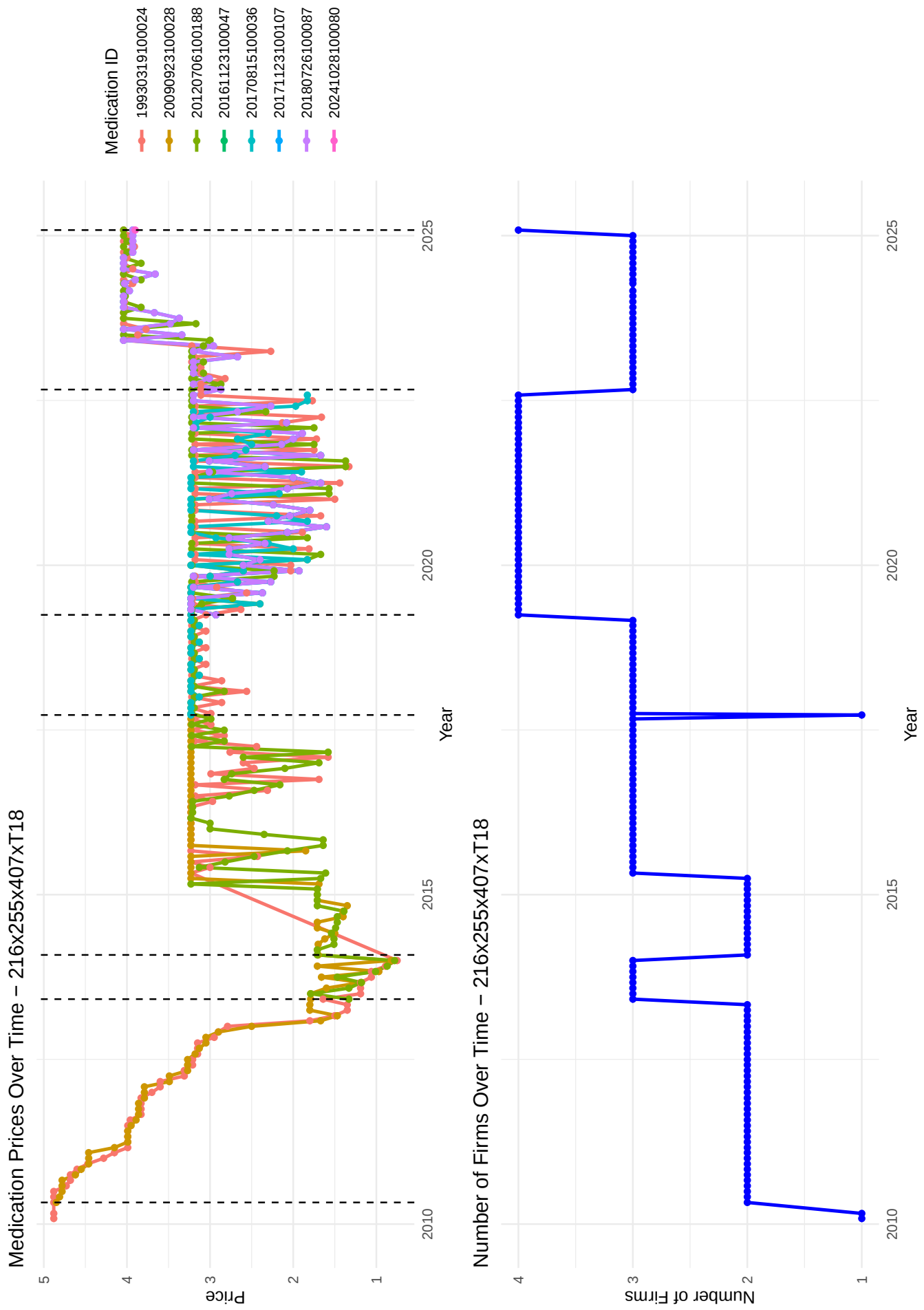


Figure 7: Example of substitution market

bid-rotation. In the second step, I will compare price levels and price dispersion during, and after bid-rotation, using two types of pre/post analysis – a fixed-effects panel regression and an event study.

The overlapping permutations test serves two purposes. First, it will be used to test the prediction **P1** stated in section 1. Second, the result from this test will provide a list of time segments across sub-markets where evidence of bid-rotation is found. This list will then be used to construct the sample used in the pre/post analysis which will be used to test prediction **P2** and **P3**.²⁶ In this section I first lay out the design of my overlapping permutations test. Then I provide formulations of my fixed-effects panel regression and event-study design.

4.1 Overlapping permutations test

4.1.1 General idea

The situation where firms each month present bids to DPBA to gain the product-of-the-month status gives rise to a sequence of winners (Firm 1, Firm 1, Firm 2, Firm 3, etc.) for each sub-market. If a pattern of cyclically repeated winners is observed we have reason to suspect bid rotation and therefore, tacit collusion. One way of detecting such a pattern is to use an overlapping permutations test. The basic procedure is to cut the original sequence into partitions which overlap each other, and then look for signs of bid-rotation within each partition. This allows us to investigate the dependencies between consecutive outcomes which relates directly to the idea of bid-rotation, and the cyclical pattern in winners that it gives rise to. If the entire sequence of winners has a cyclical pattern, then this will show up in the partitions as well. Moreover, the overlapping partitions let you zoom in on temporal structures at smaller scales compared to when analyzing the entire sequence.²⁷

The number of active firms in a sub-market at a given time is central to testing prediction **P1**. Given the theory presented in section 2.3, we can expect collusion, and thus bid rotation, to be more common in markets with fewer firms. And since the number of firms in a given sub-market changes over time, the test will be implemented in time periods within sub-markets where the number of firms is constant.

To get a feel for how the number of firms evolve over time in such a market I have provided an example of this in 7 found in section 3. To repeat some notation: a sub-market in my sample is denoted by M (1852 in total). Each period where the number of firms is stable within a sub-market is called a "sub-market segment" and is denoted by M_T (depicted by two dashed lines in figure 7). The period of each of these sub-market segments is measured in months and denoted by T . The number of active firms with each sub-market segment is denoted by F .

I will test for rotational bidding among two, three, four, five and six firms. I will call the number of the pattern n , so $n \in \{1, 2, 3, 4, 5, 6\}$. The pattern will determine the length of the partitions which I will call k . For each n , the length of the partitions will be $k = 2n$. This means that for two-firm bid-rotation, I will form partitions of length 4, for three-firm bid rotation the length of the partitions

²⁶The predictions are:

P1: Collusive behavior will be found more frequently in markets with few firms.

P2: The breakdown of collusive behavior will decrease the price in the market.

P3: The breakdown of collusive behavior will increase the variation in prices in the market.

²⁷The procedure could potentially also give more statistical power since each partition is one unit of analysis.

will be 6, and so on.

The choice of length of the partitions comes down to balancing two things. Too short partitions run the risk of missing temporal dependencies in the data – the procedure becomes too "zoomed in". Too long partitions create the opposite risk, a too "zoomed out" procedure where the partitions pick up time dependencies that should not be included (Hall, Horowitz, and Jing 1995). The reason for my choice of setting the length of the partitions to $k = 2n$ is that with this length, any pattern of bid rotation will repeat two cycles within a permutation. As an example, a permutation of length 4 could look like (Firm 1, Firm 2, Firm 1, Firm 2), repeating the two firm bid-rotation between Firm 1 and 2, two times. I believe this length of $k = 2n$ strikes a balance between taking temporal dependencies into account while not being too long.

Depending on the choice of rotational pattern to test for, some sub-market segments M_T cannot be tested. The most obvious case is sub-market segments where the number of firms $F = 1$. If only one firm is active, no rotational bidding can occur. If $F = 2$, then we can test for two-firm bid-rotation. If $F = 3$, then we can test for both two-firm and three-firm bid rotation. In general, for a given sub-market segment M_T with an active number of firms equal to F , we can test any pattern $n \leq F$.

Testing all patterns $2 \leq n \leq F$ is important since some firms who produce an original brand can charge higher prices than those producing generic drugs. This is because some customers prefer the original brand over the generic(s). In Granlund and Rudholm 2012 the authors find that 83% of patients agreed to substitution when they had an option. In the data used by Granlund and Bergman 2018, 10.5% of patients opposed substitution. This indicates that a non-trivial part of patients preferred the original brand.

There could be other reasons for not choosing the product of the month. While products are identical in the four dimensions defining a substitution group (substance $\times$ strength $\times$ form $\times$ package), they can differ in other aspects. Some packaging is harder to open than others, which is why pharmacists may advice against a particular product for elderly people, even if it is the "product-of-the-month". In some cases, doctors don't prescribe a substance $\times$ strength $\times$ form $\times$ package combination but an actual product. In that case pharmacists can't recommend generic substitution. Granlund and Bergman 2018 report that in their data, physicians opposed substitution for 3.4% of the packages and pharmacists for 2.0%. The estimates in (Bergman, Granlund, and Rudholm 2017) are similar.

Thus there could be cases of sub-market segments where three firms are active but the firm who produce the original drug face a different demand than the other two. It may be preferred by physicians, patients or pharmacists. In such a case, the two generics firms could have incentives to engage in tacit collusion to keep prices higher while the third has no such incentives.²⁸ The firm producing the original brand could then still charge higher prices than the colluding firms would under their collusion scheme.

The choice of rotational pattern also precludes some sub-market segments M_T from testing. Since the length of the partitions is set to $k = 2n$, a sub-market segment M_T needs to have at least $T \geq k = 2n$ time periods. This means that the test for two-firm bid-rotation can only be carried out on segments which is at least four months long. For a three-firm rotation, segments need to be at least six months long, a four-firm rotation requires eight months, and so on.

²⁸This precise structure of two colluding firms and one originator constantly setting a higher price is discussed in (Granlund and Rudholm 2023) (see footnote 2 in the introduction).

The test is carried out on all sub-market segments M_T within all sub-markets M . Here, there is a distinct difference between my approach and that of (Cletus 2016) and (Granlund and Rudholm 2023). Where they consider bid-rotation on the sub-markets as a whole, I turn my attention to these sub-market time-segments. This is theoretically motivated. One core element of the theory presented in section 2.3 is the effect of the number of firms on colluding behavior. This gives reason to believe that the data-generating process differs depending on how many firms are active. It is therefore essential that I define these sub-market time-segments based on the number of active firms and conduct a separate test within each. It is also necessary that the test are flexible enough to handle several segments with differing number of active, and possibly colluding firms.

For all sub-market segments M_T within a sub-market M , there is a sequence of the winners of the auction for the product-of-the-month denoted by $\{W_t\}_{t=1}^T$. The variable W_t denotes the identity of the winning firm which means that $W_t \in \{1, 2, \dots, F\}$. As an example, if there are $F = 3$ active firms in a sub-market segment M_T , and the sub-market segment spans over twelve months, such a sequence could look like

$$\{W_t\}_{t=1}^{12} = (2, 3, 1, 3, 1, 2, 3, 2, 1, 1, 2, 3)$$

which would mean that in the first month in that sub-segment, firm two won the auction, then firm three won, followed by firm one, and so on.

Next we form partitions P_t of the sequence which will overlap each other by one time period. The length of the partitions will depend on the rotational pattern that is being tested. For this example with $F = 3$, I will test both two-firm and three-firm bid-rotation. Thus two tests will be conducted on this sub-market segment and I will therefore form two sets of partitions.

For the two-firm test the length of the partitions will be $k = 2n = 4$. For the example sequence above of length $T = 12$, we can thus form 9 partitions of length 4.²⁹ As an example we can look at the first and last two partitions for a two-firm test on the above sequence.

$$\begin{aligned} P_1 &= (W_1, W_2, W_3, W_4) = (2, 3, 1, 3) \\ P_2 &= (W_2, W_3, W_4, W_5) = (3, 1, 3, 1) \\ &\vdots \\ P_8 &= (W_8, W_9, W_{10}, W_{11}) = (2, 1, 1, 2) \\ P_9 &= (W_9, W_{10}, W_{11}, W_{12}) = (1, 1, 2, 3) \end{aligned}$$

For the three firm test, the length of the partitions will be $k = 2n = 6$. We can form 7 such partitions on the example sequence. The sequence of partitions in this case becomes $\{P_t\}_{t=1}^7$ ³⁰ and the first and last two partitions in that sequence would look like this

²⁹These 9 partitions are also a sequence of the form $\{P_t\}_{t=1}^9$.

³⁰In general, for a rotational pattern $n \in \{1, 2, 3, 4, 5, 6\}$ we can form $T - k + 1$ partitions given that $T \geq k$. The general form of a sequence of partitions of length k made from a sequence of winners $\{W_t\}_{t=1}^T$ with $T \geq k$ can then be written as $\{P_t\}_{t=1}^{T-k+1} = \{W_t, W_{t+1}, \dots, W_{t+k-1}\}_{t=1}^{T-k+1}$.

$$\begin{aligned}
P_1 &= (W_1, W_2, W_3, W_4, W_5, W_6) = (2, 3, 1, 3, 1, 2) \\
P_2 &= (W_2, W_3, W_4, W_5, W_6, W_7) = (3, 1, 3, 1, 2, 3) \\
&\vdots \\
P_6 &= (W_6, W_7, W_8, W_9, W_{10}, W_{11}) = (2, 3, 2, 1, 1, 2) \\
P_7 &= (W_7, W_8, W_9, W_{10}, W_{11}, W_{12}) = (3, 2, 1, 1, 2, 3)
\end{aligned}$$

For a two-firm test on a sub-market segment with $F = 3$ and $T = 12$, the first partition of the sequence $\{P_t\}_{t=1}^9$ can result in $3^4 = 81$ possible combinations of the digits 1, 2 and 3.³¹ When we move on to the second partition, there is however only 3 possible partitions. This is because the first and second partitions, P_1 and P_2 , share the elements W_2 , W_3 and W_4 . The only new element in P_2 is W_5 , which in this example only can take on three different values. This relationship between the partitions is what allows us zoom in on temporal structures in the data since the partitions become serially correlated.

4.1.2 The statistical test

The next step is to use the setup described thus far to formulate a hypothesis test. The test should be able to determine if a rotational pattern differs from what we would expect under normal price competition. To gather evidence of bid-rotation taking place we count the number of partitions that display such a pattern. A two-firm rotational pattern within a permutation would in general look like $\bar{\omega} = (i, j, i, j)$ for $(i, j) \in (1, \dots, F)$ and $i \neq j$. We now take this pattern and try to match it against the permutations in the sequence.

For the example above with two-firm rotation we can see that $P_2 = \bar{\omega} = (i, j, i, j)$. In this case, $F = 3$, $P_2 = (3, 1, 3, 1)$ and thus $i = 3$ and $j = 1$. No other rotational pattern is found in any other of the 9 partitions of the sequence $\{W_t\}_{t=1}^{12}$, so the total count for this sequence of winners is 1.

The general formula of summing over all permutations takes the form

$$N_{\bar{\omega}} = \sum_{t=1}^{T-k+1} \mathbb{1}(P_t = \bar{\omega}) \quad (11)$$

where $\mathbb{1}$ denotes the indicator function. For the three-firm rotation test, the same steps are applied but the pattern is now defined as $\bar{\omega} = (i, j, l, i, j, l)$ for $(i, j, l) \in (1, \dots, F)$ and $i \neq j \neq l$. Likewise alterations are made for four-, five- and six-firm rotation patterns.

"Fuzzy" scoring Finding a large $N_{\bar{\omega}}$ indicates a sequence of winners that is very cyclical, or "rotation-like". It could therefore pose as strong evidence for bid-rotation. There could however exist sequences that look very cyclical but don't get a large $N_{\bar{\omega}}$. To exemplify the problem, consider the following two sequences:

³¹Each partition is an element of $\{1, 2, \dots, F\}^k$ which means that we are considering all possible k -tuples of firm identities that can occur of a partition of k consecutive months. Given $F = 3$ and $k = 4$ there are 81 possible such tuples.

$$\begin{aligned}\{W_t\}^{(a)} &= (1, 2, 1, 2, 3, 2, 1, 2) \\ \{W_t\}^{(b)} &= (1, 2, 1, 2, 1, 2, 1, 2)\end{aligned}$$

Using the counting method described above, the sequence $\{W_t\}^{(a)}$ would yield $N_{\bar{\omega}} = 1$ while $\{W_t\}^{(b)}$ would yield an $N_{\bar{\omega}} = 5$: the maximum value for a sequence of length 8. Even though the sequence $\{W_t\}^{(a)}$ $\{W_t\}^{(b)}$ differs in just one place compared, their counts vary greatly. The problem here is that both series could come from the same data generating process, i.e they could both be the outcome of bid-rotation among firms one and two. The only difference could be that the sequence $\{W_t\}^{(a)}$ is slightly more noisy. This means that for some reason, firm three happened to win the auction in month five, even though firms one and two coordinated their prices by bid-rotation.

The evidence for bid rotation in $\{W_t\}^{(a)}$ is of course smaller than for $\{W_t\}^{(b)}$ but the difference of $N_{\bar{\omega}} = 5$ v.s $N_{\bar{\omega}} = 1$ seems exaggerated. To remedy this I've added a "fuzzy" scoring function. It is a non-linear transformation of the Hamming distance $d(t)$ from rotation pattern to the observed pattern.³²

$$S_t(P_t) = \left(\frac{k - d(t)}{k} \right)^p \text{ for } p \geq 1 \quad (12)$$

The function takes in a permutation and then calculates the number of positions where the partition deviates from the rotation pattern under investigation (The Hamming distance). For a two-firm rotation pattern, the length of each partition is equal to $k = 2n = 4$ which means that the partition can be off in the investigated pattern in 4 places. Thus $k = 4$ and $d(t) \in (1, 2, 3, 4)$. The deviation is then transformed non-linearly. The parameter p governs how sensitive the function should be to deviations, i.e potential noise. To see how it works, consider the first three permutations of the sequence $\{W_t\}^{(a)}$.

$$\begin{aligned}P_1(\{W_t\}^{(a)}) &= (1, 2, 1, 2) & \text{with } d(t) &= 0 \\ P_2(\{W_t\}^{(a)}) &= (2, 1, 2, 3) & \text{with } d(t) &= 1 \\ P_3(\{W_t\}^{(a)}) &= (1, 2, 3, 2) & \text{with } d(t) &= 1\end{aligned}$$

In the first permutation, the pattern deviates from $\bar{\omega} = (1, 2, 1, 2)$ in zero places so the deviation variable becomes $d(t) = 0$. In the second and third permutations, the pattern deviates on one place so $d(t) = 1$.³³ A deviation of $d(t) = 0$ is a "perfect match" and by equation (12) it gives the same score as equation (11) would, irregardless of choice of the parameter p , while the scores $S_t(P_2)$ and $S_t(P_3)$ depends on the choice of the parameter p .

³²The exact mathematical formulation and details of how to calculate this deviation is found in appendix B.

³³The pattern is defined as $\bar{\omega} = (i, j, i, j)$ for $(i, j) \in (1, 2)$ in this case. The flexible design makes it so that both the pattern $\bar{\omega} = (1, 2, 1, 2)$ and $\bar{\omega} = (2, 1, 2, 1)$ is considered as the two-firm bid-rotation pattern in the calculation of the deviations. This is what allow us to identify the deviation in P_2 at the fourth index. The pattern $P_2 = (2, 1, 2, 3)$ differs from $\bar{\omega} = (2, 1, 2, 1)$ in exactly one spot.

In table 2 we can see how much deviations are penalized given different values of p . In the table, partitions of length $k = 4$ have been considered. In the table we see that a perfect match always gets the score of 1 and a complete mismatch always gets the score of 0. For $p = 1$ the score decreases with 0.25 for each mismatch. For $p = 2$ the steepness increases and already with two mismatches, the score is reduced by 75%. With $p = 3$ all mismatches of 2 and above gets almost zero in score.

Table 2: Fuzzy scoring for various Hamming distances $d(t)$ and exponents p .

$d(t)$	$p = 1$	$p = 2$	$p = 3$
0	1	1	1
1	0.75	0.56	0.42
2	0.50	0.25	0.13
3	0.25	0.06	0.02
4	0.00	0.00	0.00

Note: In this table, different scores have been calculated using the equation $S_t(P_t) = \left(\frac{k-d(t)}{k}\right)^p$ for partitions of length $k = 4$.

Numbers are rounded.

To calculate the score for an entire sequence, we just sum over all $S_t(P_t)$. We thus define a new variable which is expressed as

$$N_S = \sum_{t=1}^{T-k+1} S_t(P_t) \quad (13)$$

Table 3 shows how the choice of parameter p affects the score N_S for the entire sequence $\{W_t\}^{(a)}$. The score N_S for sequence $\{W_t\}^{(b)}$ is 5 for all choices of p since all permutations of that sequence are perfect matches. Thus we should compare the numbers in table 3 that number. The score which results from using $p = 1$ is 4 which is quite close to the score of a perfect sequence which is 5. With $p = 2$, the sum of scores is 3.25 or 65% of that of the perfect sequence. For $p = 3$ the score is 2.7 which is 54% of the score for the perfect sequence.

Sequence $\{W_t\}^{(a)}$ represents a pattern that might be the outcome of two possibilities. Either, it is a bid-rotation pattern with noise, i.e, the "3" in the sequence is the result of chance. Or, it is *not* bid-rotation pattern and what looks like a bid-rotation pattern of 1's and 2's has come about through randomness. What the scoring function should do is to balance the risk of disregarding too many of the first cases (Type I errors), and the risk of accepting too many of the second case (Type II errors).

Based on the results in table 2, I believe that setting $p = 2$ is a reasonable choice. It gives low scores to partitions that have two mismatches which I think is important. In such a partition, only half of the positions match the pattern. I believe that the evidence of bid rotation is small in such a case. Also, with $p = 2$ the scoring function assigns a score of $S_t(P_t) = \frac{1}{2}$ to a partition that has one mismatch. This seems reasonable since I don't believe that this should be treated as strong evidence, but still not be disregarded entirely.

Further, I believe that a decrease in the sum of the scores from 5 to 3.25 seems as a reasonable representation of how much less we should believe that there is evidence of bid-rotation in sequence

Table 3: Fuzzy scoring for the sequence $\{W_t\}^{(a)}$ using different values for parameter p

$\{W_t\}^{(a)}$	(1, 2, 1, 2, 3, 2, 1, 2)					
t	1	2	3	4	5	
$P_t(\{W_t\}^{(a)})$	(1, 2, 1, 2)	(2, 1, 2, 3)	(1, 2, 3, 2)	(2, 3, 2, 1)	(3, 2, 1, 2)	
$d(t)$	0	1	1	1	1	
$S_t(P_t ; p = 1)$	1	0.75	0.75	0.75	0.75	$N_S = 4$
$S_t(P_t ; p = 2)$	1	0.56	0.56	0.56	0.56	$N_S = 3.25$
$S_t(P_t ; p = 3)$	1	0.42	0.42	0.42	0.42	$N_S = 2.7$

Note: In this table, the scores $S_t(P_t)$ for each of the five partitions $\{P_t\}_{t=1}^5$ of length $k = 4$ formed from the sequence $\{W_t\}^{(a)}$ have been calculated. The scoring equation $S_t(P_t) = \left(\frac{k-d(t)}{k}\right)^p$ has been calculated using parameter values $p = (1, 2, 3)$. Numbers are rounded

$\{W_t\}^{(a)}$ compared to the evidence in $\{W_t\}^{(b)}$. In my opinion, the other results in 3 decreases the strength of the evidence to little in the $p = 1$ case, or to much in the $p = 3$ case. Nevertheless, the choice of $p = 2$ is arbitrary. Without a more formal justification for this choice, I will run the test with different values of p to investigate the impact of my choice further. These tests are presented in section 4.1.3.

