

Dynamic Local Average Treatment Effects in Time Series^{*}

ALESSANDRO CASINI

University of Rome Tor Vergata

ADAM McCLOSKEY

University of Colorado at Boulder

LUCA ROLLA

University of Rome Tor Vergata

RAIMONDO PALA

University of Rome Tor Vergata

22nd July 2026

Abstract

This paper discusses identification, estimation, and inference on dynamic local average treatment effects (LATEs) in instrumental variables (IVs) settings. First, we show that compliers—observations whose treatment status is affected by the instrument—can be identified *individually* in time series data using smoothness assumptions and local comparisons of treatment assignments. Second, we show that this result enables not only better interpretability of IV estimates but also direct testing of the exclusion restriction by comparing outcomes among identified non-compliers across instrument values. Third, we document pervasive weak identification in applied work using IVs with time series data by surveying recent publications in leading economics journals. However, we find that strong identification often holds in large subsamples for which the instrument induces changes in the treatment. Motivated by this, we introduce a method based on dynamic programming to detect the most strongly-identified subsample and show how to use this subsample to improve estimation and inference. We also develop new identification-robust inference procedures that focus on the most strongly-identified subsample, offering efficiency gains relative to existing full sample identification-robust inference when identification fails over parts of the sample. Finally, we apply our results to heteroskedasticity-based identification of monetary policy effects. We find that about 60% of observations are compliers (i.e., cases where the variance of the policy shifts up on FOMC announcement days), and we fail to reject the exclusion restriction. Estimation using the most strongly-identified subsample helps reconcile conflicting IV and GMM estimates in the literature.

JEL Classification: B41, C12, C32.

Keywords: Compliers, Conditional inference, Exclusion, LATE.

^{*}We thank Isaiah Andrews, Josh Angrist, Jushan Bai, Elisa Facchetti, Jean-Jacques Forneron, Eva Janssen, Hiro Kaido, Toru Kitagawa, Sokbae (Simon) Lee, Daniel Lewis, Emi Nakamura, Serena Ng, Geert Mesters, Pepe Montiel Olea, Aureo de Paula, Andrew Patton, Mounu Prem, Zhongjun Qu, Bernard Salanié, Jón Steinsson, Jim Stock, Stephen Terry, Alex Torgovitsky and Kaspar Wüthrich for useful comments. We thank seminar participants at Bristol, BU, Cambridge, Columbia, Michigan, Michigan State, Oxford and UCL and conference participants at the NBER Summer Institute 2026. Casini acknowledges financial support from EIEF through 2022 EIEF Grant and from the Italian Ministry of Education and Research through FIS-2024-02854 grant. McCloskey acknowledges support from the National Science Foundation under Grant SES-2341730. The codes are available at <https://github.com/alessandro-casini/DynamicLATE-ts>.

1 Introduction

Economists work hard to extract plausibly exogenous variation in order to identify causal effects. Many identification strategies used in applied work either rely directly on instrumental variables (IVs) or can be reframed in terms of IV identification. This holds also in dynamic settings where, for example, external IVs may be constructed using a narrative approach or heteroskedasticity is exploited to yield additional identifying equations. Since [Imbens and Angrist \(1994\)](#), it has been well-known that IV-based approaches identify the local average treatment effect (LATE)—the average treatment effect for the sub-population of compliers, i.e., those whose treatment status is influenced by the policy intervention (the instrument).

This paper makes four main contributions. First, we show that compliers can be identified at the individual (time-period) level. Second, we demonstrate that the exclusion restriction can be tested. Third, we introduce a dynamic programming method to detect the most-strongly identified subsample when instrument relevance is time-varying, and we show how this subsample can be used to improve estimation. Fourth, we develop identification-robust inference procedures based on the strongly-identified subsample that are more efficient than existing full sample procedures when identification fails over parts of the sample. We discuss each contribution in turn.

In the LATE framework, the sub-population of compliers is unobserved. This means that although a LATE can be identified, the specific sample observations this effect represents is unknown. This limitation is often described informally as the inability to observe an observation’s treatment status under both the intervention and non-intervention scenarios. From a practical interpretability perspective, this presents a challenge that has been widely discussed in the literature [see, e.g., [Angrist, Imbens, and Rubin \(1996\)](#), [Heckman \(1996\)](#), [Heckman and Vytlacil \(2005\)](#) and [Imbens \(2010\)](#)]. Some progress has been made by [Imbens and Rubin \(1997\)](#) and [Abadie \(2003\)](#) who show that the proportion of compliers and some of their statistical characteristics can be identified provided the treatment is binary.

We can list several reasons why identifying the compliers individually is important for both policy analysis and empirical applications in time series settings. A lack of individual complier identification prevents researchers from determining when a policy transmission mechanism is active. For example, researchers cannot determine whether compliance occurs during recessions rather than expansions, nor can they identify which recessionary episodes give rise to compliance. Likewise, they cannot determine whether compliance occurs during periods of financial distress or constrained monetary policy, or which specific episodes gen-

erate compliance. In turn, this hinders the ability of policymakers to target interventions toward similar time periods. Moreover, identifying compliers can improve forecasting, as it allows one to predict *ex ante* which time periods are most affected by the policy. Whereas [Abadie \(2003\)](#) identifies complier-specific characteristics without identifying which sample observations are compliers, we identify whether individual time periods are compliers by using local comparisons to recover the relevant counterfactual treatment assignment. Finally, as emphasized by [Heckman and Vytlacil \(2007\)](#), it enhances external validity, since the LATE may differ from the average treatment effect in the broader population; without knowing who the compliers are, it is impossible to assess how far the estimated causal effect can be generalized beyond that subpopulation.

This paper considers IV identification in dynamic settings and shows how compliers can be identified individually by exploiting the structure of time series data. We first show that the notion of compliers can be equivalently rewritten in terms of an inequality involving the difference in means of the potential treatment under different instrument values. Under assumptions of continuity over time in the mean of the potential treatment assignment process—conditional on a fixed hypothetical value of the instrument—it is possible to recover counterfactual values by averaging observations in a neighborhood around a given time point.