The "fuzzy" scoring function presents us with a way to quantify the cyclicity of winners in a sequence, i.e how much a sequence looks like bid-rotation. We now need a way to make a comparison between this count and what we could expect under normal price competition. For this, I turn to the Stationary block bootstrapping method of re-sampling.

Stationary block bootstrapping The main goal of the use of re-sampling is to determine if a cyclical pattern in winners is substantially more cyclical than what we could expect under price competition. If it is, then we have strong evidence for bid-rotation. The specific method chosen, the Stationary block bootstrapping (Politis and Romano 1994), provides us with a way to make such a comparison.

The basic idea of the Stationary block bootstrap approach is to resample (with replacement) the sequence under consideration many times to generate a number of synthetic sequences $\{W_t^{(r)}\}_{t=1}^T$. This can be seen as a way to simulate many samples being drawn from the underlying distribution. For each synthetic sample we then calculate its score $N_S^{(r)}$ and then compare the score for the original sequence, N_S , to the distributions of scores for the synthetic sequences $N_S^{(r)}$.

Re-sampling is done by the following procedure. We pick a starting point in the sequence at random from the uniform distribution with support over $\{1, 2, \dots, T\}$. We call this first starting position I_1 . The probability of choosing any $I \in \{1, 2, \dots, T\}$ is then $\mathbb{P}(I_1 = I) = \frac{1}{T}$ and is thus equal for all indexes in the sequence. Now we pick another number at random which we call L_1 . This number L is picked from the geometric distribution $\mathbb{P}(L_i = m) = (1-g)^{m-1}g$ where $g \in [0, 1]$ and $m \in \{1, 2, \dots\}$.³⁴ We use these two numbers to form our first block b_{I_1, L_1} of the first re-sampled sequence $\{W_t^{(1)}\}_{t=1}^T$.

³⁴The role and choice of the parameter g will be discussed in the coming section 4.1.3 on robustness test

Below is a demonstration.

Consider again our example sequence of winners

$$\{W_t\}_{t=1}^{12} = (2, 3, 1, 3, 1, 2, 3, 2, 1, 1, 2, 3)$$

Let's say that our first random number $I_1 = 7$, and our second number $L_1 = 3$. This means that our first block should start with the observation found at index 7 in our original sequence and include the 3 subsequent observations in order.³⁵ To spell it out: Our first block would be $b_{I_1, L_1} = (W_7, W_8, W_9, W_{10},) = (3, 2, 1, 1)$.

We then randomly select two new numbers, I_2 and L_2 by the same procedure. Let's say these are $I_2 = 4$ and $L_2 = 3$. Our next block then becomes $b_{I_2, L_2} = \{W_t\}_{t=4}^7$ which means that $b_{I_2, L_2} = (W_4, W_5, W_6, W_7,) = (3, 1, 2, 3)$. The algorithm continues like this until the sequence of blocks $\{b_{I_1, L_1}, b_{I_2, L_2}, \dots\}$ contain T (in this case $T = 12$) number of observations.

The algorithm has a circular design, meaning that if we randomly select $I_3 = 11$ and $L_3 = 4$ then the block "wraps around" the original sequence.³⁶ With these number, the third block becomes $b_{I_3, L_3} = \{W_t\}_{t=11}^{12} \cup \{W_t\}_{t=1}^2$ ³⁷ or more clearly, $b_{I_3, L_3} = (W_{11}, W_{12}, W_1, W_2,) = (2, 3, 2, 3)$. With these three example blocks we can now assemble our first synthetic sequence $\{W_t^{(1)}\}_{t=1}^{12}$.³⁸

$$\begin{aligned} \{W_t^{(1)}\} &= \{b_{I_1, L_1}, b_{I_2, L_2}, b_{I_3, L_3}\} \\ &= (W_7, W_8, W_9, W_{10}, W_4, W_5, W_6, W_7, W_{11}, W_{12}, W_1, W_2,) \\ &= (3, 2, 1, 1, 3, 1, 2, 3, 2, 3, 2, 3) \end{aligned}$$

We repeat this process R amount of times. For each synthetic sequence $\{W_t^{(r)}\}_{t=1}^{12}$ we go over the procedure with the partitions described above. We can then calculate the scoring function according to equation (13). Using all of R synthetic scores $N_S^{(r)}$ we can estimate its expected value by calculating its mean.

$$\widehat{\mathbb{E}[N_S^{(r)}]} = \frac{1}{R} \sum_{r=1}^R N_S^{(r)} \quad (14)$$

Now we normalize every $N_S^{(r)}$ to see how many percent each deviates from it's mean.

$$T_S^{(r)} = \frac{N_S^{(r)} - \widehat{\mathbb{E}[N_S^{(r)}]}}{\widehat{\mathbb{E}[N_S^{(r)}]}} \quad (15)$$

A large value of $T_S^{(r)}$ means that the particular synthetic sequence in question displayed a more cyclical pattern of winners than we would expect from the data generating process. A value close to zero would mean that its cyclicity is around average and a negative value means that its pattern is

³⁵Thus, our first block would be a sub-sequence of the original sequence of the form $b_{I_1, L_1} = \{W_t\}_{t=7}^{10}$.

³⁶This circular design is made to eliminate edge effects which are a drawback in more simplified block bootstrapping procedures. Edge effects arise in similar types of block bootstrapping methods from the fact that without a circular design, observations in the beginning and end of the sequence gets under-sampled.

³⁷The symbol $\cup$ denotes concatenation which is the "joining" of two sequences.

³⁸If the last block makes the total count of observations in the synthetic sequence supersede T , the last sampled observations within the last block are removed. This is made so that the length of the synthetic sequences match the original in length.

very "un-cyclical". We calculate the same normalized value for the original sequence observed in the data which becomes

$$T_S^{obs} = \frac{N_S - \widehat{\mathbb{E}[N_S^{(r)}]}}{\widehat{\mathbb{E}[N_S^{(r)}]}} \quad (16)$$

And formulate our null and alternative hypothesis.

$$H_0 : N_S - \mathbb{E}[N_S] = 0$$

$$H_a : N_S - \mathbb{E}[N_S] \geq 0$$

To test this hypothesis we compare our normalized mean from our original sequence T_S^{obs} to those from the simulated sequences $\{T_S^{(1)}, T_S^{(2)}, \dots, T_S^{(R)}\}$. The p-value of the test is calculated as

$$\text{p-value} = \frac{\#T_S^{(r)} \geq T_S^{obs}}{R} \quad (17)$$

This calculation compares our normalized score N_S to the distribution of simulated scores. If the score N_S is larger than most of the simulated scores $N_S^{(r)}$ then we have evidence for bid-rotation. The hypothesis test is one sided since I am only interested in finding segments which have significant occurrences of patterns closely resembling bid rotation.³⁹

The Stationary block bootstrapping provides us with a way to make inferences about bid-rotation patterns for the sub-market segments M_T . It also handles a different source of noise than the fuzzy scoring does. While the "fuzzy" scoring helps us to deal with "local noise"⁴⁰, Stationary block bootstrapping helps us deal with time dependencies in the data. Below I describe the problem and how the algorithm solves it.

Spells of cyclicity in a sequence could emerge from other things in the data generating process other than chance. Consider three firms competing fairly in a market. Let's say that Firm 1 wins the bid in period one by quite a large margin. They might consider making a higher bid in period 2 because they now have reason to believe that they can get away with winning the auction and setting a higher price. Firm 2 on the other hand might make the conclusion that they need to make a lower bid to win the market since Firm 1 went so low. The winner of the auction could switch in period two with Firm 1 making a higher bid and Firm 2 a lower. This could again result in the two firms concluding that they over/under shot their bids leading to yet another alteration of the winner in period three. This pattern could repeat for several periods.⁴¹

Still, if the cyclicity of the winners sequence is protruding, meaning a high $N_{\bar{\omega}}$ ⁴², we should still want to consider it as strong evidence. If a cyclical pattern repeats for just some months, It could be the result of chance, as my example above shows. The longer the pattern continues, the less credence

³⁹A two sided test would also find segments which had patterns which are very far away from resembling bid rotation.

⁴⁰The motivation for the scoring function was a very bid-rotation-like sequence getting a low score due to it being "off" in one place. Such a small deviation in just one month from a generally bid-rotation looking pattern is what I mean by "local noise".

⁴¹Commentator Mikael Elinder pointed this to me during a presentation of an early draft to this thesis.

⁴²In the following examples, I use $N_{\bar{\omega}}$ instead of N_S to describe how much a pattern resembles bid-rotation. This is to make the examples more readable.

we should place on the belief that this is the result of something other than price coordination. Thus, a very long spell of cyclicalality must still imply bid-rotation. I'll provide three example sequences to show this reasoning more clearly.

$$\begin{aligned}\{W_t\}^{(a)} &= (2, 1, 3, 1, 3, 1, 3, 1, 3, 1, 3, 1) \\ \{W_t\}^{(b)} &= (2, 3, 1, 3, 1, 2, 3, 2, 1, 1, 2, 3) \\ \{W_t\}^{(c)} &= (3, 2, 2, 1, 3, 1, 3, 1, 3, 1, 2, 1)\end{aligned}$$

Sequence $\{W_t\}^{(a)}$ shows a strong cyclical pattern with $N_{\bar{\omega}} = 8$ out of 9 possible matches. To me, it seems unreasonable to consider such a pattern the result of chance. Sequence $\{W_t\}^{(b)}$ is our original sequence and only has $N_{\bar{\omega}} = 1$ ⁴³ and is not indicative of any bid-rotation. Sequence $\{W_t\}^{(c)}$ has $N_{\bar{\omega}} = 4$. Here the indication is not so clear. The sequence starts off with no discernible pattern, then begins a spell of cyclicalality which seems to fizzle out at the end of the sequence. Such a pattern could be the result of two firms colluding by bid-rotation for some time. Or it could be the result of the process described in the example above (or something else for that matter).

An important feature of the Stationary block bootstrapping method is that it preserves time dependencies in the data. This means that if the underlying data generating process indeed produces some spells of cyclicalality, the method will account for this. What the algorithm does basically is that it resamples blocks of the original sequence. If a spell of cyclicalality is rather short, it will be resampled many times and thus the original sequence will not be considered unusually cyclical. The longer a cyclicalality spell is, the less likely it is to be re-sampled altogether. This means that as the length of a cyclical pattern of winners increases, it will be considered more and more unusual, meaning it will provide stronger evidence for bid-rotation.

4.1.3 Parametrization and robustness checks

For the overlapping permutations test, there are some parameters which are chosen arbitrarily by me as a researcher. They are; the length of partitions k , the number of samples in the Stationary block bootstrapping algorithm R , the penalty parameter for the scoring function p and the parameter governing the geometric distribution in the Stationary block bootstrapping algorithm g . Below I provide a discussion of my handling of each of these parameters.

The choice of length of the partitions $k = 2n$ I have already made an argument for in section 4.1.1. In short, I believe $k = 2n$ strikes a balance in taking the right temporal dependencies into account.

The number of bootstrap samples is set to $R = 10\,000$. The estimation of $\widehat{\mathbb{E}[N_S^{(r)}]}$ gets better with more samples but the effect diminishes as R increases. Also, I have to consider that the bootstrapping algorithm will run many times. I deem that $R = 10\,000$ strikes a reasonable balance between accuracy and computational feasibility.

For the penalty parameter for the scoring function p I have provided examples and argumentation for my choice of $p = 2$. Nevertheless I will carry out the Overlapping permutations test for values of $p = 1$, $p = 2$ and $p = 3$ and compare the results. This serves as the first part of my robustness check for the test.

⁴³This sequence was generated at random so the result stems from mere chance.

For the parameter g governing the geometrical distribution, I have chosen the value $g = \frac{1}{6}$. The expected value of this parameter is the average length of each block in the bootstrap sample. The expected value of the variable L_i is $\mathbb{E}[L_i] = \frac{1}{p}$ which is equal to 6 when $p = \frac{1}{6}$. In practice, this means that the average length of the blocks in the bootstrapped sample is six months. A large value of g will result in very short blocks and a small value will give very long blocks. Too long blocks run the risk of resulting in too little variation in the re-sampled sequences.

A balance needs to be struck so that time dependencies are preserved while we get sufficient variation in the bootstrapped samples. I believe $g = \frac{1}{6}$ does this. Still, I will carry out the test with g equal to $\frac{1}{3}$, $\frac{1}{6}$, $\frac{1}{9}$ and $\frac{1}{12}$.⁴⁴ This will be the second part of my robustness check. Taken together, I will run the test for a total of twelve times for each combination of values of p and g . For each test I set the p-value to 0.05. This is to investigate the behavior of the test under different parameterizations.

4.1.4 Choice of p-value and output of the test

The choice of p-value comes down to the trade off between Type I errors and Type II errors. Since I conduct the test on each of the 11 843 sub-market segments, we could happen to flag a market where no bid-rotation takes place by mere chance. This would constitute a type II error which motivates the use of a lower p-value.

On the other hand, the fact that I'm only flagging entire segments and not parts of segments makes the test conservative. One segment could contain a true spell of bid-rotation but not get flagged. This could happen if there are other time-periods with the segment that displays a non-rotational pattern. Thus the test might fail to detect such segments which would constitute a Type I error which motivates the use of a higher p-value.

To handle this I will consider three p-value for the test, 10%, 5% and 1%. The outcome of each test will then be compared to the results of the tests carried out by Cletus 2016 and Granlund and Rudholm 2023 and see if the fraction of identified sub-markets is in the same ballpark. A reasonably chosen p-value should not make my test identify many more, or far fewer sub-markets than these authors do.

There will therefore be three outcomes of the test considered for further analysis. Each of these outcomes will be the result of the different p-values. The test for each chosen p-value will produce a list of flagged sub-market segments. This list will be used to construct samples used for the pre/post analysis. Both the fixed-effects panel regression and the event-study will be estimated on each of the samples. For each sample I also take a sub-sample. It will be defined by the flagged segments for which the following segment had an increase in the number of firms.

4.2 Pre/post analysis

With the pre/post analysis, the basic idea is to compare the price level and price dispersion in segments where bid-rotation occurs to the segment following it. This section relates directly to the testing of predictions **P2** and **P3**. The theory presented in section 2.3 and 2.4 predicts that prices will be higher under collusion and price variance lower. I will therefore develop two models to compare prices and prices dispersion under and after tacit collusion (which is signaled by bid-rotation).

⁴⁴This will correspond to having blocks in the bootstrapped samples of average length equal to 3 (a quarter), 6 (six months), 9 (three quarters) and 12 (a years).

This analysis will be descriptive in nature. The overlapping permutations test identifies breaks in the market competition regime but it doesn't tell us why the regime changed. Because prospective and incumbent firms base their decisions on the very prices, profits, and expectations that arise within the market itself, entry and exit are fundamentally endogenous. Therefore there is no exogenous variation to exploit for a causal analysis. The endogenous reasons for a new firm entering a sub-market and disrupting the bid-rotation pattern preclude any causal analysis of this change.

4.2.1 Fixed-effects panel regression

The overlapping permutations test provides us with a list of sub-market segments flagged for bid-rotation. If a sub-market segment is flagged for bid-rotation we denote it M_i^* . The sub-market segment following it is denoted by M_i^{+*} . We start by restricting our sample so that it only contains these two types of sub-market segments. Thus each sub-market left in our sample contains a pair of sub-market segments (M_i^*, M_i^{+*}) . We can then define an indicator variable D_i as

$$D_i = \begin{cases} 1, & \text{if } M_i = M_{T,i}^{+*} \\ 0, & \text{if } M_i = M_i^* \end{cases} \quad (18)$$

The outcome variable for price is defined as

$$P_i^{norm} = \frac{\mu_i^*(P) - \mu_i^{+*}(P)}{\sigma_i^*(P)} \quad (19)$$

where $\mu_i^*(P)$ is the mean price in the segment flagged for bid-rotation, $\mu_i^{+*}(P)$ is the mean price in the segment following the bid-rotation segment and $\sigma_i^*(P)$ is the standard deviation of the price in the bid-rotation segment. For each pair of (M_i^*, M_i^{+*}) we normalize the price to the bid-rotation sub-market segment with the variable P_i^{norm} . This variable is measured in number of standard deviations from the mean price under bid-rotation.

The outcome variable to measure price dispersion, the Coefficient of variation, is defined as

$$CV^{norm} = \frac{CV^{+*}}{CV^*} = \frac{\sigma_i^{+*}(P) / \mu_i^{+*}(P)}{\sigma_i^*(P) / \mu_i^*(P)} \quad (20)$$

which is the coefficient of variation in the segment following bid-rotation, divided by the coefficient of variation in the bid-rotation segment. The coefficient of variation (CV) is the ratio of a data sets standard deviation to its mean, usually expressed as a percentage. It gauges relative variability, so a higher CV indicates greater dispersion around the mean. This variable CV^{norm} measures the ratio of the coefficient of variation in the competition segment to that in the collusion segment.

We then formulate the following fixed effects regression model

$$Y_i = \beta D_i + \alpha_i + \varepsilon_i \quad (21)$$

where Y_i is one of the two outcome variables and α_i is market fixed effects. The coefficient β captures the average within-sub-market change in Y_i when moving from bid-rotation ($D_i = 0$) to non-bid-rotation ($D_i = 1$). We estimate this model separately for $Y_i = P_i^{norm}$ and for $Y_i = CV^{norm}$. The standard errors are clustered at the substance level.

A negative β for $Y_i = P_i^{norm}$ implies that, on average, the non-bid-rotation mean price is lower (in bid-rotation-segment standard deviations) than the bid-rotation segment mean price. Conversely, a positive $Y_i = CV^{norm}$ implies that the non-bid-rotation-segment coefficient of variation is larger than under bid-rotation.

This equation will be estimated on three samples produced by the overlapping permutations test. Each sample is the result of a different choice of p-value for the test. The p-values chosen are $p = 0.1$, $p = 0.05$ and $p = 0.01$. For each sample, a sub-sample is created where I only consider bid-rotation segments M_i^* which had a following segment M_i^{+*} in which the number of firms increased. The estimation is then carried out on all of these three sub-samples as well.⁴⁵

The fixed-effects specification compares mean prices, or coefficients of variation, before and after collusion within each sub-market, netting out time-invariant heterogeneity. However, because it treats the transition to competition as a single binary shift, it fails to control for any underlying time-varying trends in prices that might coincide with, or predate the end of collusion.

Despite this limitation, the raw fixed-effects comparison has a practical value as a transparent first pass. By imposing minimal structure, it provides a clear summary of the average withinmarket shift in prices immediately surrounding the collusion termination. The fixed-effects test offers a parsimonious measure of what happened without over-relying on identifying assumptions about how prices evolve in time.

4.2.2 Event-study design

The event study design has a similar setup as the fixed-effects panel regression. The overlapping permutations test provides us with a list of sub-market segments flagged for bid-rotation. If a sub-market segment is flagged for bid-rotation we denote it M_i^* . For the event-study design, we consider *the last month* within that flagged sub-market segment T as a treatment date.⁴⁶ This month is where the bid-rotation segment ended and another, non-bid-rotation segment began.⁴⁷ Treatment is an absorbing state and treatment adoption is specific for each sub-market denoted by E_i . Hence, treatment status changes at most once at the sub-market-specific time E_i , from 0 to 1. Treatment adoption is indicated by the binary variable: $\Delta T_{i,t} = T_{i,t} - T_{i,t-1} = \mathbb{1}[E_i = t]$.

Treatment adoption ($\Delta T_{i,t} = 1$) is the point in time when a sub-market changes its competition regime. Time periods prior to this point in time are the untreated state *conditioned on* that these time periods lie within the bid-rotation segment. A never-treated sub-market is one where no bid-rotation pattern was observed for any sub-market segment.

The effect window is restricted to include leads ($\underline{l}$) and lags ($\bar{l}$) of twelve months (see paragraph 4.2.2 for a discussion of this choice). Thus we defined the *binned treatment adoption indicator* $D_{i,t}^l$ as

$$D_{i,t}^l = \begin{cases} \sum_{s=-\infty}^{\underline{l}} & \text{if } l = \underline{l} \\ \Delta T_{i,t-l} & \text{if } \underline{l} < l < \bar{l} \\ \sum_{s=\bar{l}}^{\infty} & \text{if } l = \bar{l} \end{cases} \quad (22)$$

⁴⁵See section 4.1.3 for a discussion on why these samples and sub-samples were constructed in this way

⁴⁶Remember that all sub-market segments have different number of time periods $t = (1, 2, \dots, T)$

⁴⁷Technically, another bid-rotation could start here. I've checked for this and only found a few such instances. I have thus not bothered by making any sub-analysis when including/excluding this cases. See the results in section 5.1 for more on this topic.

The outcome variables for the event-study thus slightly different compared to the fixed panel regression. The goal of the event-study design is to estimate the dynamic effects of this treatment on price levels and price dispersion. Therefore I define the outcome variable measuring as

$$P_{i,t}^{(norm)} = \frac{P_{i,t} - \mu_i^*(P)}{\sigma_i^*(P)} \quad (23)$$

where the difference to equation (19) is that we now normalize the price variable in the segment and not the mean price. This is so that we can dynamically trace the change in price for each month under bid-rotation and after it. The variable $P_{i,t}^{(norm)}$ is measured in standard deviations of the price under bid-rotation. For the variable measuring dispersion I define

$$CV_{i,t}^{norm} = \frac{CV_{i,t}}{CV_{i,t}^*} = \frac{\sigma_i(P)/\mu_i(P)}{\sigma_i^*(P)/\mu_i^*(P)} \quad (24)$$

which is now the coefficient of variation of any segments following bid-rotation. The definition in equation (20) only considered segments immediately following bid-rotation. Equation (24) and (23) is calculated for all segments following the bid-rotation segment in time. This means that this method will take more of the pre- and post-period into account. The equation I will estimate is the following

$$Y_{i,t} = \sum_{l=\bar{l}}^{\bar{l}} \beta_l D_{i,t}^l + \alpha_i + \theta_t + \mathbf{X}\gamma + \varepsilon_{i,t} \quad (25)$$

where $Y_{i,t}$ is either $P_{i,t}^{(norm)}$ or $CV_{i,t}^{norm}$ as defined above. α_i is market fixed effects and θ_t is time fixed effects at a monthly level in calendar time. $\mathbf{X}$ is a set of potential covariates and $D_{i,t}^l$ is the *binned treatment adoption indicator*. I will cluster the standard errors at the substance level as I believe the error terms to be correlated at this level. As an example, we can imagine random shocks to the price of the chemical ingredients in a specific substance. This will cause correlated errors in all markets that share the same substance. As in the case of the fixed-effects panel regression, estimation will be carried out using the three samples and sub-samples produced by the overlapping permutations test.⁴⁸

Restriction of effect window As explained by Schmidheiny and Sieglöcher (2023), a researcher must impose restrictions on the effect window to implement the event study design. The choice of effect window depends on the context and the research question that the model tries to address. We must consider when the effect can be expected to have fully materialized and what a reasonable pre-treatment period is.

Restriction of the effect window is made based on the assumption that the effect is constant outside the window. Schmidheiny and Sieglöcher (2023) provide a way to use this assumption in order to assess whether the choice of size of the effect window was reasonable. If treatment effects had fully materialized by $\bar{l}$, we would expect estimates leading up to level off and converge to $\hat{\beta}_{\bar{l}}$. A pronounced drop between $\hat{\beta}_{\bar{l}}$ and $\hat{\beta}_{\bar{l}-1}$ is instead an indication that dynamic effects are still unfolding. To investigate this, we can check whether estimates leading up to the endpoint converge toward $\hat{\beta}_{\bar{l}}$ by inspecting the coefficient plots from the event-study.

I believe an effect window of $\pm$ twelve months is appropriate. It will let us trace out the effect for

⁴⁸ Again, see section 4.1.3 for a more detailed discussion.

an entire year, hopefully capturing the full adaptation on prices and price dispersion following the end of bid-rotation schemes. A pre-trend of one year will help us determine if the estimated coefficients are close to zero prior to treatment and if they display a clear jump at the treatment date. To see if this choice was reasonable I will check the coefficient plots for the event-study for converging patterns at the end of the event window, as suggested by (Schmidheiny and Siegloch 2023).

Identification Schmidheiny and Siegloch 2023 demonstrate how the length of the effect window has direct implications for the econometrical identification⁴⁹ of the model. The choice of effect window can alter what groups, with which treatment status, are available when carrying out the estimation. To identify the dynamic treatment effects β_l , we need to observe at least one treated unit for each lag and lead of the effect window. To uniquely identify the secular time trends θ_t , we need at least one control group observation for each time period t .

Thus, the choice of effect window, together with which of the three samples being used (see section 4.1.3) can affect the identification of the model. A small sample could have some lags and/or lead without a treated unit. Or it could miss a control group observation for some time period(s). A further complication to absent never-treated groups is that "... this underidentification can easily be overlooked as many statistical packages automatically drop regressors in the case of multicollinearity."⁵⁰

This is important since the different samples produced by the overlapping permutations test to provide me with samples of decreasing size. Setting the p-value at 0.1 will produce the biggest sample. Tightening the threshold to a p-value of 0.05 will flag less sub-market segments as suspected for bid-rotation. Using a p-value of 0.01 will give me the smallest sample.

5 Results

This section has three parts. In the first, I present the results from the overlapping permutations test using three different p-values presented in section 4.1.4. The flagged segments in the different sub-markets are then used to create samples on which I carry out the fixed-panel regression and the event-study presented in section 4.2. For the fixed-panel regression I present the results using all three samples in the main text. For the event-study analysis I focus on the results achieved by using the sample generated by setting the p-value of the test to 0.1. The reason for doing so is a combination of two things.

First, by comparing the fraction of flagged markets in my sample to that in (Cletus 2016) and (Granlund and Rudholm 2023), there is reason to believe that using a p-value equal to 0.1 is the best choice for my test. Using this p-value, my test flags 18.6% of the sub-markets for bid rotation which places my results between that of Granlund and Rudholm 2023 (42.6% flagged sub-markets) and Cletus 2016 (11.6% flagged sub-markets).⁵¹ Thus I consider the results using this sample to be my main results and will focus my attention here.

Second, when I change the sample used for the fixed-panel regression, the results change. This is true also for the event-study. But the results change in very similar fashion. Thus a thorough

⁴⁹Identification here is understood as the identifiability of the regression coefficients and not the identification of average treatment effects in the population.

⁵⁰(Schmidheiny and Siegloch 2023), section 2.4.2.

⁵¹If one uses my test with a p-value of 0.05 it results in only 6.3% of sub-markets being flagged and a p-value of 0.01 results in 1.5% flagged sub-markets.

exposition of all the results from the event-study doesn't add much to the analysis. The full results for all samples can be found in appendix D.

5.1 Overlapping permutations test

The overlapping permutations test has provided a set of market segments where there are signs of bid-rotation. Each segment following a flagged segment has been checked to see if there were signs of bid rotation following the initially flagged segment. There were 6 such segments found over 5 sub-markets and 4 substances. This check is important for the coming analysis of prices and price dispersion under different market regimes. Prices and price dispersion in segments that are flagged for bid rotation are to be compared to the following segments where no such signs are present. Thus the sub-markets where bid-rotation is present in subsequent segments need to be removed before such an analysis is conducted to enable proper comparisons. To make the presentation more stringent, I have chosen to exclude this small sample from the analysis at this point already.

Removing these sub-markets has one implication for the exposition of the results. Since we now know that we only kept sub-markets which had only one bid-rotation segment, each flagged segment thus corresponds to one flagged sub-market. This is important for comparing the percentage of flagged markets in this test to the percentages in (Cletus 2016) and (Granlund and Rudholm 2023).

Figures 8, 9 and 10 display the results from the overlapping permutations test. The parameter values are set to $g = \frac{1}{6}$ and $p = 2$ and the p-values at 10%, 5% and 1%. Robustness checks for different choices of g and p are presented in section 5.4. The figures show the frequency of segments where one of the five tests has detected signs of bid rotation by number of active firms in the segments. The most common result is that two-firm bid rotation has been detected in market segments with two active firms. The next most common result is two firm bid rotation on markets with three firms active. The third most common result is three firm bid rotation in markets with three active firms. This result holds for all of the three significance levels.

Bid rotation among four and five firms are found to some extent at the 10% level while being very uncommon at the 5% level and completely absent when the significance level is at 1%. Six firm bid rotation is found very rarely across significance levels of 10% and 5% and not at all at the 1% level.

A further breakdown of the test is presented in 4. This table presents if the segment following the bid rotation segment saw an increase, decrease or no change in the number of active firms and how many collusion segments that was the last segment in time. It also reports the total number of flagged segments and the percentage of flagged segments out of the total number of sub-markets in the data. Since sub-markets with repeated flags have been removed from the sample, the percentage of flagged segments is equal to the percentage of sub-markets where bid-rotation is suspected.

In section 4.1.3 the nature of the overlapping permutations test was discussed. Since approximately 12 000 segments have been tested for bid rotation we have the potential problem of mass testing, which warrants the use of a lower p-value. On the other hand, there were other properties that made the test overly conservative, which warrants the use of a higher p-value. To get a sense of which p-value balances the Type I and Type II errors we can compare the outcome of the test to previous literature.⁵²

Both (Cletus 2016) and (Granlund and Rudholm 2023) conducted test of bid rotation in the same context as I do. Out of the 924 sub-markets studied by Cletus 2016, he flags 107 as suspected when

⁵²In this context, a Type I error would be incorrectly flagging a market for bid rotation. A Type II error would be a market that is not flagged when it should be.

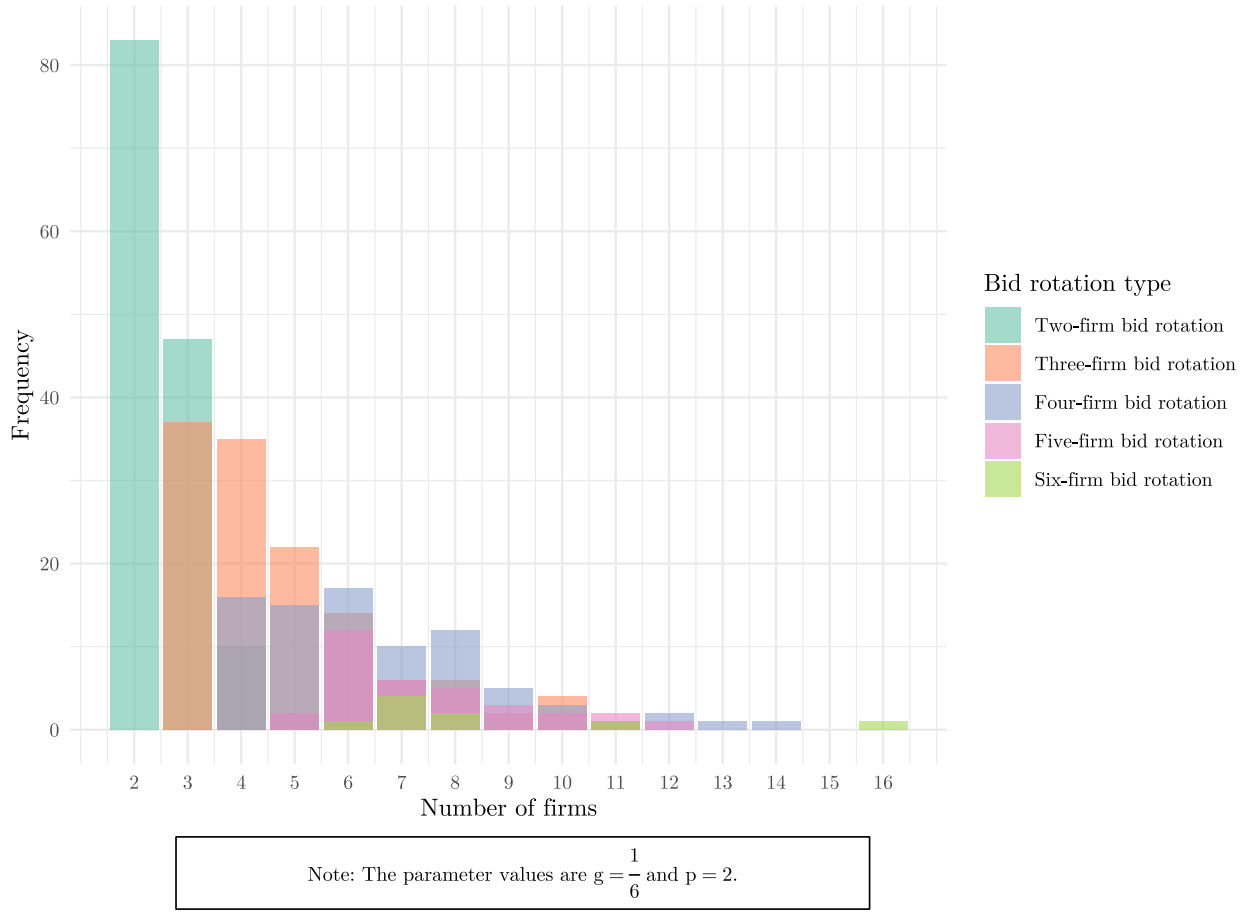


Figure 8: Overlapping permutations test – P-value:10%

checking for patterns of two- and three-firm bid rotation which results in 11.6% of sub-markets flagged for two- or three-firm bid rotation. The study by Granlund and Rudholm 2023 documents 646 distinct exchange-groups where a clear, long-lasting bid-rotation scheme is present. The patterns are for two-three- and four-firm bid rotation and the total number of sub-markets studied is 1 515. This results in 42.6% of sub-markets flagged as suspected of bid rotation.

For the test conducted in this thesis, we get the following results. When using a p-value of 10% the test finds 344 suspected sub-markets which means that 18.6% of markets were flagged. For a p-value of 5% the results are 117 flagged sub-markets which correspond to 6.3%. When the p-value is set to 1%, 28 sub-markets are flagged or 1.5%. The results can be found in table 4 and the number of investigated sub-markets are found in table 1.

5.2 Fixed-effects panel regression

In this section I present the results from the fixed-effects panel regression. The estimation has been carried out on three samples produced by the overlapping permutations test where certain segments have been flagged for bid-rotation. Each sample is the result of a different p-value chosen for the test, namely 10%, 5% and 1%. For each of these samples, a sub-sample has been created where only the sub-markets where the number of firms increased in the segment following the flagged segment are

Table 4: Results overlapping permutations test

P-value = 10%		P-value = 5%		P-value = 1%	
Total flagged	344	Total flagged	117	Total flagged	28
Pct. flagged	18.6%	Pct. flagged	6.3%	Pct. flagged	1.5%
Pct. inc.	42.2%	Pct. inc.	44.4%	Pct. inc.	53.4%
Increase	145	Increase	52	Increase	15
Decrease	96	Decrease	29	Decrease	4
Last segment	103	Last segment	36	Last segment	9

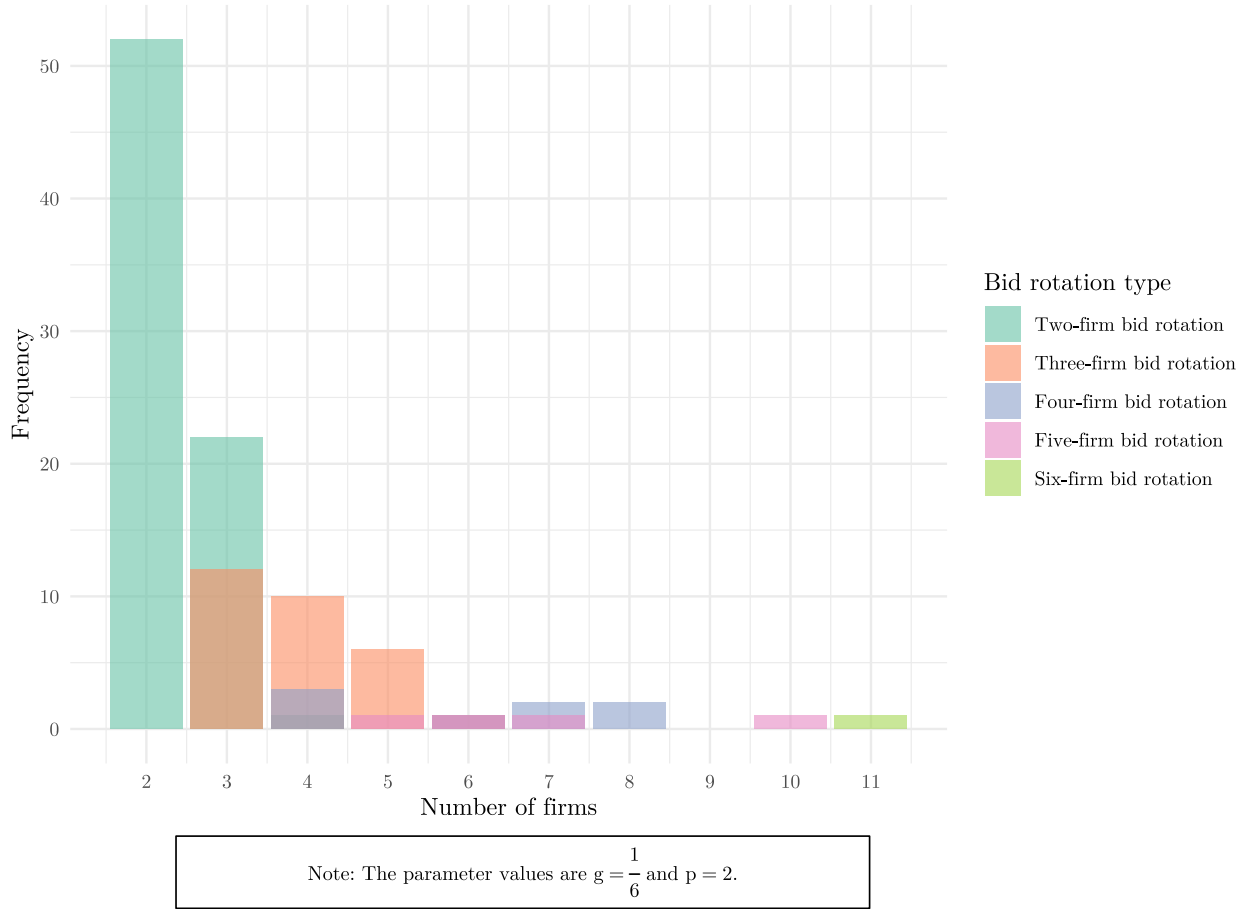


Figure 9: Overlapping permutations test – P-value:5%

included.⁵³ The outcome variables are the normalized price and normalized coefficient of variation that was presented in section 4.2.2. First I present the results from the full samples and then the results from the sub-samples.

Table 5 presents the results from the fixed-effects panel regression for all sub-markets flagged at the 10%, 5% and 1% level. The first column shows that the price in the segments following those flagged for bid-rotation at the 10% is 0.39 standard deviations lower. The result is statistically significant at the 1% level. The second column shows that price dispersion (as measured by the coefficient of

⁵³Table 4 presents how many sub-markets were flagged in total and how many of these had the number of firms increase in the subsequent segment for the different p-values.

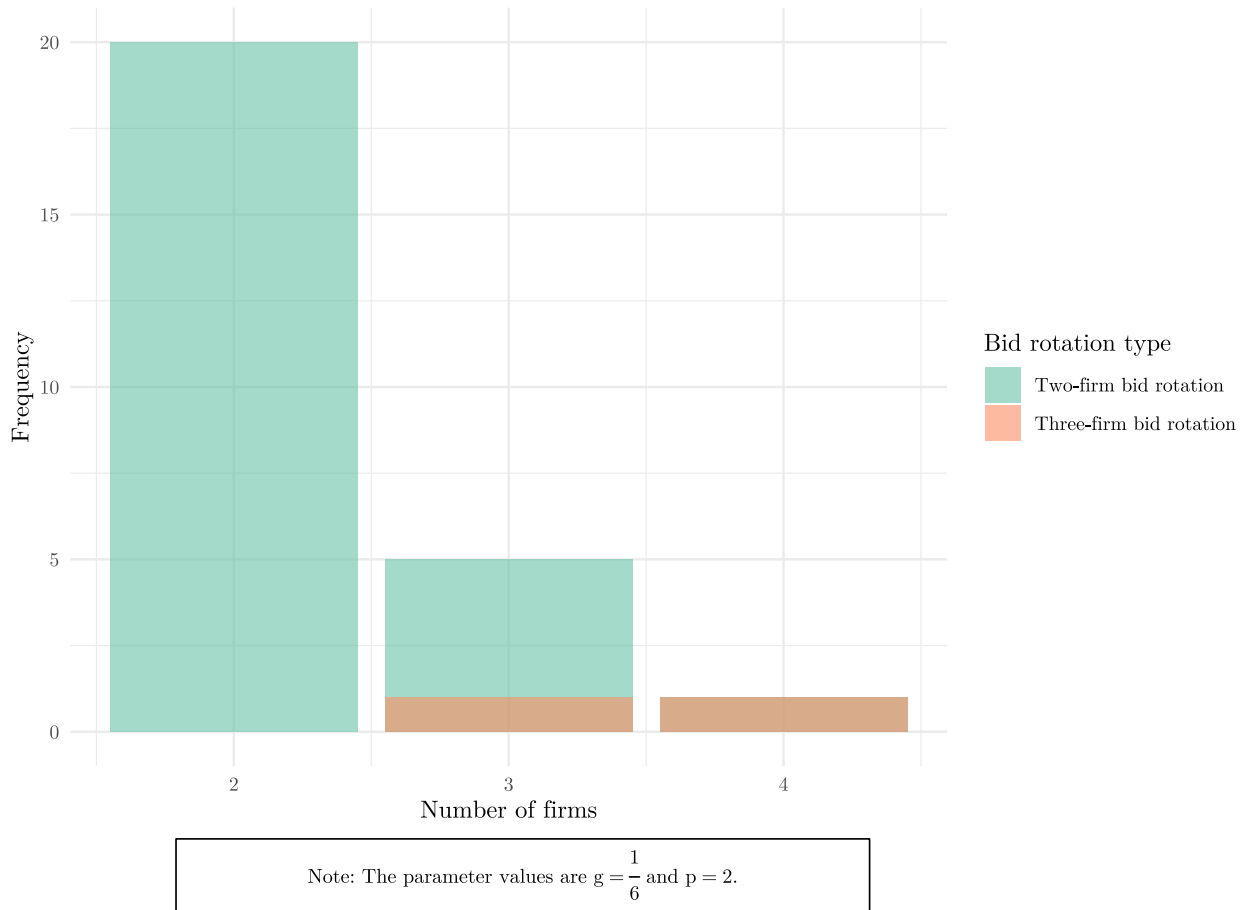


Figure 10: Overlapping permutations test – P-value:1%

variation) in the segments following those flagged for bid-rotation at the 10% is 48 percent higher than in the bid-rotation segments. The result is statistically significant at the 5% level.

Looking at columns three and four we can see the same results but calculated for segments following those flagged for bid-rotation at the 5% level. The point estimate for the price is negative, indicating a lower price following the bid-rotation segments. For the coefficient of variation the point estimate is positive indicating an increase in price dispersion following bid-rotation. The magnitudes are somewhat larger than in columns one and two, while the standard errors are substantially larger and the estimates are not statistically significant.

Columns five and six show the results for segments following those flagged for bid-rotation at the 1% level. The point estimate for the price is negative, with a much larger magnitude than those found in columns one and three. This estimate taken at face value means that the price decreased by 4 standard deviations when bid-rotation ended. The point estimate is relatively large in magnitude compared to its standard error. This indicates some suggestive evidence of an association, but the uncertainty around the estimate remains substantial. The point estimate for the coefficient of variation (column 6) is negative and large in magnitude, indicating a large decrease in price dispersion. The standard error is almost as large as the point estimate and it is thus statistically non-significant.

Lastly, the number of observations used for the estimations of the different models decreases quite substantially when the sample is changed. For the sample based on segments flagged at the 10% level we have almost four times as many observations as for the sample based on segments flagged at the

Table 5: Two period comparison – Full sample

Dep. Var. : Model:	P-value = 10%		P-value = 5%		P-value = 1%	
	P_i^{norm} (1)	CV^{norm} (2)	P_i^{norm} (3)	CV^{norm} (4)	P_i^{norm} (5)	CV^{norm} (6)
<i>Variables</i>						
Competition	-0.3889*** (0.1435)	0.4779** (0.2216)	-0.4843 (0.5173)	0.5441 (0.5932)	-4.064 (2.730)	-6.952 (6.642)
<i>Fixed-effects</i>						
Market	Yes	Yes	Yes	Yes	Yes	Yes
<i>Fit statistics</i>						
Observations	162,145	162,145	40,978	40,978	9,156	9,156
R ²	0.27266	0.19279	0.29414	0.16769	0.42195	0.31337
Within R ²	0.00237	0.00070	0.00089	0.00020	0.02044	0.00986

Clustered (Substance level) standard-errors in parentheses

*Signif. Codes: ***: 0.01, **: 0.05, *: 0.1*

5% level. The number of observations is again diminished by a factor of four when we compare the sample based on segments flagged at the 5% level to that on the 1%. This reflects the fact that the number of flagged sub-markets decreases substantially when a lower p-value is used (see table 4).

Table 6 presents the fixed-panel regression restricted to the sub-sample where market segments following the flagged segment saw an increase in the number of firms. In the first column we see that the normalized price in the periods following a suspected bid-rotation episode in the 10 % group is 0.62 standard deviations lower than in the preceding periods. The estimate is precise (s.e. = 0.224) and statistically significant at the 1 % level. Column 2 shows a positive coefficient of 0.37 for the coefficient of variation, but the associated standard error of 0.33 means the change in price dispersion is statistically indistinguishable from zero.

Tightening the flag to 5 % (columns 3 and 4) reduces the sample size to roughly one quarter of the original. The point estimate for price more than doubles in magnitude to 1.08, suggesting a larger decline, yet the standard error also increases markedly (s.e. = 0.75), rendering the result statistically insignificant. For price dispersion the sign turns negative (0.94) but, with a standard error of 1.26, the estimate likewise remains far from significance.

At the most stringent 1 % threshold (columns 5 and 6) the sample shrinks to just under 7 000 observations. The estimated effect on price rises substantially to 5.16 standard deviations, while its standard error increases to 3.58. The wide confidence interval implies considerable uncertainty and the estimate is not statistically significant. A similar pattern emerges for dispersion, where the point estimate of 8.84 is almost matched by its standard error of 8.66.