The economic intuition behind the continuity assumption is that the policy transmission mechanism evolves gradually. For example, consider heteroskedasticity-based identification of monetary policy effects [cf. [Rigobon \(2003\)](#) and [Nakamura and Steinsson \(2018\)](#)] where the instrument indicates whether there was an FOMC announcement on each date in the sample and the treatment variable is equal to the variance of a short-term interest rate. Here compliers are defined as observations for which the volatility of the policy variable (change in short-term interest rate) increases if and only if there is an FOMC announcement. Suppose there is an FOMC announcement on a particular date. We then observe the treatment assignment under the realized instrument value but do not observe the counterfactual treatment assignment that would have occurred in the absence of the announcement. Continuity implies that the mean treatment assignment under the counterfactual state can be approximated using nearby non-announcement dates. More generally, continuity allows information from neighboring periods to recover the unobserved potential treatment assignments. Economically, this assumption reflects the idea that, e.g., firms do not suddenly stop responding to short-term interest rates, nor do institutional features of financial markets abruptly change from one day to the next. As a result, the complier status of the date in question can be es-

timated and tested by comparing local means of the treatment, one corresponding to nearby dates for which an announcement occurred and the other corresponding to nearby dates for which it did not.¹ Applying our identification results and tests to the heteroskedasticity-based identification of monetary policy, we find that about 60% of observations are compliers while the non-compliers are primarily concentrated in the early zero lower bound period, when the central bank could no longer lower interest rates and forward guidance was not aggressive. Here identification of the complier dates allows one to determine when the policy transmission mechanism is active and assess time variation in policy effectiveness.

Our identification argument is not confined to the case of FOMC announcements with high-frequency data. We show that it also extends to narrative identification approaches based on low-frequency data commonly used in empirical macroeconomics, such as [Ramey’s \(2011\)](#) identification of fiscal shocks, [Mertens and Ravn’s \(2013\)](#) identification of tax changes, and [Romer and Romer’s \(2004\)](#) identification of monetary policy shocks. For recent methodological work on narrative-based IV approaches see [Barnichon and Mesters \(2025\)](#).

Identification of compliers is not only valuable in its own right. The second main contribution of our paper is to show that this enables us to test the exclusion restriction, a key condition for valid IV estimation that is typically untestable in practice. By identifying compliers, and thus also non-compliers, we show that the exclusion restriction can be tested using a t -test that compares the average outcomes of non-compliers across different instrument values. This idea differs from that of [Kitagawa \(2015\)](#), who uses a nonnegativity condition on conditional distributions of observables to construct a Kolmogorov–Smirnov test of instrument validity for a binary treatment and a discrete instrument. Whereas Kitagawa’s approach is based on partial identification, our framework relies on exact identification and does not require either a binary treatment or a discrete instrument.

The third contribution of our paper is to establish new results and introduce new methods on instrument relevance, a key condition for LATE identification that requires a nontrivial correlation between the instrument and the endogenous variable. We begin by analyzing the problem of weak instruments, entailing low correlation between the endogenous variable and instrument, in empirical work through a survey of articles using IVs published from 2019 to 2022 in five leading journals: *American Economic Review*, *Econometrica*, *Journal of Political Economy*, *Quarterly Journal of Economics*, and *Review of Economic Studies*. Our sample in-

¹Our identification results immediately apply to cross-sectional settings with spatial data provided that the temporal distance between observations is interpreted as geographical distance, and analogous continuity assumptions are imposed over space.

cludes 1,560 specifications from 18 papers, with 199 involving time series and 1,361 involving panel data.² The left panels of Figure 1 show histograms of full sample first-stage F -statistics for the specifications in our survey, truncated above 100 for visibility. Many F -statistics concentrate around the χ_1^2 critical values and fall below the conventional thresholds of 10 and 23.1 suggested by Staiger and Stock (1997) and Montiel Olea and Pflueger (2013), raising serious concerns about weak instruments. These findings align with those of Andrews, Stock, and Sun (2019), who analyze cross-sectional studies. For example, we find that 75% of time series and 72% of panel data specifications have first stage F -statistics below 24. The median F -statistic is 12.63 for time series and 9.29 for panel data.³

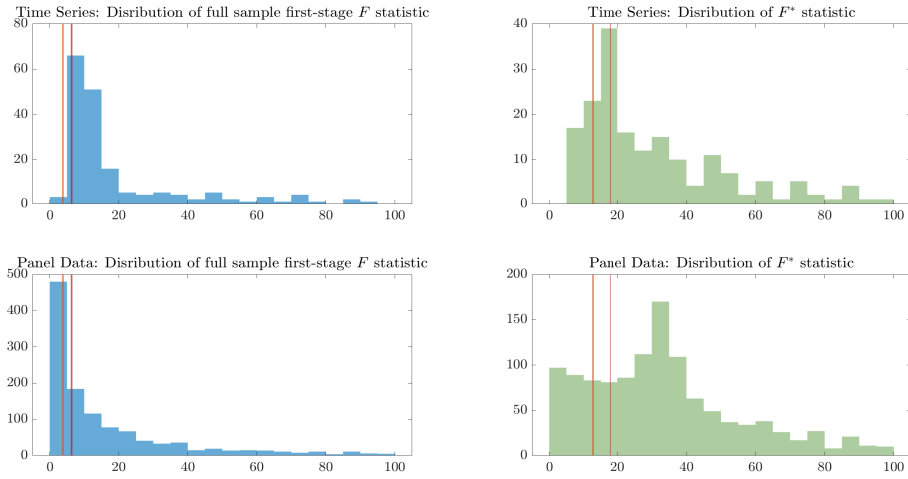


Figure 1: Distributions of the first-stage F (left panels) and F^* statistics (right panels). The top panels apply to time series specifications and the bottom panels apply to panel data specifications. The orange and red vertical lines correspond to the 5% and 1% level asymptotic critical values of the first-stage F (χ_1^2 for left panels) and F^* statistics (8.28 and 11.63) for right panels) under identification failure.

When identification fails or is weak, IV estimators can be severely biased for LATEs and conventional inference methods are rendered invalid. These problems have prompted extensive research on detecting weak instruments and constructing identification-robust confidence sets.⁴ However, there has been little work on estimation and inference in a general LATE

²See the supplement for the full list of papers and inclusion criteria.

³For panel data specifications we consider each cross-sectional unit individually to enable comparison to our proposed time series test shown on the right panels.

⁴See, e.g., Andrews et al. (2006), Kleibergen (2002), Moreira (2003) and Staiger and Stock (1997).

setting when identification may be stronger over subsamples. We develop a framework for identification, estimation, and inference on LATEs that accommodates time-varying instrument relevance. Within this framework, we propose a first-stage F -test to detect whether identification fails over all nontrivial subsamples. To solve the computationally intensive problem of searching for maximal identification strength among all possible sample partitions, we employ dynamic programming. This optimization is more complex than that in the structural break literature since evaluating identification strength requires more than comparing parameter changes across regimes.

In an attempt to understand the sources of weak IVs we plot the histograms of the F^* statistic proposed in this paper (cf. Section 4) in the right panels of Figure 1. The statistic F^* searches for the subsample with maximal identification strength among all possible subsamples of size at least $\pi_L T$.⁵ The idea is that while the IVs may appear weak in the full sample, they may be strong in a possibly large subsample. Figure 1 shows that this is indeed often the case. The red vertical lines in Figure 1 mark the 95th percentile of the asymptotic distributions of the F and F^* statistics under the null of identification failure. Although its quantiles are larger, the F^* statistics have substantially more mass in the upper quantiles of their null distribution. This has at least three implications. First, it confirms substantial time variation in the instruments' strength. Second, strong identification appears to be frequently present in a sizeable subsample even when the instruments appear weak in the full sample. The median F^* is 27.22 for time series and 33.81 for panel data specifications. These are substantially higher than their full sample counterparts and this difference cannot be simply attributed to the different null asymptotic distributions of the two test statistics given that the difference in the asymptotic critical values is relatively small while the empirical distributions of the two test statistics are markedly different. About one half of the specifications that appear to suffer from weak IVs in the full sample seem better characterized by strong IVs in the subsample with maximal identification strength. Third, the subsamples where instruments appear strong tend to be large. From an empirical perspective, this is encouraging: although weak instruments in the full sample are common, researchers can often succeed in identifying large subsamples where instruments appear strong.

Motivated by this survey evidence, we construct consistent estimators of LATEs when subsamples are strongly-identified. It is commonly believed that if IVs are strong only in some portion of the sample, the full sample IV estimator remains consistent for a LATE.

⁵We set $\pi_L = 0.6$ in Figure 1. We discuss the choice of π_L below.

We show that this belief is unwarranted unless the LATE of interest is time-invariant (i.e., homogeneous). If this condition fails, one can at best identify a LATE corresponding to the strongly-identified subsample. Even when the LATE is homogeneous, the full sample IV estimator may still be severely biased if instruments are irrelevant over parts of the sample.

Our approach differs from that of [Magnusson and Mavroeidis \(2014\)](#) and [Antoine and Boldea \(2018\)](#), who use time variation in IV strength to add moment conditions in a GMM context, enabling more efficient inference and estimation. In contrast, we exploit this time variation to identify the subsample where IVs are strongest and base our estimation on this subsample. This insight allows for *consistent estimation even when subsamples suffer from identification failure*.⁶ If the parameter of interest is heterogeneous, our estimator remains valid but is interpretable only within the strongly-identified sub-population.

We apply our methodology to the heteroskedasticity-based identification strategy used to estimate the causal effects of monetary policy from high-frequency data [e.g., [Nakamura and Steinsson \(2018\)](#)]. The key identification condition for this strategy is that the volatility of the daily changes in short-term interest rates is higher on FOMC announcement days than on non-FOMC days. [Lewis \(2022\)](#) provides evidence of weak full sample identification and shows that IV and GMM estimates even differ in sign. We find that identification is substantially stronger over a subsample comprising 80–90% of the data, with the excluded subsample centered around the financial crisis, during which volatility was high even on non-FOMC days. Estimation using the most strongly-identified subsample yields IV and GMM estimates that have the same sign and similar magnitudes. We recommend reporting the most strongly-identified subsample estimates in addition to the full sample estimates when strong full sample identification may be in question.

Although our new methods are able to find the most strongly-identified subsample, this subsample may still fail to be strongly-identified. For our final theoretical contribution, we develop identification-robust inference procedures using the most strongly-identified subsample. We propose versions of the Anderson-Rubin, Lagrange Multiplier, and conditional likelihood ratio tests, which depend only on this subsample. These tests are more efficient than their full sample counterparts, which include noise from regimes suffering from identification failure. When instruments are strong throughout the sample, our tests coincide with the conventional ones. When instruments are irrelevant over parts of the sample, our

⁶Another major difference from [Magnusson and Mavroeidis \(2014\)](#) is that we address the computational challenge for the case of multiple breaks in the first-stage coefficient. [Magnusson and Mavroeidis \(2014\)](#) did not attempt to address this issue and refer to it as “computationally demanding”.

tests achieve higher efficiency by focusing on stronger segments. In the worst case, when IVs are weak everywhere, our methods are no less efficient than existing ones. While there is a trade-off between using fewer observations and more strongly identified subsamples, simulations show that our tests have higher power, indicating that the efficiency loss from a smaller sample size is outweighed by the gain in identification strength.

The paper is organized as follows. Section 2 introduces the potential outcome framework and dynamic causal effects, and presents identification results. Section 3 discusses issues pertaining to heteroskedasticity-based identification of monetary policy. Section 4 presents an F -test for full sample identification failure. Estimation and inference robust to weak identification are discussed in Sections 5-6. An empirical application is considered in Section 7. Section 8 concludes. The supplements Casini, McCloskey, Rolla, and Pala (2026b, 2026a) include the Monte Carlo simulations, proofs and additional results.