By comparing table 5 to table 6 we see that the magnitudes for the price variable increase when we use the sub-sample. This means that when we only consider markets where the number of firms increased after collusion, prices decrease more when bid rotation ends. For the coefficient of variation, the results are mixed. The estimates for p-value 0.1 indicate that price dispersion increases after

Table 6: Two period comparison – Sub-sample

Dep. Var. : Model:	P-val = 10%		P-value =5%		P-value =1%	
	P_i^{norm} (1)	CV^{norm} (2)	P_i^{norm} (3)	CV^{norm} (4)	P_i^{norm} (5)	CV^{norm} (6)
<i>Variables</i>						
competition	-0.6180*** (0.2239)	0.3661 (0.3290)	-1.080 (0.7481)	-0.9424 (1.260)	-5.163 (3.579)	-8.841 (8.663)
<i>Fixed-effects</i>						
Market	Yes	Yes	Yes	Yes	Yes	Yes
<i>Fit statistics</i>						
Observations	87,276	87,276	23,712	23,712	6,901	6,901
R ²	0.33981	0.25273	0.36493	0.23510	0.42439	0.30032
Within R ²	0.00419	0.00028	0.00331	0.00043	0.02594	0.01254

Clustered (Substance level) standard-errors in parentheses

*Signif. Codes: ***: 0.01, **: 0.05, *: 0.1*

bid-rotation, but to a lesser extent when we only look at cases where the number of firms increase. For p-value 0.05, price dispersion increases when bid rotation ends when we allow for any change in the number of firms, but decreases when we only allow increases in the number of firms. For a p-value of 0.01, the dispersion in prices decreases in both cases, but more so when we only look at an increase in the number of firms. Importantly, only the estimates when we use a sample based on a p-value of 0.1 are statistically significant, meaning that we should treat the others with caution.

5.3 Event study

Below I present coefficient plots and tables for the event-study estimation. Treatment date refers to the month when a segment flagged for bid-rotation ended and a non-flagged segment started. As with the two-period test, the event study has been carried out on samples generated by the overlapping permutations test. Three different samples have been generated using p-values of 10%, 5% and 1%. In this section I will focus the presentation on the results from the sample generated by a p-value set to 10%. The other results are found in appendix D. Sub-samples have also been created where I only consider markets where flagged segments were followed by segments which saw an increase in the number of firms. I will first present the results obtained by using the full sample and after that I will present the results using the sub-sample.

In figure 11 we see the coefficient plot for the normalized price. The coefficients for the pre-trend are all statistically insignificant. When the treatment month occurs the price drops by 0.14 standard deviations and continues to decrease for the subsequent months. In the last period the price is almost 1 standard deviation lower than the price under bid-rotation. The uncertainty in the estimates gradually increases as we move further away from the treatment date. Even so, all estimates are statistically significant. There seems to be a quite clear of a convergence of the point estimates for the last three event-months.

In figure 12 we see the coefficient plot for the normalized coefficient of variation. The estimates for the pre-trend are all small and statistically insignificant. At the time of treatment, the estimates make a clear jump indicating an increase of 100 percent in the coefficient of variation. As we move away from the treatment date, the estimates increase along with the standard errors. The estimates remain statistically significant for the entire post-treatment period. The point estimates increase slightly in the last 3 event-months.

In columns (1) and (2) in table 7 we can read off the exact results for the coefficient estimates along with their standard errors. Here we see that the point estimates in the pre-treatment period are close to zero with small standard errors. For the price variable, the point estimates following treatment range from (-0.1376) to (-0.9593) . The point estimate of fixed-effects panel regression of (-0.3889) thus falls within this range. The range for the point estimates for the coefficient of variation is (1.007) to (2.510) which places the the point estimate of fixed-effects panel regression of (0.4779) outside of this range.

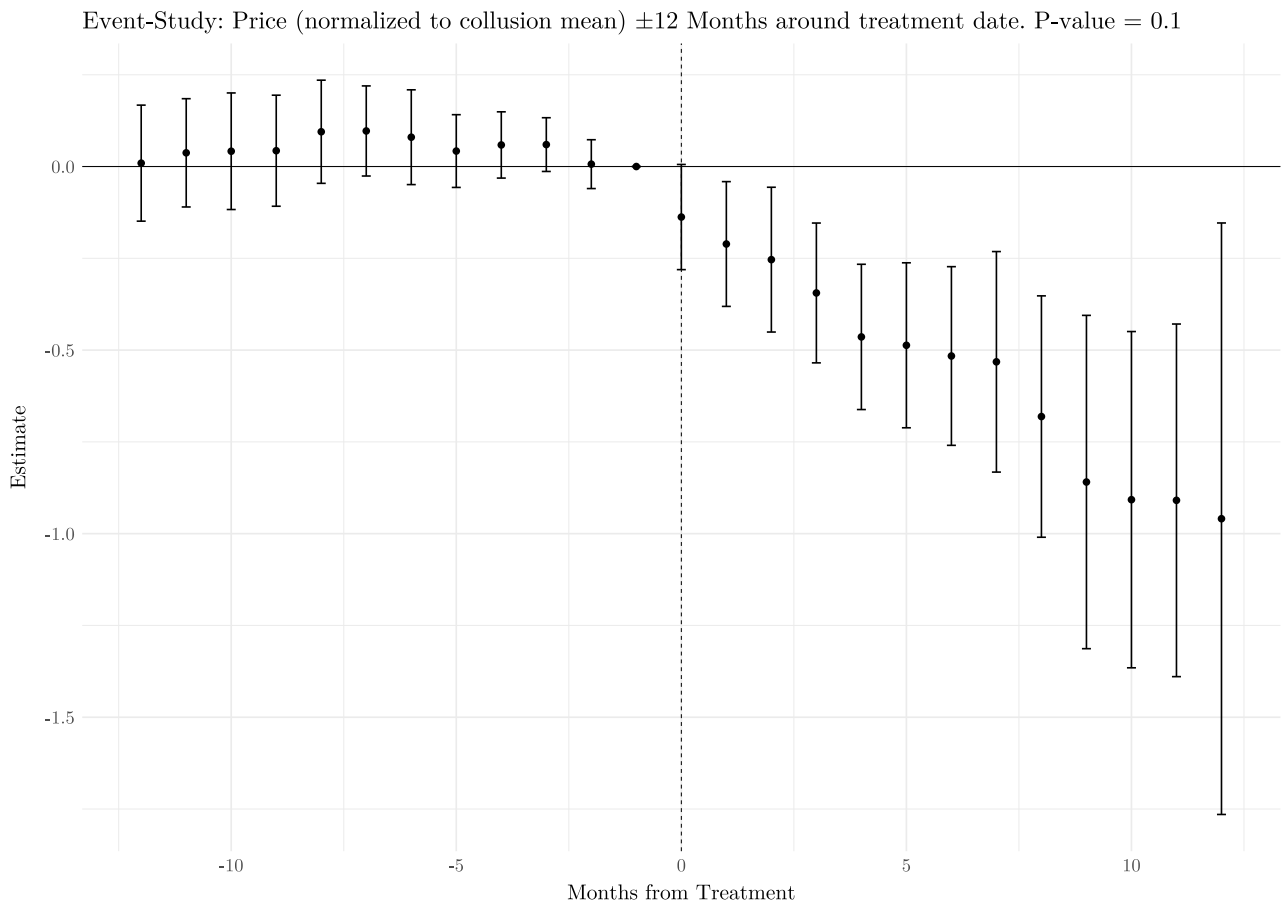


Figure 11: Event study coefficient plot – P-value 10%

Event-Study results: CoV (normalized to collusion CoV): ± 12 Years around treatment date. P-value = 0.1.

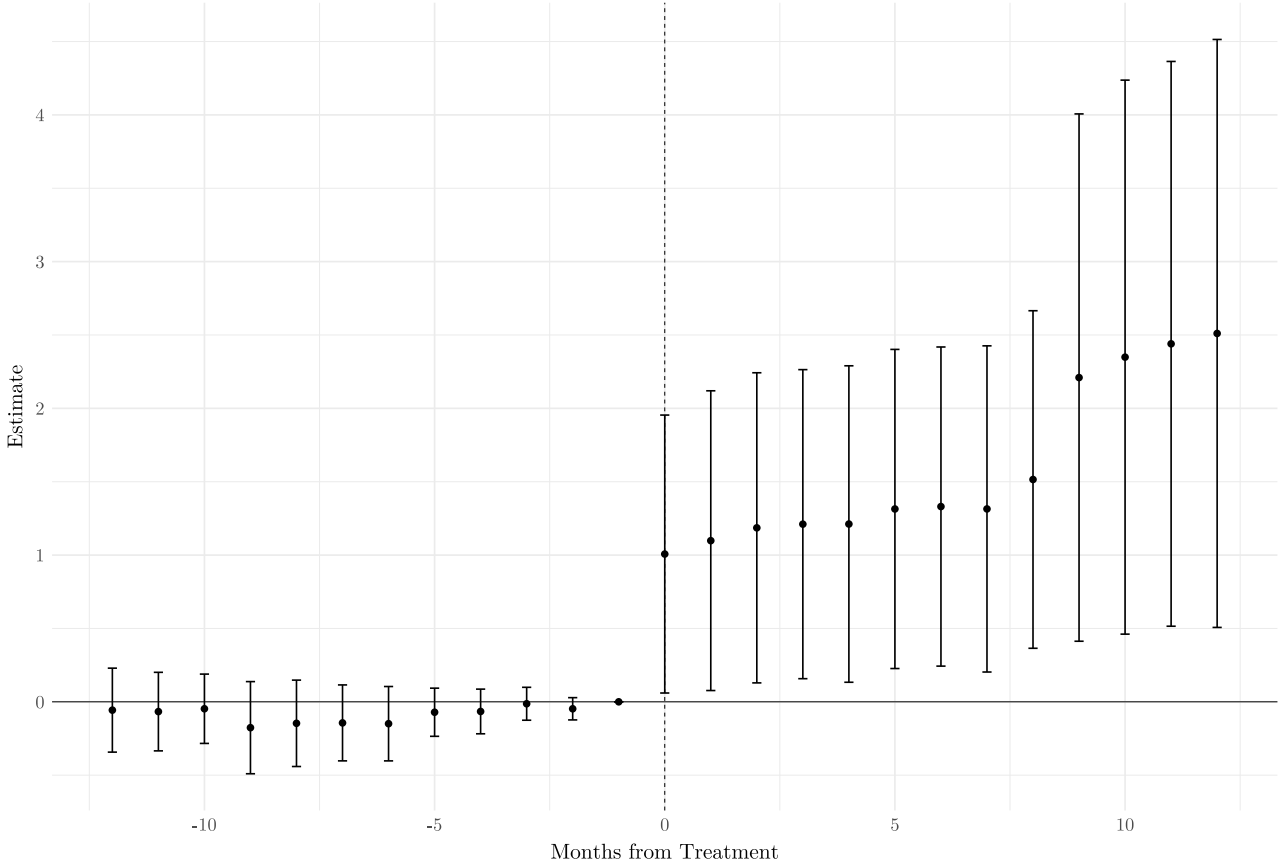


Figure 12: Event study coefficient plot – P-value 10%

In figure 13 we see the coefficient plot for the normalized price when the analysis is restricted to those market segments where the number of active firms increases once the segment flagged for bid-rotation ends. The coefficients during the twelve months leading up to treatment hover very close to zero, ranging from (-0.04) to (0.14) , and none of them are statistically significant.

At the treatment month, the point estimate drops sharply by (-0.47) , and the decline deepens steadily over the subsequent months. Standard errors are modest in the immediate post-treatment window (≈ 0.11 - 0.25), so all point estimates up to event time = 8 are highly significant at the 1% level. From month 9 onward the standard errors widen (≈ 0.41 - 0.70), yet the estimates remain statistically significant – typically at the 1 % level. After month 8 the series appears to level off: the successive estimates cluster between -1.56 and -1.66 .

Figure 14 display the coefficient plot for the normalized coefficient of variation. Again the pattern is similar to that produced when using the full sample. The plot shows a pre-trend close to zero and a clear jump at the treatment date. As for the results on prices, the magnitudes of the estimates are larger when using the restricted sample. The initial jump at the treatment date indicates an increase in price dispersion of almost 250%. The estimated price dispersion then increases for the whole post-treatment period with the last estimate in the time period indicating almost a 500% increase.

Event-Study: Price (normalized to collusion mean) ± 12 Months around treatment date. P-value = 0.1.
Only segments where number of firms increase after collusion

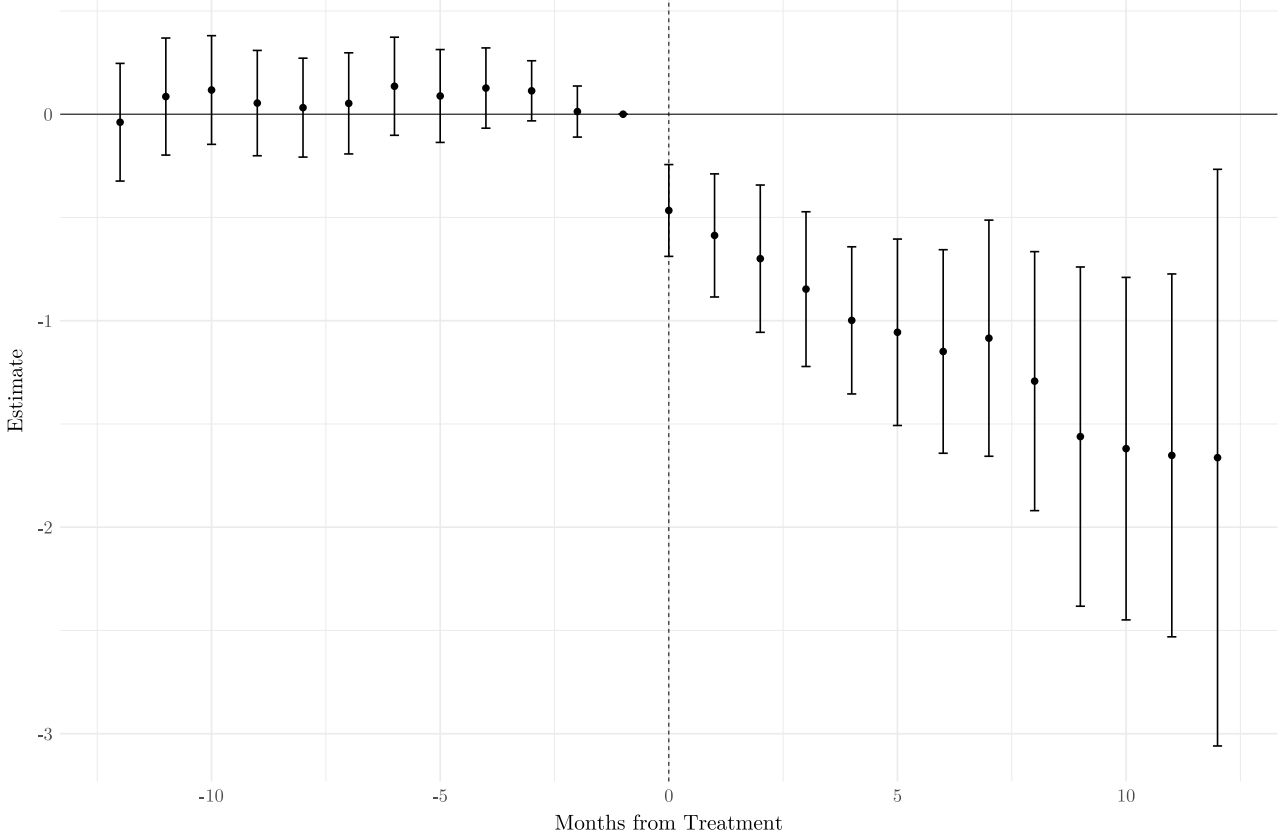


Figure 13: Event study coefficient plot – P-value 10%

Turning our attention to table 7 we can compare the results of the event-study estimation to that of the fixed-panel regression when using the restricted sample. For the price variable, we get an estimate of (-0.6180) from the fixed-panel regression which lies within the range of (-0.4662) to (-1.663) from the event-study. For the coefficient of variation, the point estimate from the fixed-panel regression lies outside the range of the event study. The estimate from the fixed-panel regression is also statistically insignificant, while all the point estimates from the event-study all are statistically significant at the 5% level.

By further studying the results for the fixed-panel regression (tables 5 and 6) and those for the event-study (table 7), we can compare the effect of changing the sample. For prices, the effect is similar: the point estimates increase in magnitude for both methods and remain statistically significant. For the coefficient of variation, the effect is different depending on the method. When changing to the sub-sample, the fixed-panel regression produces an estimate that is smaller in magnitude and statistically insignificant. For the event-study, the change of sample results in larger point estimates across the whole post-treatment period. Statistical significance remains unchanged when using the sub-sample.

So far, I have only presented the results produced by using the sample generated by using a p-value of 010% for the overlapping permutations test. However, all models have been estimated on the two other samples as well (p-values of 5% and 1%) with further estimation carried out on the sub-samples described in table 4. I will now turn to comparing these results across different choices of p-values and the different sub-samples. The tables and coefficient plots presenting these results have been left out of the main text and can be found in appendix D.

Event-Study results: CoV (normalized to collusion CoV): ± 12 Years around treatment date. P-value = 0.1.
Only segments where number of firms increase after collusion

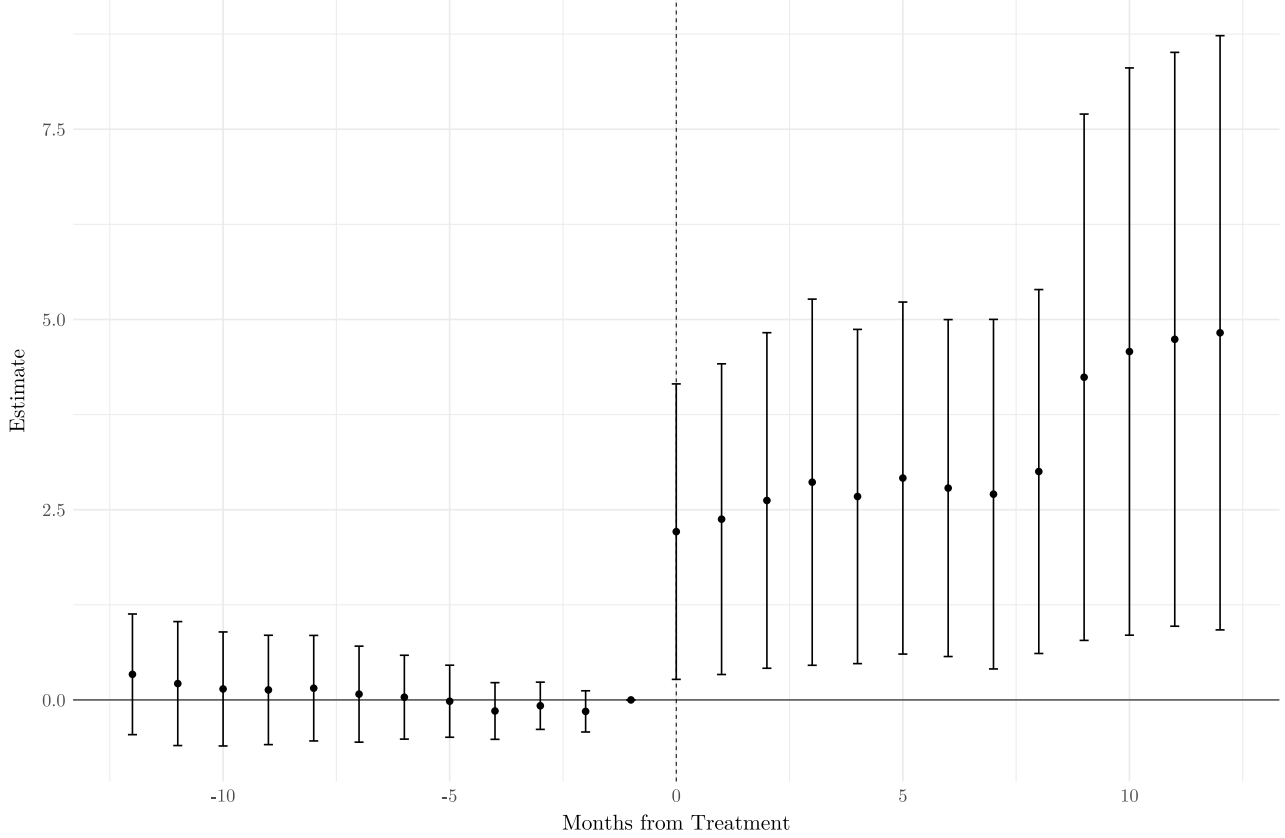


Figure 14: Event study coefficient plot – P-value 10%

When I change the sample from a p-value of 10% to 5% the magnitude increases for the point estimates for both variables. Looking at prices, the uncertainty in the estimate increases substantially and the point estimates become statistically insignificant. Still, the trend of decreasing prices for the post-period remains. The estimates turn negative at the treatment date, but weakly so, and thus no clear jump is visible. When the sub-sample is used, magnitudes increase further and a jump at the treatment date becomes discernible, although the point estimate remains statistically insignificant. All point estimates are more uncertain which is visible by the widened confidence bands.

For the coefficient of variation, the jump at the treatment date is much clearer. All point estimates for the post-treatment period are statistically significant at least at the 10% level. When the sub-sample is used these patterns are amplified. Point estimates increase further, the jump becomes even clearer and the significance level improves.

When I change the sample from a p-value of 5% to 1%, the point estimates for the price variable become very noisy, for both the pre-trend and treatment period. There is some semblance of a more stable pattern around the zero line in the pre-trend and a negative trend in the treatment period. Still, uncertainty is very large which speaks against strong interpretations. When the sub-sample is used, a similar noisy pattern emerge. However, the magnitudes of the point estimates have more than doubled, along with the standard errors. This again speaks against any strong interpretation.

For the coefficient of variation, the patterns are clearer. The pre-trend hovers closely to the zero line and there is a clear jump at the treatment time. Some point estimates in the treatment period are close to being significant in the treatment period. When we use the sub-sample instead, point estimates increase and confidence bands widen. The shape of the point estimates in the pre-trend

and post treatment period are roughly unchanged. A stable pre-trend with point estimates closer to zero is followed by a clear jump at the treatment date. Uncertainty is lower when compared to the estimation of the price variable but still high enough to caution against strong interpretations.

5.4 Robustness checks

The overlapping permutations test has been carried out for a number of different choices for the parameters p and g . For each test, the p-value was set to 5% and the number of resamples to $R = 10\,000$. The plots for the robustness checks are found in appendix C.

From the plots two clear results emerge. First of all, the geometric parameter used in the Stationary block bootstrapping method has a large impact on the test. Setting this value to $g = \frac{1}{3}$ makes the test flag the most segments across all choices for the penalty parameter p . Decreasing this parameter value decreases the number of flagged segments.