2 Identification of Dynamic Causal Effects

A growing literature in macroeconomics uses IVs to identify dynamic causal effects when the policy variable of interest is endogenous.⁷ Many existing identification approaches can be reframed in terms of IVs, either derived from the modeling approach [e.g., heteroskedasticity-based identification as in Rigobon (2003) and Nakamura and Steinsson (2018)] or through external IVs constructed using a narrative approach [cf., Montiel Olea, Stock, and Watson (2021)]. For example, Romer and Romer (1989) study the FOMC minutes to pinpoint dates when monetary policy actions were arguably exogenous. This allows the construction of exogenous variables that can be interpreted as IVs for some structural shock of interest.⁸

We adopt a potential outcomes framework, as introduced by Rubin (1974) and extended to time series settings by Angrist and Kuersteiner (2011) and Rambachan and Shephard (2021). Let the stochastic process $V_t = (Y_t, X_t, D_t, Z_t)$ be defined on the probability space $(\Omega, \mathcal{F}, \mathbb{P})$, where Y_t is a vector of outcome variables, D_t is a policy variable, X_t is a vector of other exogenous and/or lagged endogenous variables, and Z_t is a vector of instruments. Let $\vec{X}_t = \{\dots, X_{t-1}, X_t\}$ denote the covariate path up to time t , with analogous definitions for $\vec{Y}_t$, $\vec{D}_t$ and $\vec{Z}_t$. Let the policy-relevant information set at time t denoted by $\mathcal{F}_t = \sigma(\vec{V}_t)$

⁷See, e.g., Gertler and Karadi (2015), Mertens and Ravn (2013), Plagborg-Møller and Wolf (2022) and Stock and Watson (2018).

⁸See Ramey (2011) for unanticipated defense spending shocks, Nakamura and Steinsson (2018) and Romer and Romer (2004) for monetary policy shocks, Känzig (2021) and Kilian (2009) for oil market shocks, and Romer and Romer (2010) for tax shocks.

where $\sigma(\tilde{V}_t)$ is the σ -algebra generated by the history of V_t , $\tilde{V}_t = (\vec{Y}_{t-1}, \vec{X}_t, \vec{D}_{t-1}, \vec{Z}_{t-1})$.

Policy decisions, D_t , depend on past observable variables and the contemporaneous outcome through a systematic component and on idiosyncratic information available to the policy-maker (i.e., the random component). The systematic component is a time-varying non-stochastic function of the observed random variables $\tilde{V}_t$, contemporaneous outcome Y_t , and the contemporaneous instrument Z_t . The idiosyncratic information is represented by a scalar stochastic shock e_t that is not observed by the researcher. In a SVAR context, e_t is the structural shock to the policy variable D_t . For example, if the monetary authority follows a simple Taylor rule for the nominal interest rate, then D_t is the policy interest rate and $\tilde{V}_t$ includes inflation, output and the natural rate of interest.

We define two types of potential outcomes. The first, $Y_t((\epsilon_{1:t}), (z_{1:t}))$, denotes the counterfactual values of Y_t under hypothetical sequences of the policy shocks $\epsilon_{1:t}$ and instruments $z_{1:t}$, where $a_{1:t} = \{a_s\}_{s=1}^t$.

Definition 2.1. A generalized potential outcome, $Y_t((\epsilon_{1:t}), (z_{1:t}))$, is defined as the value assumed by Y_t if $e_s = \epsilon_s$ and $Z_s = z_s$ for $s = 1, \dots, t$.

This definition excludes dependence on future shocks or instruments. The potential outcome process should not be confused with the observed outcome $\{Y_t\}_{t \geq 1} = \{Y_t(e_{1:t}, Z_{1:t})\}_{t \geq 1}$. For $h \geq 0$ and any given ϵ and z , write the time- $t+h$ potential outcome along the path $((e_{1:t-1}, \epsilon, e_{t+1:t+h}), (Z_{1:t-1}, z, Z_{t+1:t+h}))$ as

$$Y_{t,h}(\epsilon, z) = Y_{t+h}((e_{1:t-1}, \epsilon, e_{t+1:t+h}), (Z_{1:t-1}, z, Z_{t+1:t+h})),$$

where $Y_{t,h}(e_t, Z_t) = Y_{t+h}$. Definition 2.1 captures the property that $Y_{t,h}(\epsilon, z)$ also depends on policy shocks that occur between time $t+1$ and $t+h$. The notation $Y_{t,h}(e, z)$ focuses on the effect of a single policy shock on current and future outcomes akin to the idea underlying an impulse response. When the potential outcomes do not depend on the instruments, $Y_{t,h}(\epsilon, z) = Y_{t,h}(\epsilon)$, and for $\epsilon \neq \epsilon'$, $Y_{t,h}(\epsilon) - Y_{t,h}(\epsilon')$ for $h = 0, 1, \dots$ are the dynamic causal effects of a policy shock on the outcome. In a SVAR setting, one is often interested in these dynamic causal effects which are in fact the impulse responses.

The second potential outcome that we discuss, $Y_t^*((d_{1:t}), (z_{1:t}))$, is defined as the counterfactual values of Y_t under hypothetical sequences of treatments $d_{1:t}$ and instruments $z_{1:t}$. The distinction with $Y_t(\epsilon_{1:t}, z_{1:t})$ is that this formulation focuses on causal effects of the policy variable D , not the policy shock e . For $t \geq 1$, we assume that $d_t \in \mathbf{D}$, $z_t \in \mathbf{Z}$ for some sets

D and **Z**. In many applications outside SVARs, the causal effects of the policy are of interest. Think about the slope of demand functions, price elasticities, response coefficients or reaction functions of, for example, asset prices to monetary policy, and so on. Typically these causal effects are analyzed using event-studies, quasi-experiments, IV regressions, etc. The recent literature on causal effects in time series [e.g., [Rambachan and Shephard \(2021\)](#)] focuses on the identification of causal effects of the structural shocks. In this paper, we consider identification of causal effects of the policy variable. We illustrate the difference between these two causal effects and an application to SVAR using the following two examples.

Example 2.1. Consider the following system of simultaneous equations,

$$Y_t = \beta D_t + \eta_t \quad \text{and} \quad D_t = aY_t + e_t, \quad (2.1)$$

where the first equation is the demand curve, the second is the supply curve, Y_t and D_t are the observed price and quantity, and η_t and e_t are the structural shocks. The parameter β captures the slope of the demand function, which corresponds to the causal effect $\partial Y_t^*(d) / \partial d = \beta$. On the other hand, in a SVAR context one may be interested in the impulse response of Y_t given a shock to supply e_t . Solving for the reduced-form of (2.1),

$$Y_t = \frac{\beta}{1 - a\beta} e_t + \frac{1}{1 - a\beta} \eta_t,$$

shows that the lag-0 impulse response is $dY_t(e) / de = \beta / (1 - a\beta)$, which differs from β .