A large parameter value means that the sampled blocks that make up the synthetic sequence of winners are short on average. The average length of blocks is equal to the reciprocal of the parameter g . This means that setting $g = \frac{1}{3}$ produces blocks with average length of three months. This will make it so that rather short collusion spells in an original sequence of winners – those that we are testing – will have a lower probability of being re-sampled. This will cause the original sequence of winners to stand out more in relationship to the re-sampled sequences. Thus, segments will have a greater chance of being flagged for bid-rotation when g is large. Conversely, setting g to a lower value, say $\frac{1}{12}$, produces blocks with average length of 12 months. Across 10 000 re-samples, the chance of reproducing a spell of bid-rotation increases in such a case. The shorter the spell, the more times it will be re-sampled. Thus the test becomes more stringent when g decreases.

Second, there is a clear pattern finding bid-rotation more often in markets with fewer active firms across all parametrization tested. This is a good sign since it means that the results align with the expectations from theory, regardless of parameter specification. We expect to find collusion, and thus bid-rotation, in markets with fewer firms. This is clear from equation (8) in section 2.3.

The parameter p seems to have the smallest impact on the test results. The overall pattern when comparing figures 16, 17 and 18 is similar. The main difference is that the number of flagged sub-markets is somewhat lesser when increasing the parameter p . This is also in line with the construction of the test.

The parameter p was introduced to allow handling noise in the sequence of winners. The algorithm calculates the Hamming distance between the investigated bid-rotation pattern and the sequence it is being applied to – synthetic or real – and then applies a non-linear "penalty" based on this difference. The idea is to regard patterns with high cyclicity in winners – albeit with some deviations from a perfect match to the investigated pattern – as giving evidence of bid-rotation. Since the algorithm is applied to both the real and the synthetic sequence(s), it will penalize both from being off. This penalization will be greater for the synthetic sequences since the re-sampling method introduces more noise. Thus synthetic sequences will get lower scores as p increases, and as a consequence, the real sequence will be considered more cyclical in comparison. This results in a more frequent rejection of the null hypothesis which is why the test flags more markets when p increases.

Taken together, the robustness checks show that the test behaves as expected. The changes to the parameters g and p make the test change its output in ways that align with its construction and what we expect from the data generating process.

Table 7: Event-Study results

Dependent Variables: Model:	Full sample		Sub-sample	
	aip_norm (1)	Cov_norm (2)	aip_norm (3)	Cov_norm (4)
<i>Variables</i>				
event_time = -12	0.0092 (0.0800)	-0.0568 (0.1449)	-0.0389 (0.1435)	0.3364 (0.3991)
event_time = -11	0.0373 (0.0747)	-0.0667 (0.1356)	0.0854 (0.1426)	0.2149 (0.4101)
event_time = -10	0.0417 (0.0804)	-0.0473 (0.1196)	0.1171 (0.1325)	0.1443 (0.3771)
event_time = -9	0.0431 (0.0766)	-0.1763 (0.1590)	0.0538 (0.1284)	0.1314 (0.3618)
event_time = -8	0.0947 (0.0712)	-0.1465 (0.1491)	0.0318 (0.1206)	0.1544 (0.3491)
event_time = -7	0.0969 (0.0621)	-0.1437 (0.1311)	0.0524 (0.1233)	0.0761 (0.3176)
event_time = -6	0.0799 (0.0654)	-0.1491 (0.1283)	0.1353 (0.1196)	0.0360 (0.2774)
event_time = -5	0.0422 (0.0502)	-0.0711 (0.0831)	0.0882 (0.1132)	-0.0170 (0.2383)
event_time = -4	0.0588 (0.0457)	-0.0657 (0.0771)	0.1266 (0.0979)	-0.1453 (0.1874)
event_time = -3	0.0598 (0.0371)	-0.0132 (0.0568)	0.1132 (0.0733)	-0.0766 (0.1563)
event_time = -2	0.0065 (0.0337)	-0.0475 (0.0384)	0.0128 (0.0622)	-0.1505 (0.1363)
event_time = 0	-0.1376* (0.0726)	1.007** (0.4797)	-0.4662*** (0.1119)	2.212** (0.9774)
event_time = 1	-0.2110** (0.0861)	1.098** (0.5173)	-0.5871*** (0.1500)	2.376** (1.028)
event_time = 2	-0.2536** (0.0999)	1.186** (0.5352)	-0.6996*** (0.1795)	2.621** (1.110)
event_time = 3	-0.3443*** (0.0964)	1.211** (0.5334)	-0.8471*** (0.1886)	2.861** (1.211)
event_time = 4	-0.4641*** (0.1002)	1.212** (0.5462)	-0.9984*** (0.1794)	2.673** (1.106)
event_time = 5	-0.4869*** (0.1138)	1.314** (0.5508)	-1.056*** (0.2272)	2.916** (1.165)
event_time = 6	-0.5161*** (0.1233)	1.331** (0.5508)	-1.149*** (0.2481)	2.784** (1.115)
event_time = 7	-0.5319*** (0.1521)	1.315** (0.5630)	-1.085*** (0.2879)	2.704** (1.156)
event_time = 8	-0.6810*** (0.1665)	1.515** (0.5826)	-1.293*** (0.3157)	3.002** (1.204)
event_time = 9	-0.8593*** (0.2299)	2.210** (0.9101)	-1.561*** (0.4135)	4.241** (1.741)
event_time = 10	-0.9074*** (0.2320)	2.349** (0.9562)	-1.620*** (0.4175)	4.578** (1.876)
event_time = 11	-0.9092*** (0.2432)	2.440** (0.9747)	-1.652*** (0.4423)	4.739** (1.898)
event_time = 12	-0.9593** (0.4080)	2.510** (1.015)	-1.663** (0.7029)	4.825** (1.966)
<i>Fixed-effects</i>				
Market	Yes	Yes	Yes	Yes
Time	Yes	Yes	Yes	Yes
<i>Fit statistics</i>				
Observations	35,906	35,884	16,587	16,574
R ²	0.04464	0.04884	0.05537	0.07427
Within R ²	0.01470	0.01487	0.02662	0.02405

Clustered (substance_id) standard-errors in parentheses.
Signif. Codes: ***, 0.01, **, 0.05, *, 0.1

6 Discussion

The results of this study support the theoretical expectations outlined in the introduction. In particular, the overlapping permutations test confirms Prediction **P1** by showing that bid-rotation is much more frequently detected in markets with a small number of firms. This pattern aligns with theory: sustaining collusion is easier when there are fewer competitors, since the critical discount factor needed to enforce cooperation rises with the number of firms. The data indicate that collusive bid-rotation schemes were identified in about 18.6% of sub-markets at the 10% significance level of the test. These flagged collusive episodes predominantly come from less populated markets, consistent with the notion that limited competition (fewer rivals) facilitates coordination.

Equally telling is how these schemes respond to changes in market structure. The test results reveal that bid-rotation tends to end far more often when a new firm enters the market than when an incumbent firm exits. In roughly 42% of the detected collusive episodes (at $p=0.10$), a collusive rotation collapsed in the period immediately following an increase in the number of active firms (market entry). By contrast, only about 28% of collusion breakdowns coincided with a decrease in the number of firms (a firm exit). This asymmetry suggests that entry by a new competitor is a potent disruptor of collusion, whereas the exit of a firm – which further concentrates the market – is less likely to disturb an ongoing bid-rotation.

Once a rotation falls apart, it almost never reappears. The overlapping permutations analysis found only a handful of cases (six instances across five sub-markets) where a bid-rotation scheme seemed to restart after having previously collapsed. In other words, once collusion is broken, the market typically does not revert back to a coordinated regime during the observed period. This irreversible shift underscores the fragility of bid-rotation: when the conditions supporting collusion (like a stable group of few firms) are upset, the coordinated pricing pattern seldom, if ever, recovers. These findings from the permutation test not only validate the introductions predictions but also highlight the critical role of market entry in undermining collusive arrangements and the durable nature of the regime change once collusion ends.

The fixed-effects panel regression analysis further reinforces these insights by quantifying the impact on prices and price dispersion when bid-rotation ceases. Focusing first on the broad sample of detected collusive episodes (using the 10% significance threshold for the test), the regression results comport well with Predictions **P2** and **P3** of the theoretical model. The end of a collusive period is associated with a significant reduction in prices and a concurrent increase in price dispersion. In the first post-collusion segment, prices are on average about 0.39 standard deviations lower than they were during the collusive segment. This drop is statistically significant (at the 1% level) and reflects the aggressive price undercutting that ensues once firms stop coordinating their bids. At the same time, the dispersion of prices (measured by the coefficient of variation) rises by roughly 0.48 (in normalized units) in the post-collusion regime, a change significant at the 5% level. The higher price variance indicates that prices become more volatile and differentiated after collusion breaks – some firms bid much lower than others – consistent with a return to competitive behavior where each firm strives to win the product-of-the-month-auction every time period.

These empirical outcomes are in line with the theoretical expectation that breaking a bid-rotation

will lower the level of prices (toward more competitive, near-cost levels) and increase price dispersion (as bidding becomes more aggressive less coordinated and cost-shocks translate more fully to market prices). The fact that both effects are observed with conventional significance in the full sample underscores that the collapse of collusion has a tangible pro-competitive impact on the market. Importantly, the magnitude of these effects is even more pronounced in the sub-samples where a new firm enters at the moment collusion ends.

When we restrict the analysis to those collusive episodes that were followed by an increase in the number of firms, the fixed-effects estimates indicate an even larger price decline. In this entrant sub-sample, the price drop after the bid-rotation collapses is around 0.62 standard deviations, substantially bigger than the average effect in the full sample. This pronounced decline (significant at the 1% level in the regression) suggests that the competitive pressure introduced by a new entrant drives prices even lower than a collusion breakdown due only to internal factors.

In other words, entry amplifies the price competition effect: the incoming firm likely bids aggressively to capture market share by winning the auction, forcing incumbents to respond with lower bids themselves once the collusive understanding unravels. The increase in price dispersion in this scenario is also notable. Intuitively, an added competitor would contribute to more varied pricing. Further, the change in regime makes cost shocks transition more fully into market prices. However, the fixed-effects estimates for dispersion in the entry sub-sample are noisier than for prices. The coefficient on post-collusion price dispersion remains positive (indicating a rise in dispersion) but is not statistically significant in the panel regression for the entrant sub-sample. This imprecision is unsurprising given the smaller sample size of these specific events, yet the positive point estimate is directionally consistent with the idea that an entering firm widens the range of prices (for instance, if the entrant undercuts all others, the spread between lowest and higher bids grows). It is also informative to compare these findings with the complementary event-study approach, which, as discussed below, shows a clear and significant jump in dispersion when entry occurs.

Overall, the panel evidence from the sub-sample reinforces the conclusion that firm entry is a catalyst that not only precipitates the end of collusion (as shown in the permutation test) but also intensifies the competitive aftermath of that transition. The primary effect – prices falling and variability rising – is driven by the collapse of the collusive regime itself, but the presence of a new competitor accentuates this effect beyond the baseline change. When we tighten the definition of collusion to more stringent significance levels (requiring p-values of 5% or 1% in the overlapping test), the fixed-effects regression results become less stable and sometimes counterintuitive. This is largely due to the much smaller sample sizes and limited variation available at these stricter thresholds.

Using a 5% cutoff, only 117 sub-markets were flagged for bid-rotation (versus 344 at 10%), and at 1% the number drops to just 28. Consequently, the post-collusion price and dispersion estimates based on these reduced samples come with much larger standard errors and lose statistical significance. For example, at the 5% level, the point estimates still suggest a price decrease and a dispersion increase after collusion (the signs remain in line with theory), but neither effect is significant by conventional criteria. When only the most extreme collusion cases are considered ($p=0.01$), the regression yields a very large estimated price drop – on the order of 4 standard deviations – but this estimate is

so imprecise that we cannot distinguish it from zero impact. In fact, with so few events, one or two idiosyncratic markets can skew the results; the enormous 4σ price effect is likely capturing peculiarities of a tiny subset of observations rather than a generalizable impact. Likewise, the estimated effect on price dispersion in the 1% sample turns negative (suggesting lower dispersion after collusion, which contradicts our expectations) and is statistically indistinguishable from no effect.

Such anomalies underscore the dangers of over-interpreting results from limited data. The drop in sample size is severe – the number of observations in the regression is roughly four times smaller at 5% than at 10%, and smaller by another factor of four at 1% – leaving few sub-markets for a meaningful comparison. Additionally, the few collusive episodes that meet the stricter criteria might not be representative; they could be unusually persistent or involve particular market conditions, leading to estimates that don't generalize well. In sum, when the analysis is confined to only the highest-confidence collusion cases, the regression results become noisy and even counterintuitive due to low power and comparability. This finding highlights a trade-off: a very conservative detection of collusion yields cleaner identification of collusive instances, but it so restricts the sample that detecting the effects of collusion breakdown (the price and dispersion changes) becomes statistically difficult.

The evidence at the 10% level is therefore the clearest and most reliable, whereas the 5% and 1% samples illustrate the limits of inference when only a few markets are available as evidence. The event-study analysis provides a dynamic perspective on these transitions and largely corroborates the regression findings, especially for the main sample of collusion episodes ($p=0.10$).

The event-study graphs track average outcomes in the periods before and after the end of a bid-rotation, effectively treating the breakdown as a treatment event. Consistent with the models predictions, the event studies show little to no change in prices prior to the breakdown, followed by a sharp change at the moment collusion ends. In the price series, there is a flat pre-trend – prices remain at an elevated, steady collusive level in the months leading up to the event – and then a clear downward jump immediately after the bid-rotation collapses. This indicates that nothing else was systematically driving prices down before the collusion ended; the dramatic decline occurs right when the coordinated bidding ceases, reinforcing the interpretation that the collapse of collusion drives the price reduction.

Similarly, the price dispersion shows a stable pattern before the event and then a distinct increase at the point of breakdown (i.e. the first non-collusive period). The coefficient of variation of prices rises as soon as the market shifts to open competition, reflecting that costs transition more fully to the market price and that some firms start bidding much lower than others once they are no longer constrained by the rotating bid agreement. Notably, the confidence bands in the event-study remain reasonably narrow around the break for the 10% sample, and the jump in both price and dispersion is statistically evident (with post-break effects significant at conventional levels in the plotted estimates). These dynamic results accord well with the fixed-panel estimates discussed above: they illustrate that once collusion ends, the market undergoes an abrupt regime change – prices fall and become more variable in the competitive aftermath.

This aligns with theory: the timing and direction of the changes match what a switch from collusion to Bertrand-style competition would predict. Examining the cases with firm entry through the lens

of the event study further highlights the impact of new competitors. When a new firm enters at the time the bid-rotation collapses, the event-study plots show even more pronounced effects. The drop in price at the break tends to be larger in magnitude for entry events than for collusion endings without entry, and the increase in price dispersion is likewise more marked. In the thesis results, using the entry sub-sample (where the number of firms increases), the entire post-collusion trajectory shifts more dramatically than in the full sample, yet the pattern of a stable pre-period and a jump at the event remains. In fact, a comparison of event-study outcomes confirms that introducing a new firm amplifies the competitive response: the post-entry price reduction falls roughly within the range of the fixed-panel estimate (around 0.6σ) and is clearly distinguishable in the event plot, and the dispersion measure rises to a higher plateau after entry, indicating more intense competition for the winners spot.

Importantly, these heightened effects in the entry sub-sample are estimated with no loss of statistical significance in the event-study framework. This contrasts with the panel regression, where limited observations made the dispersion increase harder to detect in the entry sub-sample, but the event study – by pooling information over time around the entry events – still finds a significant uptick in variance. The consistency of significant findings in the event-study for both price and dispersion, even with a narrowed sample, underscores that the qualitative impact of entry is robust: markets that welcome a new entrant experience a sharper break from the collusive regime, with greater price drops and variability, than those where collusion ends without an influx of new competition. These results affirm the theoretical expectation that an outsiders entry undermines collusion and invigorates price competition.

As with the regression approach, however, the event-study results grow less precise when focusing only on the rare, high-confidence collusion cases. Using the 5% or 1% detection thresholds means there are very few breakdown events to analyze, which weakens the clarity of the event-study patterns. With so few events, the confidence intervals in the event plots widen considerably, and the jump at the collusion breakdown becomes harder to distinguish from noise.

Moreover, the statistical identification⁵⁴ of the coefficients might be compromised by the lack of comparability in the smaller samples. In the empirical strategy section, it was noted that identification (in the statistical sense) can be limited for smaller samples within an event-study framework. For instance, if only a handful of collusion collapses are considered within a give treatment window, some lags and/or lead could be without a treated unit or could miss a control group observation for some time period(s) leading to underidentification. With a tiny sample, these effects can blur the estimated effect. Indeed, in the event-study for the 1% sample, the price and dispersion trajectories exhibit considerable volatility, and any pattern is obscured by wide uncertainty bands (mirroring the regressions lack of significance).

Thus, while the qualitative trend – prices trending flat then dropping post-break – is generally consistent even in these stringent samples, we cannot draw strong conclusions at the 5% or 1% levels due to identification constraints and low power. The main takeaway is that the event-study evidence is compelling at the 10% level (and for sub-cases like entry events), but becomes increasingly ambiguous as the sample of events shrinks, reinforcing the earlier point that one needs enough observations to

⁵⁴In this context, identification is understood as the identifiability of the regression coefficients and not the identification of average treatment effects in the population.

reliably assess the competitive effects of collusions end.

In light of these results, we return to the question posed in the title and the central research question: *Does market entry end bid-rotation?* and: *How does firm entry (or exit) affect price competition dynamics?* The combined evidence from the overlapping test, regression, and event study yields a coherent answer. Yes, market entry is a strong force in ending bid-rotation and spurring competition, but the fundamental driver of the improved market outcomes is the collapse of the collusive regime itself. Two key points emerge. **(a)** The primary cause of the observed price reductions and heightened price dispersion is the termination of collusion – the shift from a coordinated bidding rotation to a competitive regime. When bid-rotation ends (for whatever reason), prices tend to gravitate downward and the uniformity of bids breaks apart (price dispersion), resolving the paradox of persistently high prices that existed under collusion. **(b)** New firm entry serves as a critical mechanism that can trigger and amplify this shift: an entrant not only often precipitates the break in the rotation but also intensifies the ensuing price war, leading to even lower prices and greater dispersion than would occur from a collusion breakdown alone. In practical terms, this means that policies fostering entry (e.g. reducing barriers for generic drug suppliers in the Swedish auctions) can be highly effective in destabilizing collusive patterns and enhancing competition.

However, it is also evident that entry is not the only path to a regime change. Collusion can unravel due to other disruptions, and once it does, the market will likely experience the same competitive outcomes (lower prices, greater dispersion) as observed here. In the cases analyzed, even when looking at the broader sample, the end of bid-rotation still brought prices down significantly. Thus, while encouraging market entry is an important tool to break collusion, the broader conclusion is that any collapse of a collusive arrangement can restore competitive dynamics. Market entry should be viewed as a particularly powerful catalyst of this process, but the results remind us that the underlying regime change from collusion to competition is the fundamental factor driving price competition. Ultimately, this study highlights that changing market structure – especially through entry – is a potent antidote to collusive stability, and that once a collusive regime is broken, the market tends to remain competitive, yielding substantial benefits in terms of lower prices and more aggressive price competition going forward. **Empty**

7 Conclusion

In conclusion, the evidence from Swedens generic substitution market indicates that market entry can indeed unravel bid-rotation schemes and reintroduce vigorous price competition. The analysis shows that collusive price rotations were most prevalent when only two or three firms were active, allowing incumbents to take turns winning bids at elevated prices. Consistent with theoretical expectations, these arrangements often collapsed once a new competitor entered. When collusion broke especially due to entry prices dropped markedly and the uniformity of bids gave way to a wider dispersion of prices. In other words, the end of bid-rotation brought prices closer to competitive levels while increasing variance, effectively resolving the paradox of persistently high prices under a nominally competitive auction system.

These findings align closely with the study's initial predictions: collusion was indeed more frequent in concentrated markets and its collapse led to significantly lower average prices accompanied by greater price variability. Market entry thus emerges as a powerful catalyst for breaking collusive stability and driving a transition to competition, as the new firms' presence typically precipitates the breakdown and intensifies the ensuing price war. Policies that encourage entry—for example by reducing barriers for generic drug suppliers—can be especially effective in undermining collusion and delivering lower prices to consumers. Finally, an avenue for further research would be to model firms' entry and exit decisions more explicitly. By understanding when and why potential competitors decide to enter (or incumbents choose to exit), future work could better predict the conditions under which bid-rotation collapses, thereby deepening our understanding of how market structure dynamics interact with collusive behavior.

References

- Abrantes-Metz, Rosa M et al. (2006). “A variance screen for collusion”. In: *International Journal of Industrial Organization* 24.3, pp. 467–486.
- Aoyagi, Masaki (2003). “Bid rotation and collusion in repeated auctions”. In: *Journal of Economic Theory* 112.1, pp. 79–105. ISSN: 0022-0531.
- Athey, Susan, Kyle Bagwell, and Chris Sanchirico (2004). “Collusion and price rigidity”. In: *The Review of Economic Studies* 71.2, pp. 317–349.
- Aumann, Robert J and Lloyd S Shapley (1994). “Long-term competition: a game-theoretic analysis”. In: *Essays in Game Theory: In Honor of Michael Maschler*. Springer, pp. 1–15.
- Bergman, Mats A, David Granlund, and Niklas Rudholm (2017). “Squeezing the last drop out of your suppliers: an empirical study of market-based purchasing policies for generic pharmaceuticals”. In: *Oxford Bulletin of Economics and Statistics* 79.6, pp. 969–996.
- Berndt, Ernst R (2002). “Pharmaceuticals in US health care: determinants of quantity and price”. In: *Journal of Economic Perspectives* 16.4, pp. 45–66.
- Berndt, Ernst R, Rena M Conti, and Stephen J Murphy (2018). “The generic drug user fee amendments: an economic perspective”. In: *Journal of Law and the Biosciences* 5.1, pp. 103–141.
- Bertrand, Joseph (1883). “Review of *Theorie mathématique de la richesse sociale* and of *Recherches sur les principes mathématiques de la théorie des richesses*.” In: *Journal de savants* 67, p. 499.
- Bolotova, Yuliya, John M Connor, and Douglas J Miller (2008). “The impact of collusion on price behavior: Empirical results from two recent cases”. In: *International Journal of Industrial Organization* 26.6, pp. 1290–1307.
- Byrne, David P and Nicolas De Roos (2019). “Learning to coordinate: A study in retail gasoline”. In: *American Economic Review* 109.2, pp. 591–619.