Example 2.2. Consider the following reduced-form VAR,

$$V_t = A_1 V_{t-1} + A_2 V_{t-2} + \dots + A_p V_{t-p} + u_t,$$

where $V_t = (D_t, Y_t)'$ is $n \times 1$, D_t is a scalar, and u_t is a vector of reduced-form VAR innovations. The latter are related to structural shocks, $\varepsilon_t = (e_t, \eta_t)'$, via $u_t = B_0 \varepsilon_t$ where B_0 is a non-singular matrix. Under suitable conditions, V_t admits a moving-average representation $V_t = \sum_{j=0}^{\infty} C_j(A) B_0 \varepsilon_{t-j}$, where $C_j(A) = \sum_{i=1}^j C_{j-i}(A) A_i$ for $j = 1, 2, \dots$ with $C_0(A) = I_n$ and $A_i = 0$ for $i > p$. Then, the outcome variable admits a moving-average representation,

$$Y_t = \sum_{j=0}^{\infty} c_{ye,j} e_{t-j} + \sum_{j=0}^{\infty} c_{y\eta,j} \eta_{t-j},$$

where $c_{ye,j}$ and $c_{y\eta,j}$ are blocks of $C_j(A) B_0$ partitioned conformably to Y_t , e_t and η_t . If e_t is

the policy shock, the potential outcomes here are defined as

$$Y_{t,h}(\epsilon) = Y_{t,h}(\epsilon, z) = \sum_{j=0, j \neq h}^{\infty} c_{ye,j} e_{t+h-j} + \sum_{j=0}^{\infty} c_{y\eta,j} \eta_{t+h-j} + c_{ye,h} \epsilon.$$

The potential outcome $Y_{t,h}(\epsilon)$ tells us what Y_{t+h} would be if $e_t = \epsilon$ and it does not depend upon z since the instrument Z_t is excluded from the VAR. Here the absence of causal effects means that $c_{ye,h} = 0$ for all h , coinciding with the canonical condition that the impulse responses are identically equal to zero.

The potential outcome framework is useful because it allows the study of nonparametric conditions such that common statistical estimands (e.g., impulse responses) have a causal interpretation. [Montiel Olea, Stock, and Watson \(2021\)](#) show how to use the instrument Z_t to identify the impulse response coefficient $\phi_{r,e,h} = \partial Y_{t+h}^{(r)} / \partial e_t$ (the effect of e_t on the r th variable in Y_{t+h}). From the moving-average representation we have $\phi_{r,e,h} = \iota'_{r+1} C_h(A) B_0 \iota_1$ where ι_s denotes the s -th standard basis vector. This shows that $\phi_{r,e,h}$ depends on the A 's and the first column of B_0 . The following assumptions are needed for the identification of $\phi_{r,e,h}$: (i) $\mathbb{E}(Z_t e_t) = \theta \neq 0$ (instrument relevance) and (ii) $\mathbb{E}(Z_t \eta_t) = 0$ (instrument exogeneity). By (i)-(ii), $B_0^{(:,1)} = B_0 \iota_1$ is identified up to scale by the covariance between Z_t and the reduced-form innovations u_t : $\Gamma = \mathbb{E}(Z_t u_t) = \mathbb{E}(Z_t B_0 \varepsilon_t) = \theta B_0^{(:,1)}$. Using the scale normalization $B_0^{(1,1)} = 1$ [see [Stock and Watson \(2018\)](#) for a discussion] we have $\Gamma^{(1,1)} = \mathbb{E}(Z_t e_t) = \theta$ and $B_0^{(:,1)} = \Gamma / \Gamma^{(1,1)} = \Gamma / \iota'_1 \Gamma$. It follows that $\phi_{r,e,h}$ is identified since $\phi_{r,e,h} = \iota'_{r+1} C_h(A) \Gamma / \iota'_1 \Gamma$, where A can be estimated consistently from the reduced-form VAR and Γ can be estimated consistently by using the VAR residuals $\hat{u}_t$ in place of u_t . On the other hand, identifying the causal effects of the policy D_t here would require additional identification restrictions.

[Montiel Olea, Stock, and Watson \(2021\)](#) use shortfalls in OPEC oil production associated with wars and civil disruptions as an instrument for the oil supply shock in the SVAR of [Kilian \(2009\)](#) who investigates the effect of oil supply and demand shocks on oil production and prices. This variable is plausibly correlated with the oil supply shock and, because the shortfalls are associated with political events such as wars in the Middle East, it is plausibly uncorrelated with the demand shocks. Using the analog of the nonparametric conditions we provide below, applied to the shock e_t rather than the policy D_t , permits a causal interpretation of the impulse response even when $\mathbb{E}(Z_t e_t) = 0$ for a sub-population.

In the following, we discuss identification of causal effects of the policy via IV estimands.

2.1 Identification Conditions

We explicitly allow for endogeneity and rely on IVs. We define $D_t(z)$ as the potential treatment assignment at time t when Z_t is set equal to $z \in \mathbf{Z}$. The instrument Z_t is assumed to be (conditionally) independent of the potential outcomes $Y_{t,j}^*(d, z)$ and treatments $D_t(z)$ but correlated with the observed treatment D_t .

Assumption 2.1. (*Independence*) For all $d \in \mathbf{D}$, $z \in \mathbf{Z}$ and $t \geq 1$, we have

$$\left\{ \left\{ Y_{t,h}^*(d, z) \right\}_{h \geq 0}, D_t(z) \right\} \perp_{Z_t} \tilde{V}_t. \quad (2.2)$$

Assumption 2.1 states that, given $\tilde{V}_t$, the instrument is as good as randomly assigned. In studies of monetary policy such as Nakamura and Steinsson (2018), $Z_t = 1$ if there is an FOMC announcement on day t , and $Z_t = 0$ otherwise. Since FOMC announcements are scheduled in advance, Z_t is deterministic, and therefore Assumption 2.1 holds.⁹

The second assumption is that potential outcomes $Y_{t,h}^*(d, z)$ are a function of d but not of z . In the case of FOMC announcements, this means that potential realizations of expected output growth respond to changes in the monetary policy variable in the same way regardless of whether the change is associated with an FOMC announcement or not.

Assumption 2.2. (*Exclusion*) For all $d \in \mathbf{D}$, $t \geq 1$ and $h \geq 0$, we have

$$\left\{ Y_{t,h}^*(d, z) = Y_{t,h}^*(d, z') \right\} \mid \tilde{V}_t, \quad \text{for all } z, z' \in \mathbf{Z}. \quad (2.3)$$

In a dynamic simultaneous equations model (e.g., a SVAR) the exclusion restriction requires the instrument not to appear in the causal equation of interest. In Example 2.2, Assumption 2.2 corresponds to condition (ii), i.e., $\mathbb{E}(Z_t \eta_t) = 0$ where η_t is composed of the structural shocks other than e_t . Under Assumption 2.2 we write $Y_{t,h}^*(d, z) = Y_{t,h}^*(d)$.

Identification based on IVs requires instrument relevance or “existence of a first-stage”. The latter means that $\mathbb{E}(D_t(z) \mid \tilde{V}_t)$ is a non-trivial function of z . In cross-sectional settings, the existence of a first-stage is typically assumed to hold for all units to guarantee strong identification. Strong identification of this form often fails to hold in applications involving time series data due to temporary misspecification, bad luck, rare events or parameter instability. The analysis based on articles in five leading journals that we report earlier suggests that there are time periods for which the first-stage exists (strong identification) and others

⁹Unscheduled FOMC meetings that occur during emergencies are typically excluded from the sample.

for which it does not (identification failure). Standard first-stage F -tests are then likely to indicate weak identification since they are based on averaging these two sub-populations.

We provide a theoretical framework to address this identification problem by assuming that there are two sub-populations. One comprises a fraction $\pi_0 \in [0, 1]$ of the overall population for which the first-stage exists. For the second sub-population, which comprises a fraction $1 - \pi_0$ of the population, the first-stage does not exist. This leads to a new notion of LATE, which we name π -LATE, the LATE for the (unknown) π_0 fraction of the population for which the first-stage exists. If $\pi_0 = 1$, then one recovers LATE.

Denote by $|\mathbf{S}_{0,T}|$ the cardinality of $\mathbf{S}_{0,T}$ (i.e., the number of indices in $\mathbf{S}_{0,T}$).

Assumption 2.3. (*Partial first-stage*) Assume there exists $\mathbf{S}_{0,T} \subseteq \{1, \dots, T\}$ such that $|\mathbf{S}_{0,T}| = \lfloor \pi_0 T \rfloor$ with $\pi_0 \in (0, 1]$ and for $t \in \mathbf{S}_{0,T}$, $\mathbb{E}(D_t(z) | \tilde{V}_t)$ is a non-trivial function of z .¹⁰

Assumption 2.3 implies that there are two sub-populations: one for which the first-stage exists and one for which it does not. An average treatment effect can only be identified via IVs for the fraction π_0 of the population for which a first-stage exists.

The next assumption is monotonicity which, under heteroskedasticity-based identification of monetary policy (see Section 3), means that while for some days the FOMC announcement does not coincide with higher volatility in the policy variable, all of those days in which the announcement affects the volatility of the policy variable, volatility is shifted up.

Assumption 2.4. (*Monotonicity*) $\mathbf{D} \subseteq \mathbb{R}$. For all $t = 1, \dots, T$ and $z, z' \in \mathbf{Z}$, either $D_t(z) \geq D_t(z')$ or $D_t(z') \geq D_t(z)$ with probability one.

If $\pi_0 = 1$ (so $|\mathbf{S}_{0,T}| = T$), the condition reduces to that in Imbens and Angrist (1994).

Following Kolesár and Plagborg-Møller (2025), we impose the following assumption.

Assumption 2.5. For all $t \geq 1$ and $h \geq 0$, (i) $Y_{t,h}^*(\cdot)$ is locally absolutely continuous on $\mathbf{D}$ and (ii) $\mathbb{E} \left[\int_{\mathbf{D}} |\partial Y_{t,h}^*(d) / \partial d| dd \middle| \tilde{V}_t \right] < \infty$.

Assumption 2.5 allows D_t to be either discrete, continuous or mixed. When D_t is discrete or mixed, it is implicitly assumed that to deal with the gaps in the support of D_t one extends $Y_{t,h}^*(\cdot)$ to $\mathbf{D}$ such that the extension is locally absolutely continuous. The support of D_t is allowed to be unbounded. These conditions are weaker than counterparts imposed in the recent literature [cf. Casini and McCloskey (2025) and Rambachan and Shephard (2021)], in

¹⁰We assume that all expectations exist.

particular local absolute continuity replaces differentiability of $Y_{t,h}^*(\cdot)$ plus bounded support of D_t . It allows the application of the fundamental theorem of calculus to $Y_{t,h}^*(\cdot)$ without requiring the support of D_t to be bounded.

2.2 Identification Results

2.2.1 Identification of Causal Effects

We first discuss the case of a discrete instrument. When the first-stage does not exist for all t , it is useful to define an IV estimand corresponding to the sub-population for which it does. Let the generalized Wald estimand be defined for all $z', z \in \mathbf{Z}$ by

$$\beta_{\pi,t,h}(\tilde{v}) = \frac{\mathbb{E}(Y_{t+h} | Z_t = z', \tilde{V}_t = \tilde{v}) - \mathbb{E}(Y_{t+h} | Z_t = z, \tilde{V}_t = \tilde{v})}{\mathbb{E}(D_t | Z_t = z', \tilde{V}_t = \tilde{v}) - \mathbb{E}(D_t | Z_t = z, \tilde{V}_t = \tilde{v})}, \quad \text{for } t \in \mathbf{S}_{0,T}, \quad (2.4)$$

where $\tilde{v} \in \mathbf{V}$. This is the ratio of a reduced-form generalized impulse response to a first-stage generalized impulse response for $t \in \mathbf{S}_{0,T}$. We show that for $t \in \mathbf{S}_{0,T}$, the estimand $\beta_{\pi,t,h}(\tilde{v})$ identifies a weighted average of causal effects for the compliers. Recall that $t \in \mathbf{S}_{0,T}$ and π_0 are related by $|\mathbf{S}_{0,T}| = \lfloor \pi_0 T \rfloor$. When $\pi_0 = 1$ and there is no conditioning on $\tilde{V}_t = \tilde{v}$, $\beta_{1,t,h}$ reduces to the Wald estimand considered by [Rambachan and Shephard \(2021\)](#). For $t \notin \mathbf{S}_{0,T}$, $\beta_{\pi,t,h}$ does not identify a causal effect because the denominator of (2.4) is equal to zero.