- Caves, Richard E et al. (1991). “Patent expiration, entry, and competition in the US pharmaceutical industry”. In: *Brookings papers on economic activity. Microeconomics* 1991, pp. 1–66.
- Cletus, Jadwiga (2016). “Investigation of bid collusion within the swedish generic drugs market”. MA thesis. University of Gothenburg School of Business, Economics and Law.
- Cuddy, Emily (2020). *Competition and collusion in the generic drug market*. Tech. rep. Technical report, Mimeo.
- Danzon, PM and Carrier, MA (2021). The neglected
- Flood, Merrill M (1958). “Some experimental games”. In: *Management Science* 5.1, pp. 5–26.
- Friedman, James W (1971). “A non-cooperative equilibrium for supergames”. In: *The Review of Economic Studies* 38.1, pp. 1–12.
- Granlund, David and Mats A Bergman (2018). “Price competition in pharmaceuticals—evidence from 1303 Swedish markets”. In: *Journal of health economics* 61, pp. 1–12.
- Granlund, David and Niklas Rudholm (2012). “The prescribing physicians influence on consumer choice between medically equivalent pharmaceuticals”. In: *Review of Industrial Organization* 41, pp. 207–222.
- (2023). “A method for calculating the probability of collusion based on observed price patterns”. In: *Available at SSRN 4605725*.
- Grossman, Gene M and Elhanan Helpman (1993). *Innovation and growth in the global economy*. MIT press.
- Grossman, Gene M and Edwin L-C Lai (2004). “International protection of intellectual property”. In: *American Economic Review* 94.5, pp. 1635–1653.
- Habib, Hedi Bel (2017). “Prismodeller och prispress på läkemedelsmarknaden, Konkurrensverkets rapportserie 2017:9”. In: *Konkurrensverket*.
- Hall, Peter, Joel L Horowitz, and Bing-Yi Jing (1995). “On blocking rules for the bootstrap with dependent data”. In: *Biometrika* 82.3, pp. 561–574.
- Harrington Jr, Joseph E and Joe Chen (2006). “Cartel pricing dynamics with cost variability and endogenous buyer detection”. In: *International Journal of Industrial Organization* 24.6, pp. 1185–1212.
- Ivaldi, Marc et al. (2003). “The economics of tacit collusion”. In.
- Janssen, Aljoscha (2022). “Price dynamics of swedish pharmaceuticals”. In: *Quantitative Marketing and Economics* 20.4, pp. 313–351.
- Lakdawalla, Darius N (2018). “Economics of the pharmaceutical industry”. In: *Journal of Economic Literature* 56.2, pp. 397–449.
- Morton, Fiona Scott and Lysle T Boller (2017). “Enabling competition in pharmaceutical markets”. In: *Brookings, Brookings* 2.
- Nordhaus, William D (1969). “An economic theory of technological change”. In: *The American Economic Review* 59.2, pp. 18–28.
- Oronsky, Bryan et al. (2023). “Patent and marketing exclusivities 101 for drug developers”. In: *Recent Patents on Biotechnology* 17.3, p. 257.

- Politis, Dimitris N and Joseph P Romano (1994). “The stationary bootstrap”. In: *Journal of the American Statistical association* 89.428, pp. 1303–1313.
- Rondahl, Karolina (2022). *Immaterialrätt inom life science*. <https://www.lif.se/sa-funkar-det/immaterialratt-inom-life-science/>. Accessed: 2025-01-04.
- Schmidheiny, Kurt and Sebastian Siegloch (2023). “On event studies and distributed-lags in two-way fixed effects models: Identification, equivalence, and generalization”. In: *Journal of Applied Econometrics* 38.5, pp. 695–713. DOI: <https://doi.org/10.1002/jae.2971>. eprint: <https://onlinelibrary.wiley.com/doi/pdf/10.1002/jae.2971>. URL: <https://onlinelibrary.wiley.com/doi/abs/10.1002/jae.2971>.
- Skrzypacz, Andrzej and Hugo Hopenhayn (2004). “Tacit collusion in repeated auctions”. In: *Journal of Economic Theory* 114.1, pp. 153–169.
- Starc, Amanda and Thomas G Wollmann (2025). “Does entry remedy collusion? Evidence from the generic prescription drug cartel”. In: *American Economic Review* 115.5, pp. 1400–1438.
- TLV (2025). *Periodens vara*. <https://www.tlv.se/apotek/generiskt-utbyte/periodens-varor.html>. Accessed: 2025-01-04.
- Wallström, Sofia et al. (2012). “Pris, tillgång och service,-fortsatt utveckling av läkemedels-och apoteksmarknaden”. In: *SOU 2012:75*.
- Wouters, Olivier J, Martin McKee, and Jeroen Luyten (2020). “Estimated research and development investment needed to bring a new medicine to market, 2009-2018”. In: *Jama* 323.9, pp. 844–853.

A Data

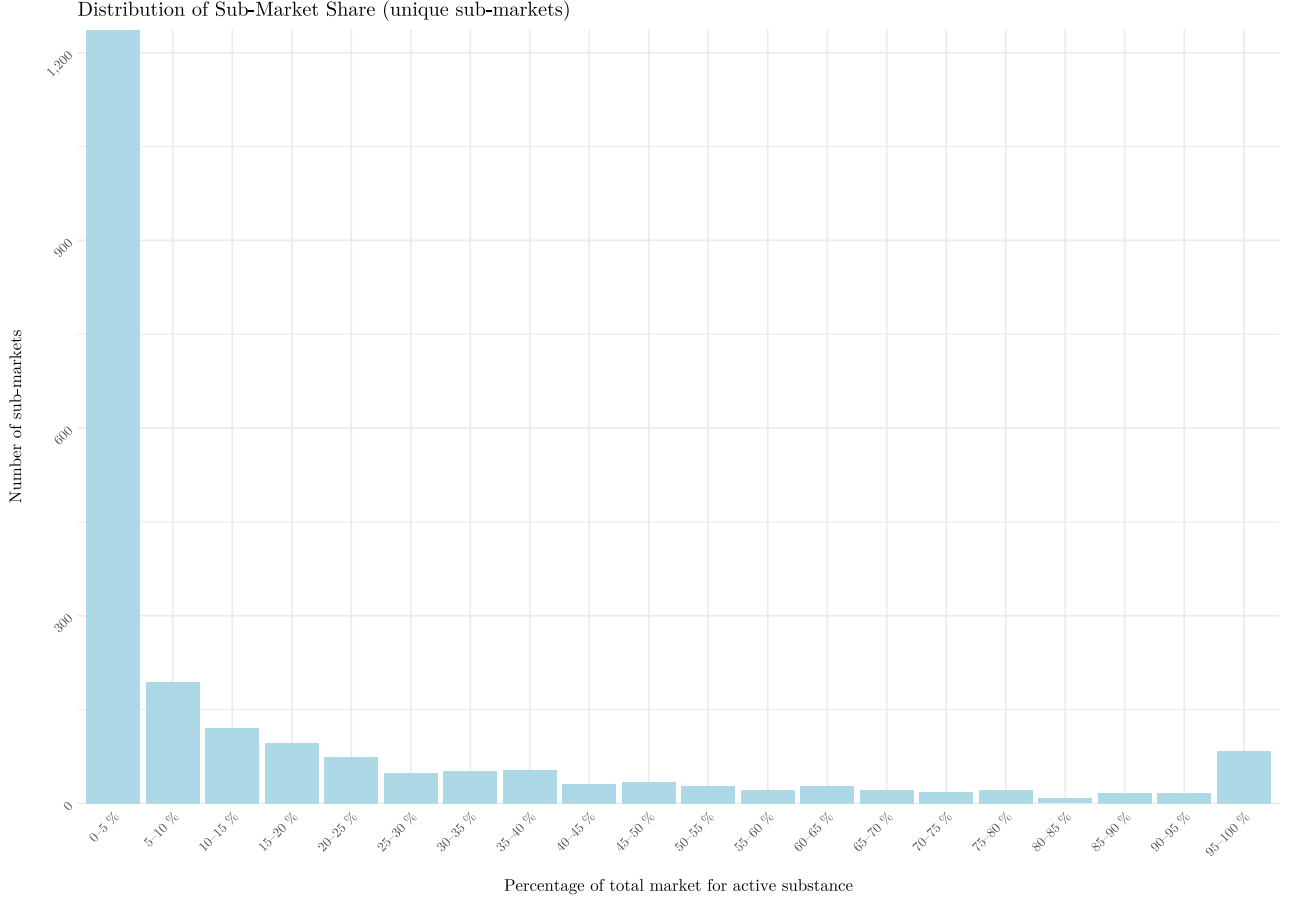


Figure 15: Market sizes – Full sample

B Empirical strategy

We take our series $\{W_t\}_{t=1}^T$ and our partitions P_t and create two candidates for perfect alteration. These are

$$C_n^{(1)}(P_t) = (W_t, \dots, W_{t+n-1}, W_t, \dots, W_{t+n-1})$$

and

$$C_n^{(2)}(P_t) = (W_{t+n-1}, \dots, W_t, W_{t+n-1}, \dots, W_t)$$

The first candidate is simply the first n elements in the partition P_t repeated twice (remember that the partition has length $k = 2n$). The second candidate has the same elements but in reverse order. The candidates are there to pick up the most pronounced pattern in the partition. This is needed for the algorithm to pick up a pattern from which deviations are found. They also provide flexibility since they allow me to not pre-code all possible bid rotation patterns. With these candidates in place I can compute the so call Hamming-distances of P_t as

$$d_j(t) = \sum_{i=1}^k \mathbb{1}\{P_{i,t} \neq C_{n,i}^{(j)}(P_t)\} \quad \text{for } j = 1, 2 \quad (26)$$

For a given partition P_t with length k the formula computes the number of mismatches between the two candidates and the actual partition. We then select the pattern that is closest to a pattern of bid rotation like so

$$d(t) = \min\{d_1(t), d_2(t)\} \quad (27)$$

Effectively, we have selected the pattern that is least far off from a bid rotation pattern ⁵⁵ and we can now assign a score S_t to it based on how closely the patterns resembles bid rotation. For this I use a non-linear transformation of the Hamming distance

$$S_t = \begin{cases} 0, & \text{if } |\{W_t, W_{t+1}, \dots, W_{t+n-1}\}| < n \\ \left(\frac{k-d(t)}{k}\right)^p, & \text{if } |\{W_t, W_{t+1}, \dots, W_{t+n-1}\}| = n, \text{ for } p \geq 1 \end{cases} \quad (28)$$

C Robustness check – overlapping permutations test

⁵⁵In case of a tie both candidates are equally far off and $d(t) = \min\{d_1(t), d_2(t)\} = d_1(t) = d_2(t)$

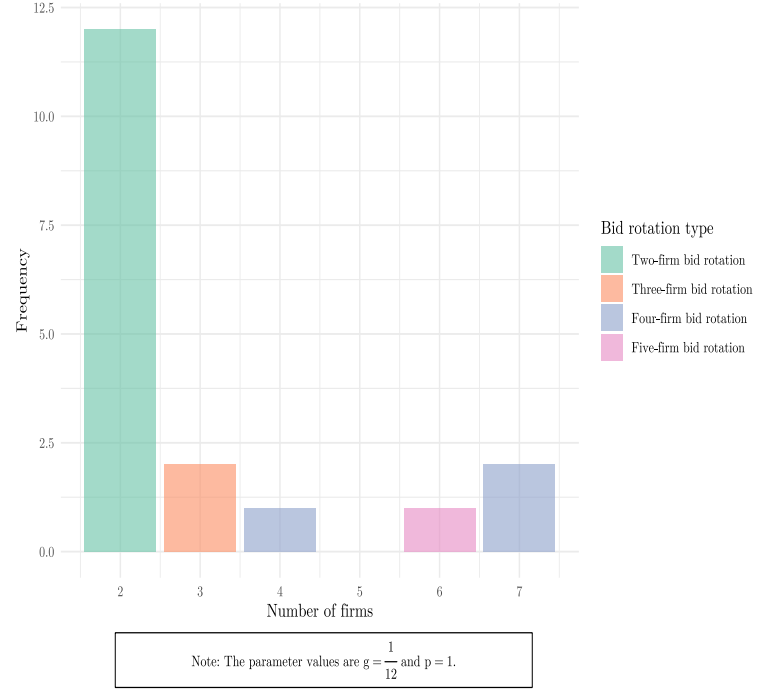
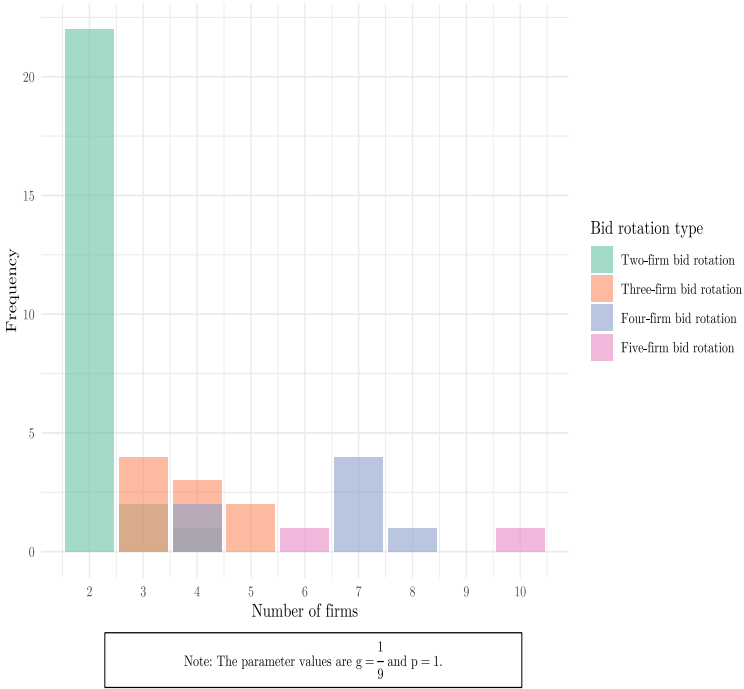
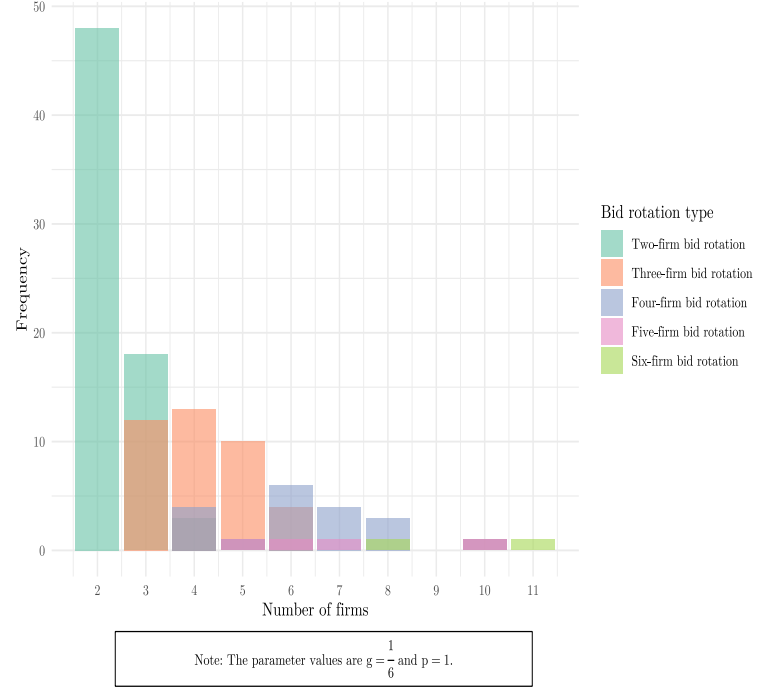
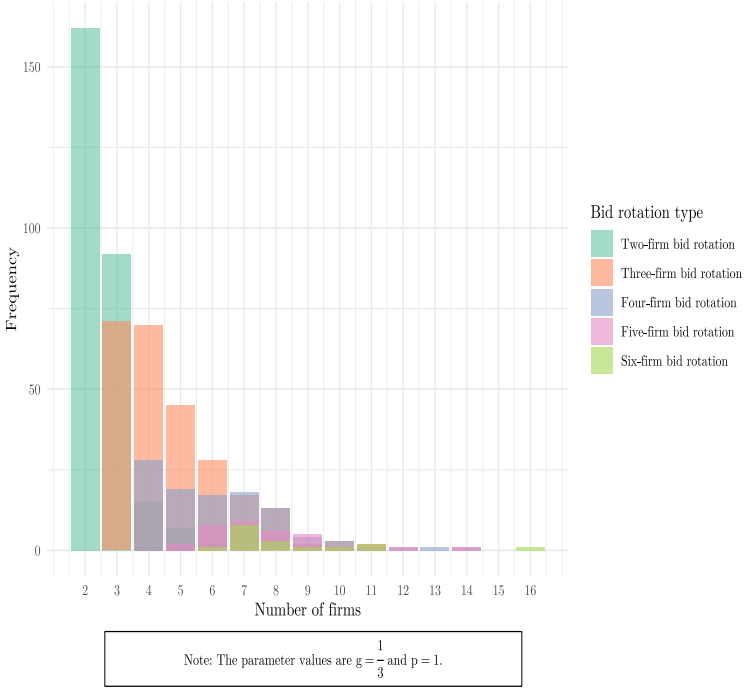


Figure 16: Robustness check with penalty parameter equal to 1

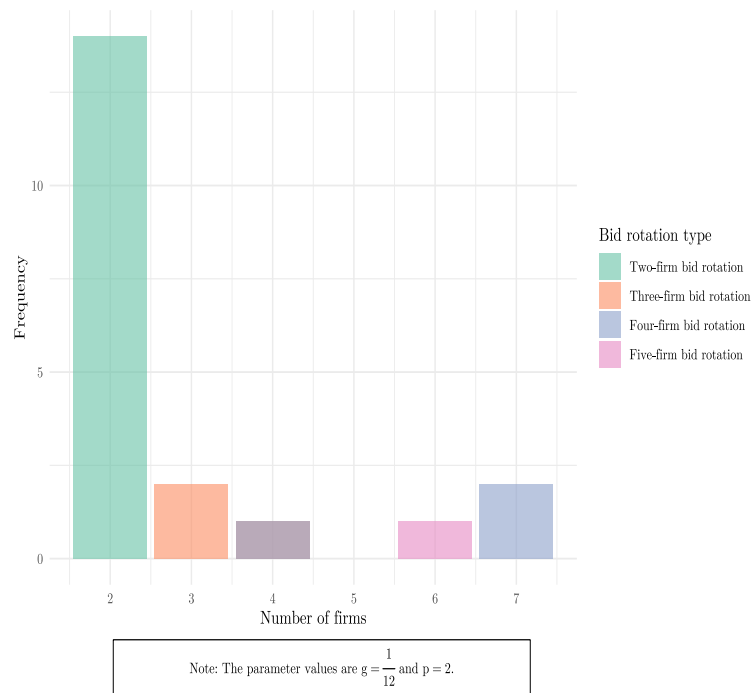
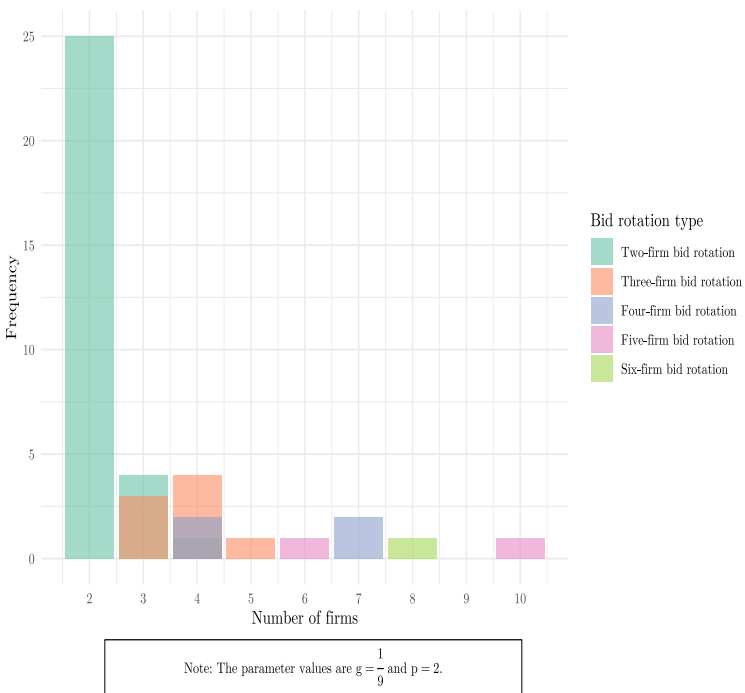
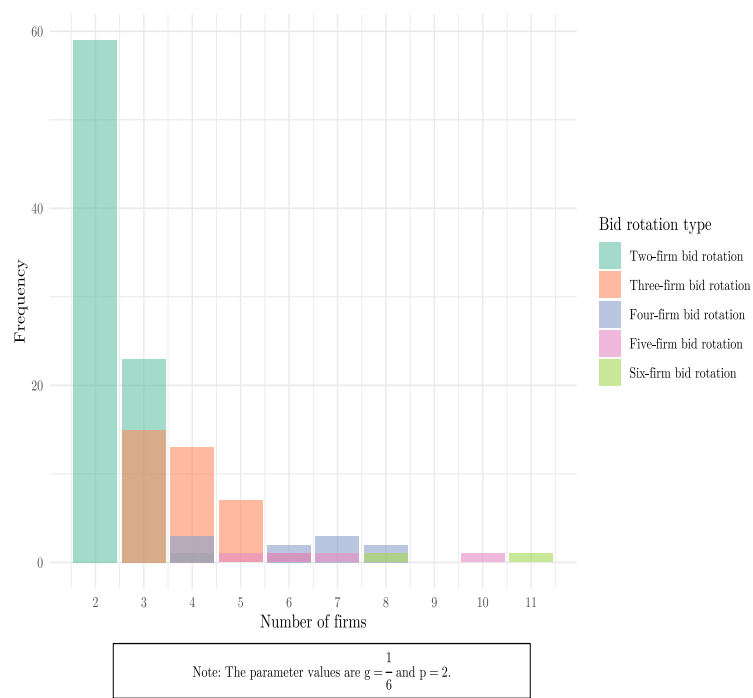
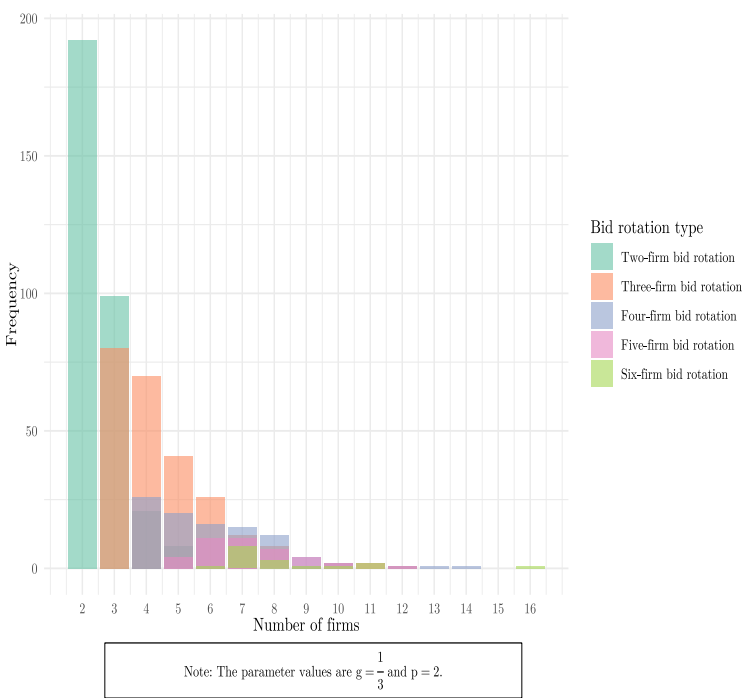