We show that for $t \in \mathbf{S}_{0,T}$, the generalized Wald estimand is equal to a weighted average of marginal effects where the latter are the derivatives $\partial Y_{t,h}^*(d) / \partial d$.

Proposition 2.1. (π -LATE) *Let Assumptions 2.1-2.5 hold. For $t \in \mathbf{S}_{0,T}$, $h \geq 0$, $\tilde{v} \in \mathbf{V}$ and $z', z \in \mathbf{Z}$, we have*

$$\begin{aligned} \beta_{\pi,t,h}(\tilde{v}) &= \int_{\mathbf{D}} \mathbb{E} \left[\frac{\partial Y_{t,h}^*(d)}{\partial d} \middle| D_t(z) \leq d \leq D_t(z'), \tilde{V}_t = \tilde{v} \right] w_t(d | \tilde{v}) dd, \quad \text{where} \quad (2.5) \\ w_t(d | \tilde{v}) &= \frac{\mathbb{P}(D_t(z) \leq d \leq D_t(z') | \tilde{V}_t = \tilde{v})}{\int_{\mathbf{D}} \mathbb{P}(D_t(z) \leq r \leq D_t(z') | \tilde{V}_t = \tilde{v}) dr} \geq 0 \quad \text{and} \quad \int_{\mathbf{D}} w_t(d | \tilde{v}) dd = 1. \end{aligned}$$

Proposition 2.1 shows that $\beta_{\pi,t,h}(\tilde{v})$ identifies a weighted average of causal effects for compliers, characterized by $D_t(z') > D_t(z)$, for observations with a first-stage, with weights $w_t(d | \tilde{v})$ determined by the (conditional) likelihood that $D_t(z) \leq d \leq D_t(z')$. We refer to the average treatment effect on the right-hand side of (2.5) as the time- t π -LATE since

it is the LATE for the observations in this sub-population, which is a fraction π_0 of the whole population. In practice, the IV estimand $\beta_{\pi,t,h}(\tilde{v})$ is characterized by two types of averaging. First, there is averaging over time. For any treatment d , the average involves only those observations whose treatment variable can be induced to change by a change in the instrument and is computed only over those observations that satisfy the first-stage (i.e., $t \in \mathbf{S}_{0,T}$). The second averaging is over different treatment values d at the same date t . This is reflected in the weight $w_t(\cdot)$, which is proportional to the conditional probability that $D_t(z) \leq d \leq D_t(z')$.

Stationarity of the conditional joint distribution of the average potential outcome and treatment assignment functions for observations with a first-stage lends further interpretability to the generalized Wald estimand. Specifically, if $\{Y_{t,h}^*(d), D_t(z)\}|\tilde{V}_t$ is identically distributed across t for all $t \in \mathbf{S}_{0,T}$, $d \in \mathbf{D}$ and $z \in \mathbf{Z}$, Proposition 2.1, immediately implies that $\beta_{\pi,t,h}$ is equal for all $t \in \mathbf{S}_{0,T}$. Given this, we can write $\beta_{\pi,t,h} = \beta_{\pi,h}$, making explicit that the generalized Wald estimand (2.4) equals a weighted average of causal effects for members of the sub-population with a first-stage, which represents a π_0 -sized fraction of the total population. Under this assumption, we refer to the average causal effect inside of the integral as π -LATE since it is a LATE for a member of the $\mathbf{S}_{0,T}$ sub-population whose treatment variable can be induced to change by a change in the instrument.

The sample counterpart to the generalized Wald estimand (2.4) involves replacing the conditional expectations with sample estimates based upon observations $t \in \mathbf{S}_{0,T}$, yielding an estimator of a causal effect. When Assumption 2.3 holds with $\pi_0 \in (0, 1)$, the full sample estimand, i.e., the ratio of the time averages of the numerator and denominator of (2.4), is a poor representative of the full sample average treatment effects because it includes observations for which the instrument is not relevant in the averaging. We caution that the usual practice of estimating the conditional expectations in (2.4) with full sample estimates will not estimate the full sample LATE, but π -LATE.

Angrist, Graddy, and Imbens (2000), Kolesár and Plagborg-Møller (2025) and Ram-bachan and Shephard (2021) consider related results. The difference here is that we do not require D_t to be continuous or that the first-stage holds for all t .¹¹

A connection to program evaluation with binary policy actions arises when we map a dynamic problem with continuous variables into one with binary policy actions and instruments. For example, consider the analysis of causal effects of monetary policy using

¹¹Sojitra and Syrgkanis (2025) focus is on the causal effect of treatment histories on long term outcomes, rather than of one-time shocks or single policy shifts on outcomes at horizon h .

heteroskedasticity-based identification [cf. Nakamura and Steinsson (2018) and Rigobon and Sack (2003)]. Define a binary instrument Z_t with $Z_t = 1$ if there is a scheduled announcement on day t and $Z_t = 0$ otherwise. The policy Δi_t typically reflects changes in short-term interest rates. Identification relies on higher volatility in Δi_t during announcement days (policy sample) compared to non-announcement days (control sample). Think about mapping $|\Delta i_t|$ into a binary treatment such that $D_t = 1$ if $|\Delta i_t| \geq \delta$ for some threshold $\delta > 0$ and $D_t = 0$ if $|\Delta i_t| < \delta$ [cf. Rigobon and Sack (2003)]. Here π -LATE captures the average treatment effect for the sub-population whose interest rate changes exceed δ only when there is an announcement (i.e., when $Z_t = 1$). Observations where $|\Delta i_t| < \delta$ regardless of announcements are “never-takers,” while those with $|\Delta i_t| \geq \delta$ regardless of announcements are “always-takers.” Under monotonicity, these groups form the non-compliers, whose responses are driven by idiosyncratic factors other than announcement-specific effects. In Section 3 we document regimes where the volatility of Δi_t is high even in the absence of announcements.