Figure 17: Robustness check with penalty parameter equal to 2

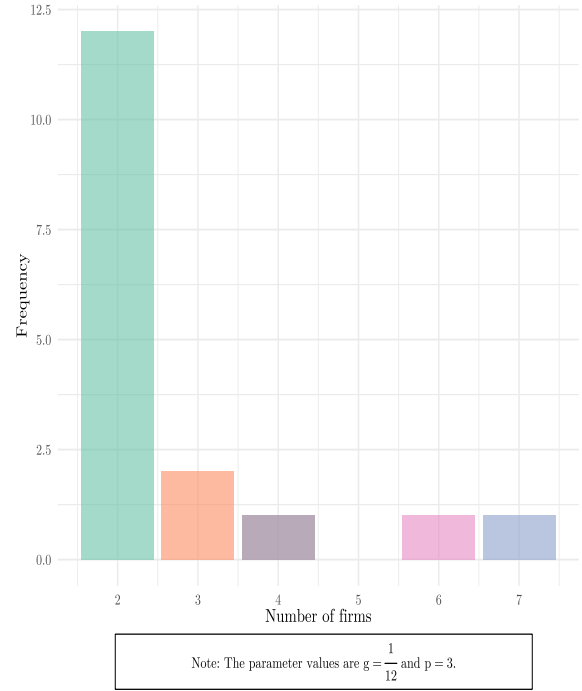
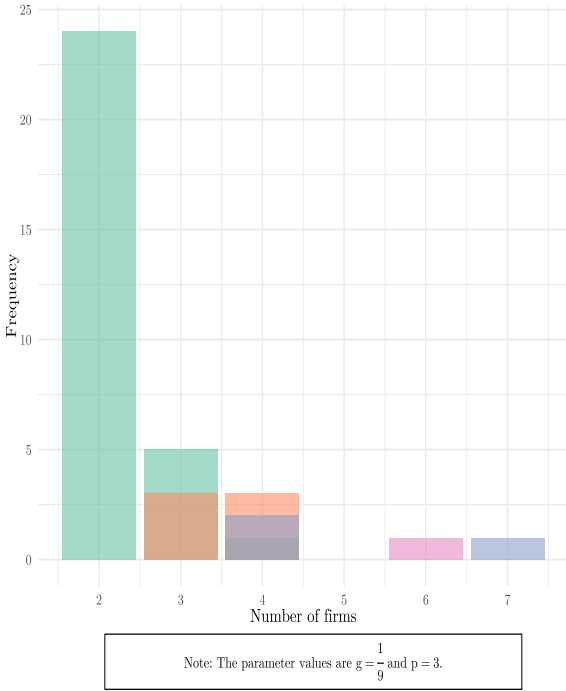
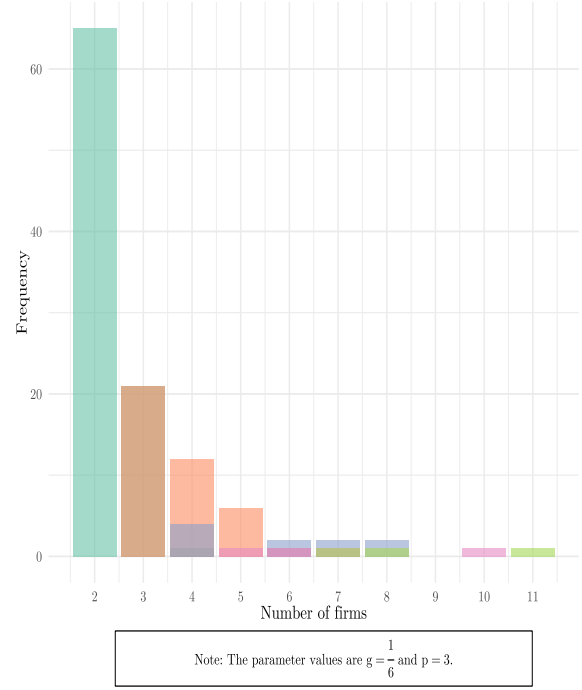
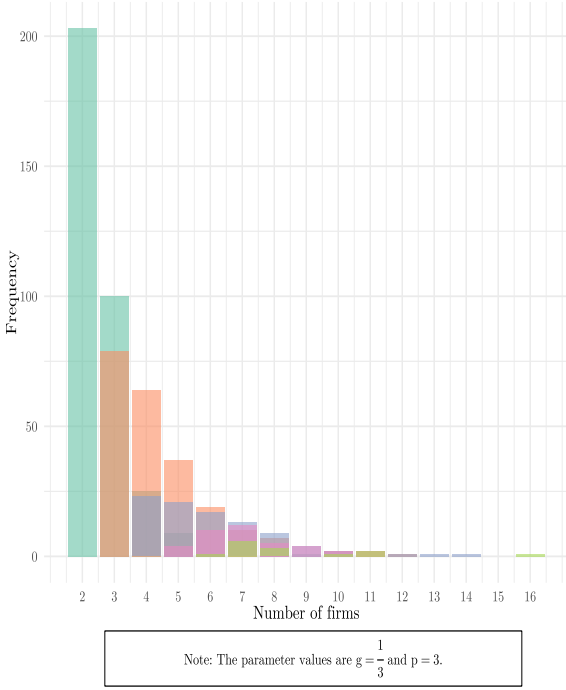


Figure 18: Robustness check with penalty parameter equal to 3

D Event-study

D.1 Coefficient plots

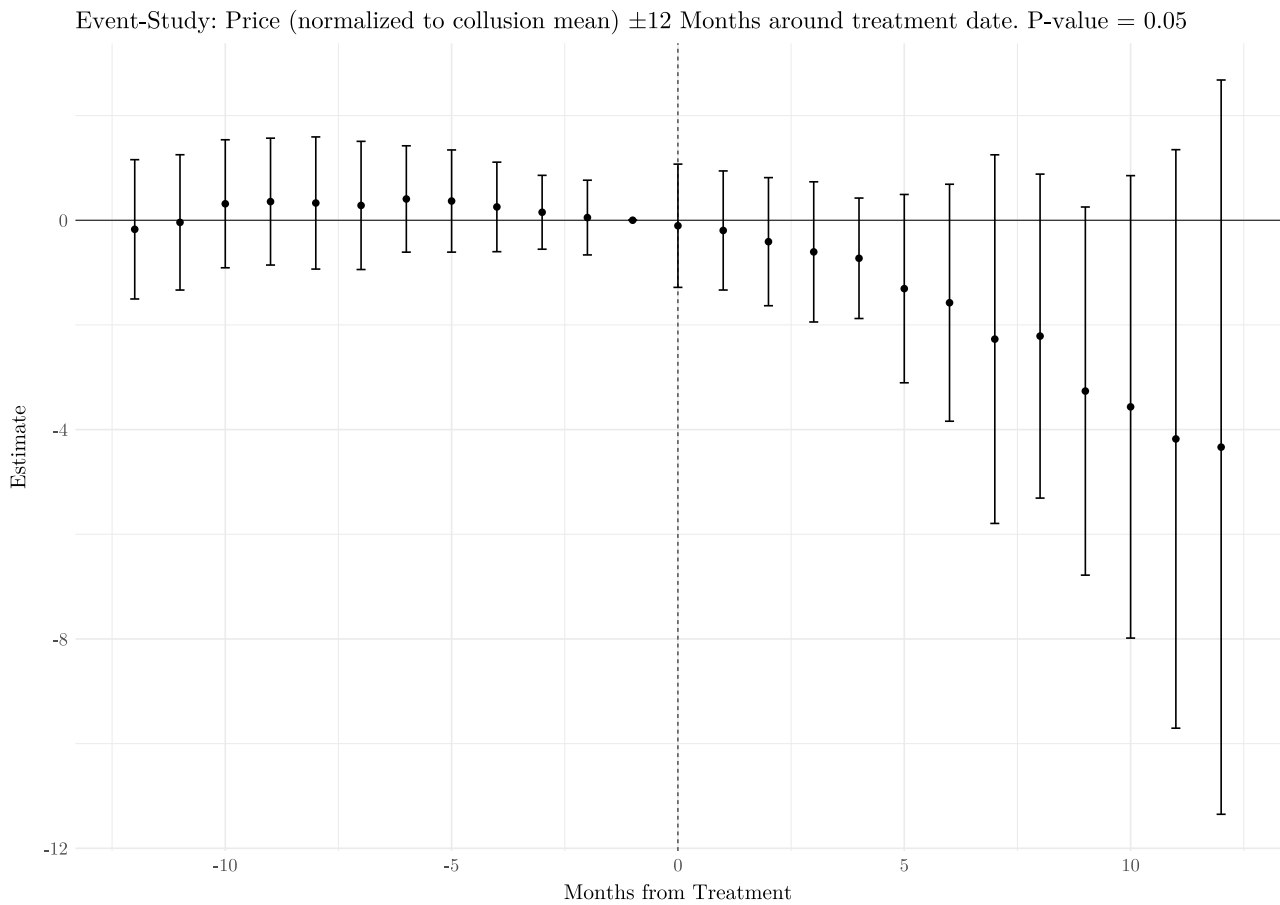


Figure 19: Event study coefficient plot – P-value 5%

Event-Study: Price (normalized to collusion mean) ± 12 Months around treatment date. P-value = 0.05.
 Only segments where number of firms increase after collusion

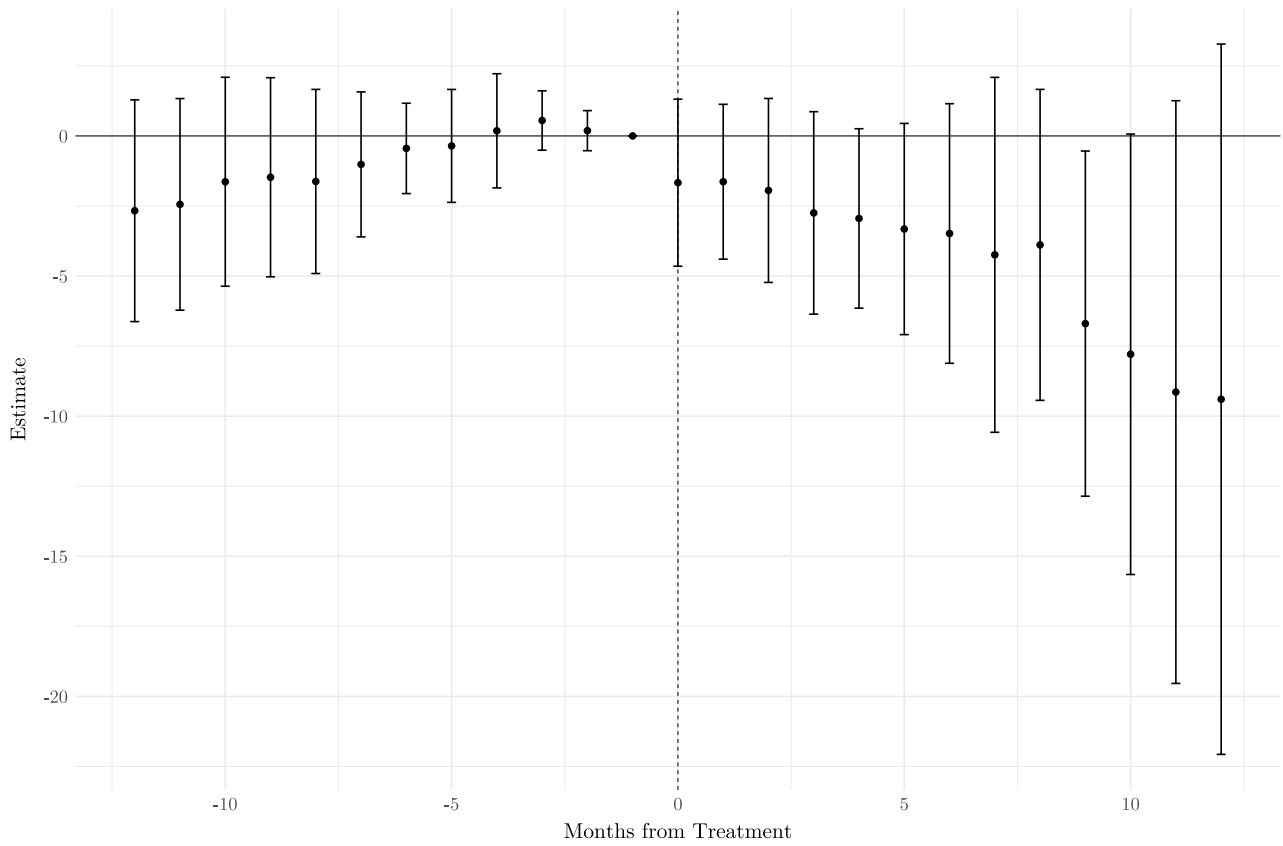


Figure 20: Event study coefficient plot – P-value 5%

Event-Study results: CoV (normalized to collusion CoV): ± 12 Years around treatment date. P-value = 0.05.

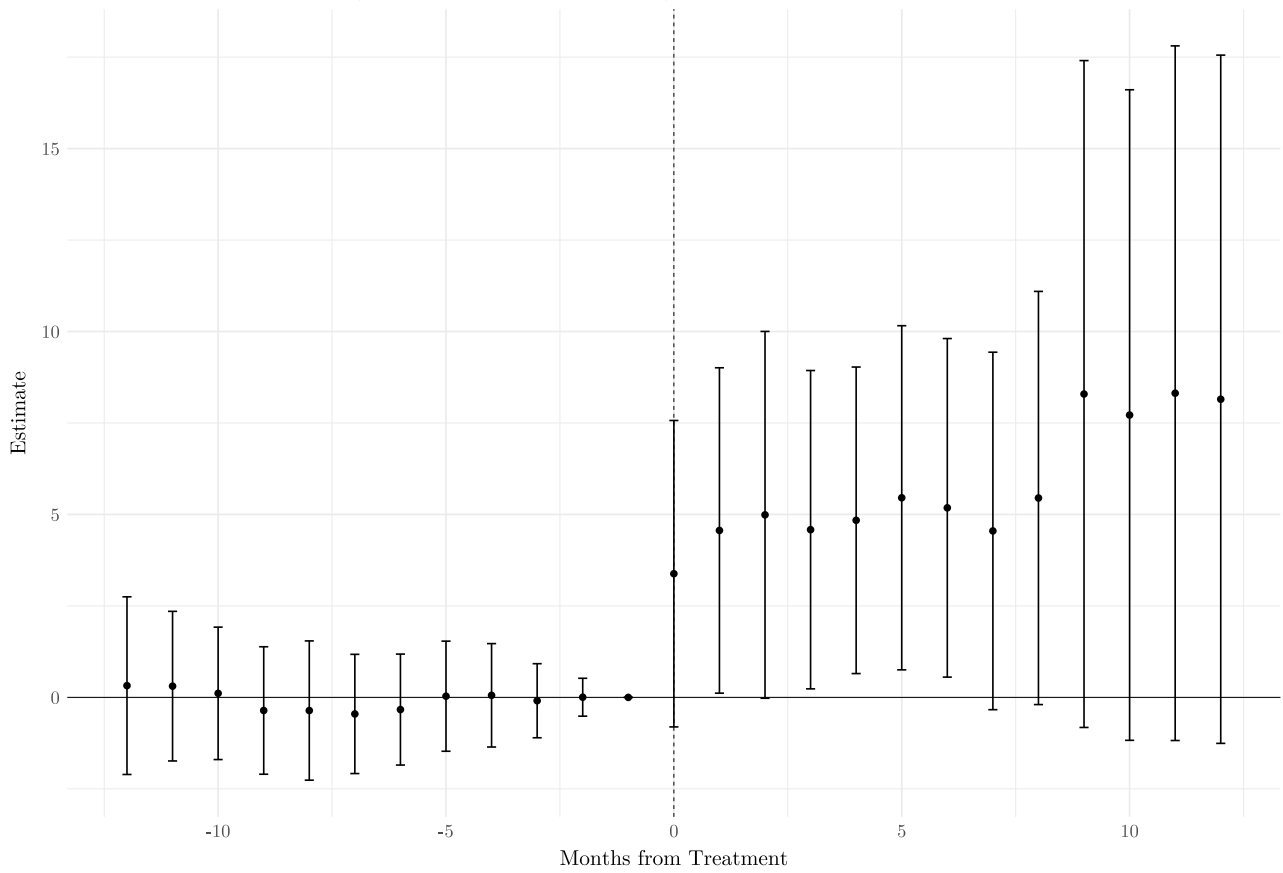


Figure 21: Event study coefficient plot – P-value 5%

Event-Study results: CoV (normalized to collusion CoV): ± 12 Years around treatment date. P-value = 0.05.
 Only segments where number of firms increase after collusion

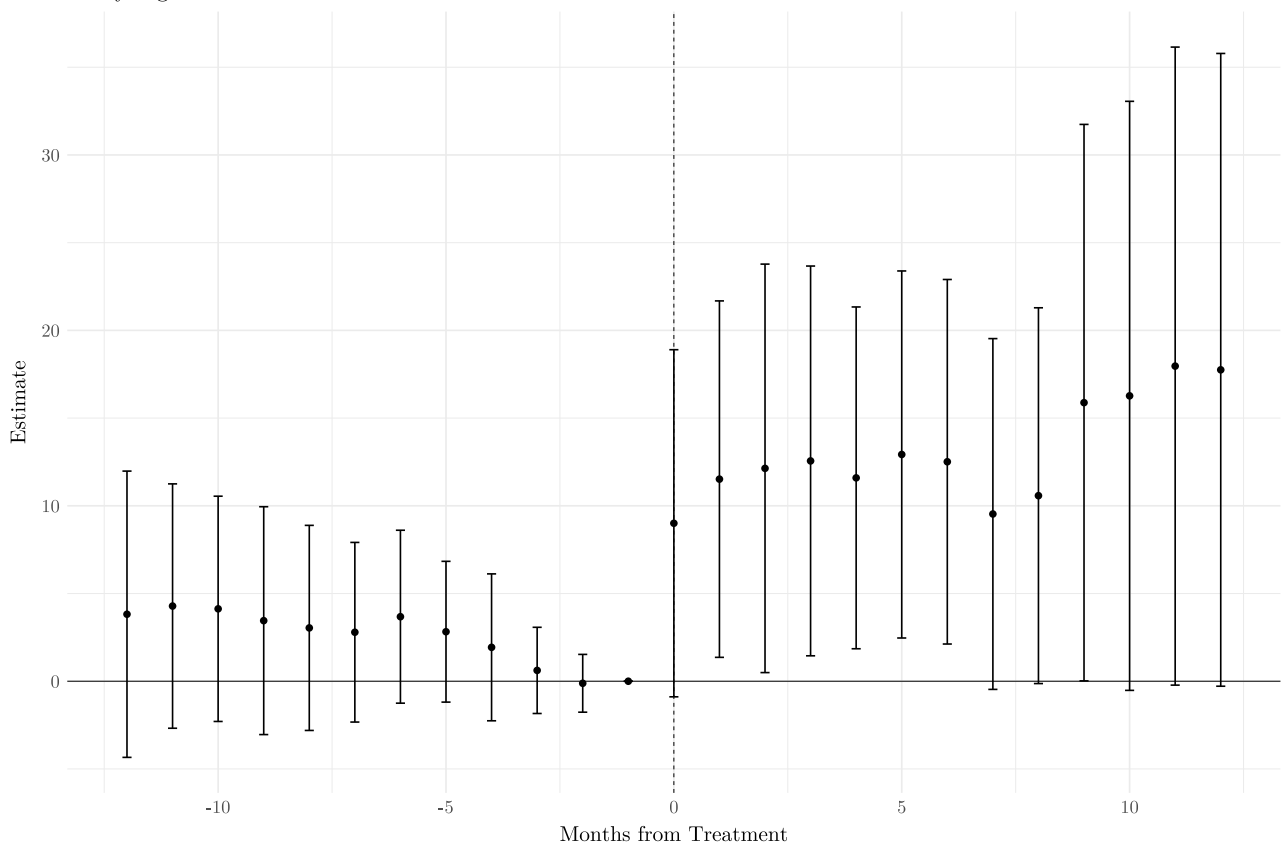


Figure 22: Event study coefficient plot – P-value 5%

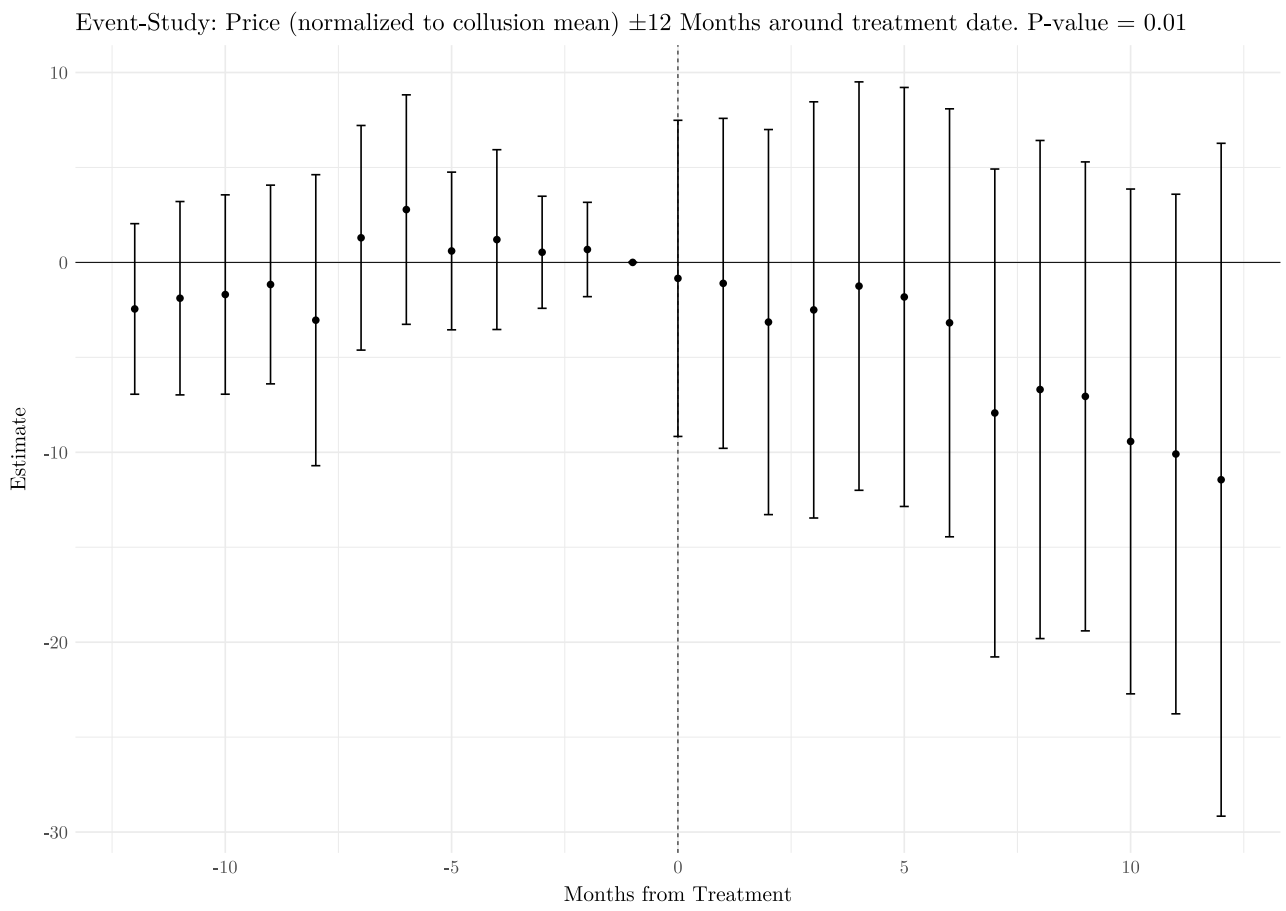


Figure 23: Event study coefficient plot – P-value 1%

Event-Study: Price (normalized to collusion mean) ± 12 Months around treatment date. P-value = 0.01.
 Only segments where number of firms increase after collusion

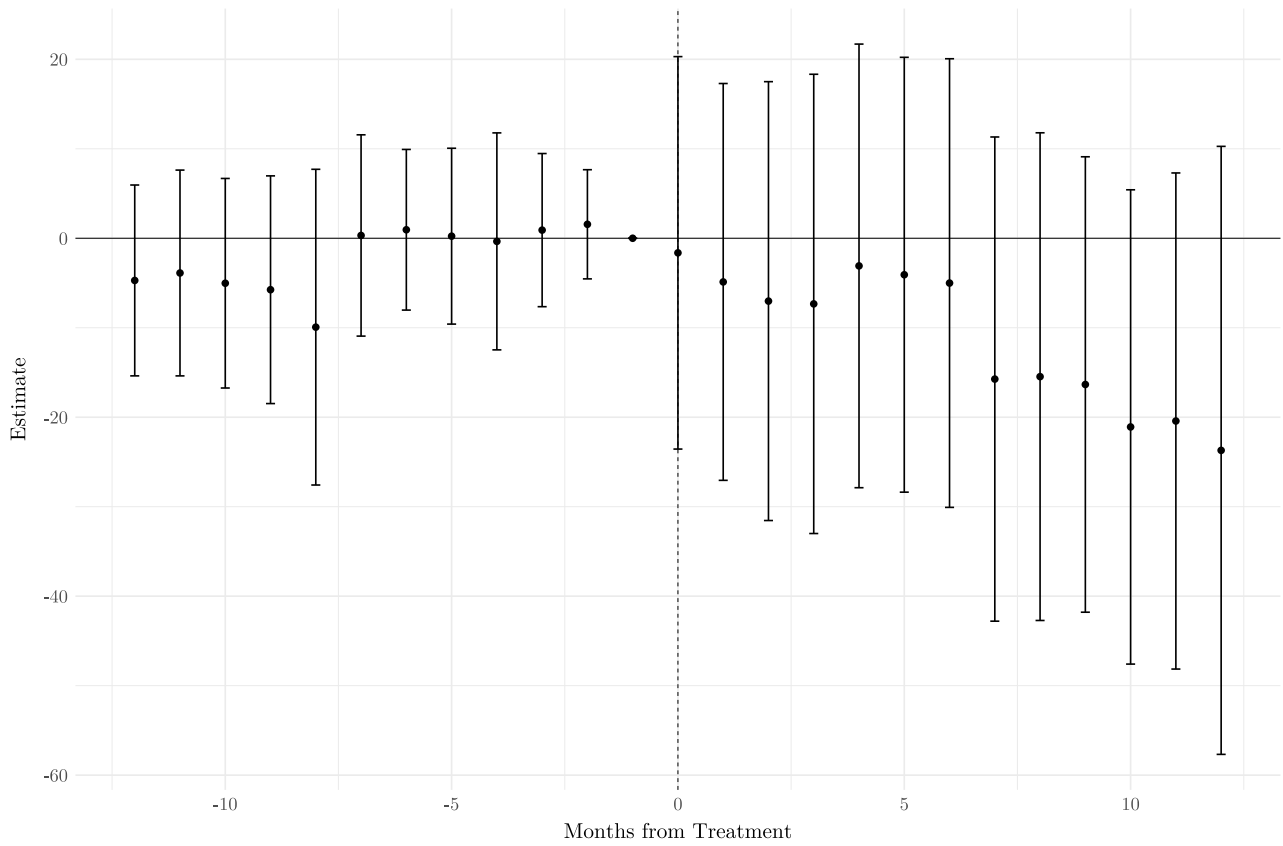


Figure 24: Event study coefficient plot – P-value 1%

Event-Study results: CoV (normalized to collusion CoV): ± 12 Years around treatment date. P-value = 0.01.

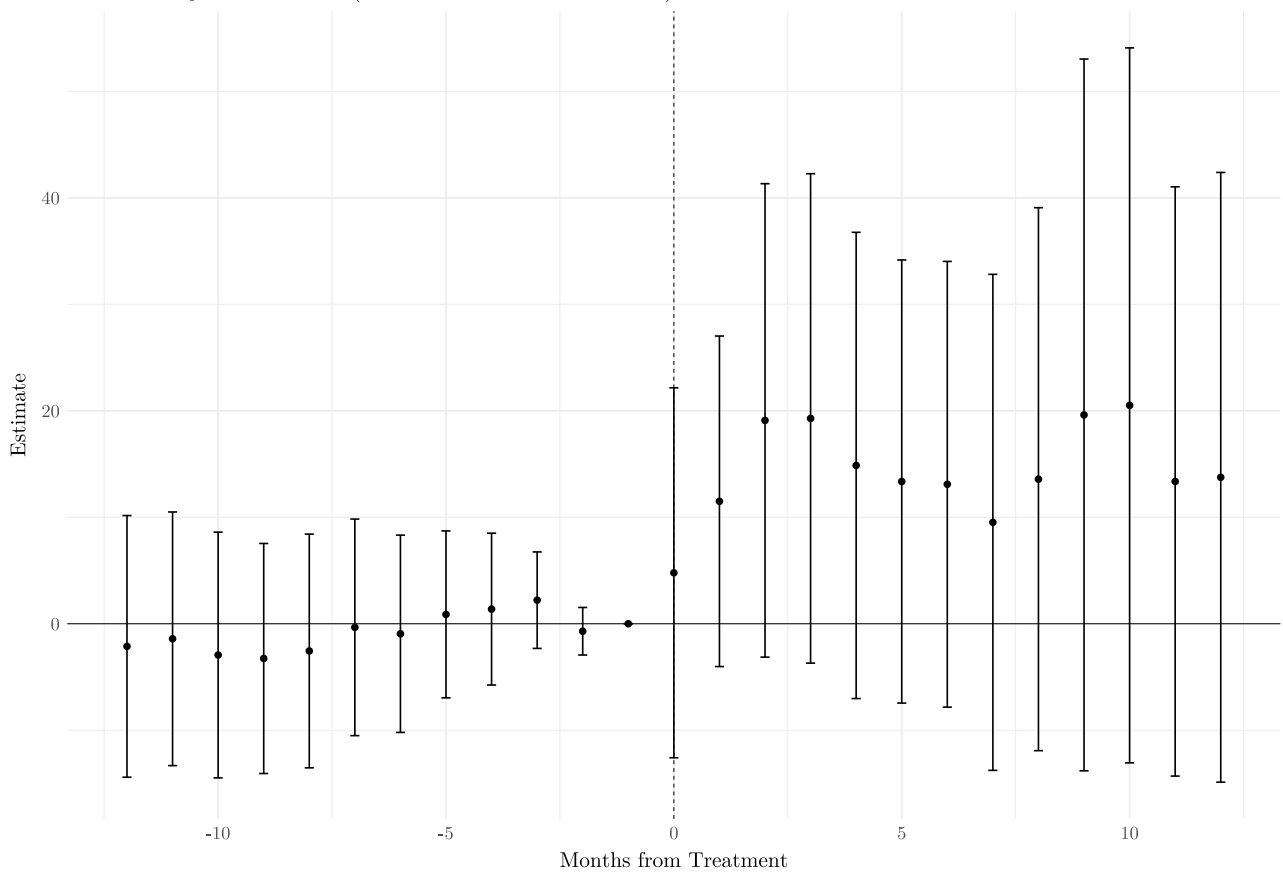


Figure 25: Event study coefficient plot – P-value 1%

Event-Study results: CoV (normalized to collusion CoV): ± 12 Years around treatment date. P-value = 0.01.
 Only segments where number of firms increase after collusion

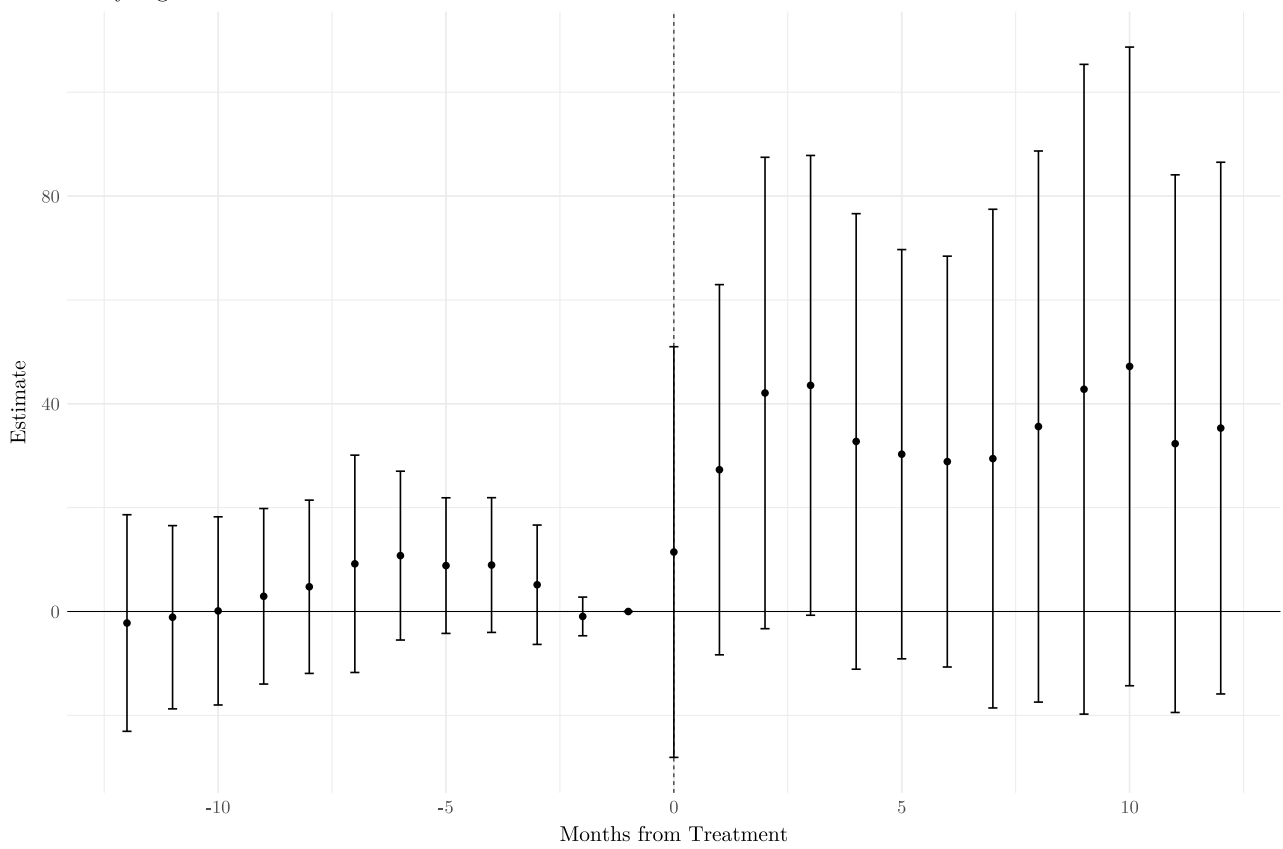


Figure 26: Event study coefficient plot – P-value 1%

D.2 Tables

Table 8: EventStudy results: Price (normalized to collusion mean) ALL SEGMENTS

Dependent Variable: Model:	(P-value = 0.1)	aip_norm (P-value = 0.05)	(P-value = 0.01)
<i>Variables</i>			
event_time = -12	0.0092 (0.0800)	-0.1738 (0.6674)	-2.451 (2.160)
event_time = -11	0.0373 (0.0747)	-0.0410 (0.6482)	-1.886 (2.448)
event_time = -10	0.0417 (0.0804)	0.3144 (0.6134)	-1.692 (2.524)
event_time = -9	0.0431 (0.0766)	0.3560 (0.6078)	-1.164 (2.516)
event_time = -8	0.0947 (0.0712)	0.3292 (0.6331)	-3.043 (3.684)
event_time = -7	0.0969 (0.0621)	0.2828 (0.6141)	1.296 (2.844)
event_time = -6	0.0799 (0.0654)	0.4062 (0.5095)	2.781 (2.906)
event_time = -5	0.0422 (0.0502)	0.3667 (0.4895)	0.6016 (1.997)
event_time = -4	0.0588 (0.0457)	0.2545 (0.4290)	1.200 (2.277)
event_time = -3	0.0598 (0.0371)	0.1522 (0.3539)	0.5335 (1.419)
event_time = -2	0.0065 (0.0337)	0.0510 (0.3580)	0.6796 (1.194)
event_time = 0	-0.1376* (0.0726)	-0.1052 (0.5905)	-0.8408 (4.003)
event_time = 1	-0.2110** (0.0861)	-0.1958 (0.5705)	-1.103 (4.178)
event_time = 2	-0.2536** (0.0999)	-0.4095 (0.6137)	-3.143 (4.876)
event_time = 3	-0.3443*** (0.0964)	-0.6050 (0.6716)	-2.502 (5.269)
event_time = 4	-0.4641*** (0.1002)	-0.7269 (0.5771)	-1.247 (5.171)
event_time = 5	-0.4869*** (0.1138)	-1.307 (0.9023)	-1.821 (5.305)
event_time = 6	-0.5161*** (0.1233)	-1.576 (1.135)	-3.181 (5.419)
event_time = 7	-0.5319*** (0.1521)	-2.271 (1.766)	-7.930 (6.177)
event_time = 8	-0.6810*** (0.1665)	-2.214 (1.552)	-6.695 (6.307)
event_time = 9	-0.8593*** (0.2299)	-3.264* (1.764)	-7.056 (5.938)
event_time = 10	-0.9074*** (0.2320)	-3.566 (2.216)	-9.429 (6.392)
event_time = 11	-0.9092*** (0.2432)	-4.179 (2.772)	-10.09 (6.579)
event_time = 12	-0.9593** (0.4080)	-4.335 (3.518)	-11.45 (8.520)
<i>Fixed-effects</i>			
segments	Yes	Yes	Yes
valid_from	Yes	Yes	Yes
<i>Fit statistics</i>			
Observations	35,906	8,638	1,657
R ²	0.04464	0.10385	0.28382
Within R ²	0.01470	0.02560	0.03422

Clustered (substance_id) standard-errors in parentheses
Signif. Codes: ***: 0.01, **: 0.05, *: 0.1

Table 9: EventStudy results: CoV (normalized to collusion CoV) ALL SEGMENTS

Dependent Variable: Model:	(P-value = 0.1)	Cov_norm (P-value = 0.05)	(P-value = 0.01)
<i>Variables</i>			
event_time = -12	-0.0568 (0.1449)	0.3210 (1.218)	-2.126 (5.909)
event_time = -11	-0.0667 (0.1356)	0.3070 (1.025)	-1.414 (5.728)
event_time = -10	-0.0473 (0.1196)	0.1104 (0.9076)	-2.936 (5.549)
event_time = -9	-0.1763 (0.1590)	-0.3584 (0.8738)	-3.264 (5.194)
event_time = -8	-0.1465 (0.1491)	-0.3595 (0.9545)	-2.558 (5.277)
event_time = -7	-0.1437 (0.1311)	-0.4527 (0.8171)	-0.3385 (4.890)
event_time = -6	-0.1491 (0.1283)	-0.3330 (0.7605)	-0.9420 (4.455)
event_time = -5	-0.0711 (0.0831)	0.0325 (0.7546)	0.8792 (3.769)
event_time = -4	-0.0657 (0.0771)	0.0565 (0.7089)	1.375 (3.430)
event_time = -3	-0.0132 (0.0568)	-0.0914 (0.5076)	2.215 (2.178)
event_time = -2	-0.0475 (0.0384)	0.0042 (0.2598)	-0.7042 (1.074)
event_time = 0	1.007** (0.4797)	3.381 (2.100)	4.786 (8.353)
event_time = 1	1.098** (0.5173)	4.562** (2.230)	11.50 (7.464)
event_time = 2	1.186** (0.5352)	4.989* (2.514)	19.10* (10.69)
event_time = 3	1.211** (0.5334)	4.584** (2.182)	19.28* (11.05)
event_time = 4	1.212** (0.5462)	4.840** (2.101)	14.87 (10.53)
event_time = 5	1.314** (0.5508)	5.456** (2.359)	13.36 (10.01)
event_time = 6	1.331** (0.5508)	5.181** (2.321)	13.10 (10.06)
event_time = 7	1.315** (0.5630)	4.550* (2.450)	9.523 (11.20)
event_time = 8	1.515** (0.5826)	5.450* (2.832)	13.58 (12.26)
event_time = 9	2.210** (0.9101)	8.292* (4.571)	19.61 (16.07)
event_time = 10	2.349** (0.9562)	7.717* (4.459)	20.51 (16.14)
event_time = 11	2.440** (0.9747)	8.315* (4.761)	13.36 (13.31)
event_time = 12	2.510** (1.015)	8.149* (4.717)	13.75 (13.77)
<i>Fixed-effects</i>			
segment_id	Yes	Yes	Yes
valid_from	Yes	Yes	Yes
<i>Fit statistics</i>			
Observations	35,884	8,630	1,654
R ²	0.04884	0.18531	0.59557
Within R ²	0.01487	0.02799	0.05823

Clustered (substance_id) standard-errors in parentheses
Signif. Codes: ***: 0.01, **: 0.05, *: 0.1

Table 10: EventStudy results: Price (normalized to collusion mean) INCREASING N FIRMS SEGMENTS

Dependent Variable: Model:	(P-value = 0.1)	aip_norm (P-value = 0.05)	(P-value = 0.01)
<i>Variables</i>			
event_time = -12	-0.0389 (0.1435)	-2.669 (1.937)	-4.715 (4.715)
event_time = -11	0.0854 (0.1426)	-2.442 (1.849)	-3.881 (5.084)
event_time = -10	0.1171 (0.1325)	-1.634 (1.826)	-5.026 (5.177)
event_time = -9	0.0538 (0.1284)	-1.475 (1.740)	-5.751 (5.625)
event_time = -8	0.0318 (0.1206)	-1.624 (1.609)	-9.935 (7.799)
event_time = -7	0.0524 (0.1233)	-1.016 (1.267)	0.3131 (4.972)
event_time = -6	0.1353 (0.1196)	-0.4448 (0.7892)	0.9494 (3.969)
event_time = -5	0.0882 (0.1132)	-0.3561 (0.9867)	0.2333 (4.344)
event_time = -4	0.1266 (0.0979)	0.1823 (0.9976)	-0.3488 (5.360)
event_time = -3	0.1132 (0.0733)	0.5514 (0.5186)	0.9108 (3.783)
event_time = -2	0.0128 (0.0622)	0.1880 (0.3503)	1.560 (2.697)
event_time = 0	-0.4662*** (0.1119)	-1.668 (1.459)	-1.629 (9.694)
event_time = 1	-0.5871*** (0.1500)	-1.634 (1.352)	-4.880 (9.803)
event_time = 2	-0.6996*** (0.1795)	-1.945 (1.607)	-7.022 (10.84)
event_time = 3	-0.8471*** (0.1886)	-2.747 (1.769)	-7.334 (11.34)
event_time = 4	-0.9984*** (0.1794)	-2.945* (1.567)	-3.089 (10.96)
event_time = 5	-1.056*** (0.2272)	-3.322* (1.845)	-4.077 (10.74)
event_time = 6	-1.149*** (0.2481)	-3.481 (2.268)	-5.007 (11.08)
event_time = 7	-1.085*** (0.2879)	-4.242 (3.101)	-15.74 (11.96)
event_time = 8	-1.293*** (0.3157)	-3.887 (2.717)	-15.46 (12.05)
event_time = 9	-1.561*** (0.4135)	-6.698** (3.016)	-16.34 (11.25)
event_time = 10	-1.620*** (0.4175)	-7.791* (3.847)	-21.09 (11.72)
event_time = 11	-1.652*** (0.4423)	-9.140* (5.091)	-20.42 (12.25)
event_time = 12	-1.663** (0.7029)	-9.395 (6.205)	-23.71 (15.02)
<i>Fixed-effects</i>			
segment_id	Yes	Yes	Yes
valid_from	Yes	Yes	Yes
<i>Fit statistics</i>			
Observations	16,587	3,935	1,046
R ²	0.05537	0.15953	0.34098
Within R ²	0.02662	0.04588	0.08104

Clustered (substance_id) standard-errors in parentheses
Signif. Codes: ***: 0.01, **: 0.05, *: 0.1

Table 11: EventStudy results: CoV (normalized to collusion CoV) INCREASING N FIRM SEGMENTS

Dependent Variable: Model:	(P-value = 0.1)	Cov_norm (P-value = 0.05)	(P-value = 0.01)
<i>Variables</i>			
event_time = -12	0.3364 (0.3991)	3.819 (3.995)	-2.217 (9.219)
event_time = -11	0.2149 (0.4101)	4.285 (3.410)	-1.105 (7.800)
event_time = -10	0.1443 (0.3771)	4.127 (3.144)	0.1178 (8.010)
event_time = -9	0.1314 (0.3618)	3.453 (3.180)	2.930 (7.472)
event_time = -8	0.1544 (0.3491)	3.039 (2.862)	4.762 (7.376)
event_time = -7	0.0761 (0.3176)	2.793 (2.507)	9.185 (9.251)
event_time = -6	0.0360 (0.2774)	3.678 (2.413)	10.76 (7.186)
event_time = -5	-0.0170 (0.2383)	2.822 (1.965)	8.840 (5.773)
event_time = -4	-0.1453 (0.1874)	1.932 (2.050)	8.946 (5.736)
event_time = -3	-0.0766 (0.1563)	0.6174 (1.203)	5.153 (5.080)
event_time = -2	-0.1505 (0.1363)	-0.1176 (0.8066)	-0.9492 (1.646)
event_time = 0	2.212** (0.9774)	9.005* (4.844)	11.45 (17.48)
event_time = 1	2.376** (1.028)	11.52** (4.975)	27.30 (15.76)
event_time = 2	2.621** (1.110)	12.13** (5.701)	42.08* (20.07)
event_time = 3	2.861** (1.211)	12.56** (5.440)	43.55* (19.57)
event_time = 4	2.673** (1.106)	11.59** (4.770)	32.75 (19.39)
event_time = 5	2.916** (1.165)	12.93** (5.122)	30.29 (17.42)
event_time = 6	2.784** (1.115)	12.51** (5.087)	28.87 (17.49)
event_time = 7	2.704** (1.156)	9.533* (4.896)	29.44 (21.23)
event_time = 8	3.002** (1.204)	10.58* (5.245)	35.62 (23.46)
event_time = 9	4.241** (1.741)	15.88** (7.766)	42.80 (27.66)
event_time = 10	4.578** (1.876)	16.27* (8.221)	47.20 (27.18)
event_time = 11	4.739** (1.898)	17.96* (8.905)	32.32 (22.88)
event_time = 12	4.825** (1.966)	17.75* (8.831)	35.32 (22.63)
<i>Fixed-effects</i>			
segment_id	Yes	Yes	Yes
valid_from	Yes	Yes	Yes
<i>Fit statistics</i>			
Observations	16,574	3,931	1,043
R ²	0.07427	0.28282	0.65069
Within R ²	0.02405	0.04519	0.14833

Clustered (substance_id) standard-errors in parentheses
Signif. Codes: ***: 0.01, **: 0.05, *: 0.1