2.2.2 Identification of Compliers and Exclusion Restriction

A practical challenge for the π -LATE framework, and LATE frameworks in general, is that the sub-population of compliers is unknown. However, in time series settings with binary instruments, we show below that one can identify the compliers individually, i.e., to determine whether each observation t is a complier. In this section, we consider a binary instrument, e.g., $Z_t = 1$ if t is an FOMC meeting day and $Z_t = 0$ otherwise. Under Assumption 2.4, assume without loss of generality that $D_t(1) \geq D_t(0)$ for all t . Then, observation t_0 is a complier if and only if $D_{t_0}(1) > D_{t_0}(0)$ with probability one—if the treatment changes in response to the instrument.

We begin with the following assumption which states that each observation is either a complier or a non-complier with certainty.

Assumption 2.6. (*Deterministic complier status*) For each t either $\mathbb{P}(D_t(1) > D_t(0)) = 1$ or $\mathbb{P}(D_t(1) > D_t(0)) = 0$.

Assumption 2.6 rules out cases where $\mathbb{P}(D_t(1) > D_t(0)) = p$ for some $p \in (0, 1)$. A non-complier cannot be characterized by $\mathbb{P}(D_t(1) > D_t(0)) > 0$. The latter probability must be zero. It allows for continuous treatment variables. For example, $D_t(0)$ may be an arbitrary continuous random variable, and $D_t(1) = D_t(0) + c$ where $c \geq 0$. Under Assumption 2.6, Lemma S.D.2 in the supplement shows that $\mathbb{P}(D_{t_0}(1) > D_{t_0}(0)) = 1$ is equivalent

to $\mathbb{E}(D_{t_0}(1)) > \mathbb{E}(D_{t_0}(0))$. This equivalence implies that compliers can be identified by comparing the expected treatment values under different instrument values.¹² Under mild smoothness assumptions that we discuss below, the latter two expected values can be estimated consistently from the sample so that we can determine whether t_0 is a complier in large samples by looking at the corresponding inequality based on sample quantities.

Let $\mathbf{P} \subset \{1, \dots, T\}$ denote the “policy sample”, the set of observations for which $Z_t = 1$ so that $D_t = D_t(1)$ for all $t \in \mathbf{P}$, and let $\mathbf{C} = \{1, \dots, T\} \setminus \mathbf{P}$ denote the “control sample”, where $D_t = D_t(0)$. It is reasonable to assume that, for a given value of the instrument, the potential treatment assignments vary smoothly over time. Suppose we wish to determine whether an observation $t_0 \in \mathbf{P}$ is a complier. Since $D_{t_0}(0)$ is not observed, under time-smoothness we approximate $\mathbb{E}(D_{t_0}(0))$ by averaging nearby observations in the control sample. Letting $N_0(t_0)$ denote the n_0 largest indices $s \in \mathbf{C}$ such that $s \leq t_0 - 1$, this implies

$$\overline{D}_{C,t_0-1,n_0} \equiv n_0^{-1} \sum_{s \in N_0(t_0)} D_s \xrightarrow{\mathbb{P}} \mathbb{E}(D_{t_0-1}(0))$$

as $n_0 \rightarrow \infty$ with $n_0/|\mathbf{C}| \rightarrow 0$ under mild conditions. This is a standard nonparametric condition: the window size n_0 must grow, but slowly relative to the size of $\mathbf{C}$, so that the averaging neighborhood becomes dense on the $\mathbf{C}$ while remaining asymptotically local. In addition, it follows that $\mathbb{E}(D_{t_0-1}(0))$ is close to $\mathbb{E}(D_{t_0}(0))$. A similar argument can be applied to $\mathbb{E}(D_{t_0}(1))$ using adjacent days in the policy sample: we have $\overline{D}_{P,t_0,n_1} \xrightarrow{\mathbb{P}} \mathbb{E}(D_{t_0}(1))$ as $n_1 \rightarrow \infty$ with $n_1/|\mathbf{P}| \rightarrow 0$, where $\overline{D}_{P,t_0,n_1} = n_1^{-1} \sum_{s \in N_1(t_0)} D_s$ and $N_1(t_0)$ denotes the n_1 largest indices $s \in \mathbf{P}$ such that $s \leq t_0$. Thus, observation $t_0 \in \mathbf{P}$ is a complier if and only if $\overline{D}_{P,t_0,n_1} - \overline{D}_{C,t_0-1,n_0} \xrightarrow{\mathbb{P}} c$ as $n_0, n_1 \rightarrow \infty$ with $n_0/|\mathbf{C}|, n_1/|\mathbf{P}| \rightarrow 0$ for $c > 0$.

Intuitively, even though $D_{t_0}(0)$ is not observed when $t_0 \in \mathbf{P}$, observations close to t_0 characterized by no FOMC announcement provide information about what $\mathbb{E}(D_{t_0}(0))$ would have been in the absence of an FOMC announcement.¹³ There are about six weeks in between any two FOMC meetings, and so $n_0 \approx 30$. Alternatively, following [Nakamura and Steinsson \(2018\)](#) the control sample could include all Tuesdays and Wednesdays that are not FOMC meeting days. Nevertheless, one can skip the observation that per-

¹²Note that this result is different from that in Lemma 2.1 in [Abadie \(2003\)](#) who shows that under several assumptions the proportion of compliers can be identified by $\mathbb{E}(D_i(1)) - \mathbb{E}(D_i(0))$ in a cross-sectional setting. He uses this lemma to show that any statistical characteristic that can be defined in terms of moments of the joint distribution of (Y_i, D_i, Z_i) is identified for compliers. He then remarks that it is not possible to identify compliers individually under these assumptions.

¹³One could also use the observations to the right of t_0 to construct $\overline{D}_{C,t_0+1,n}$, i.e., $D_{t_0+1}, \dots, D_{t_0+n}$.

tains to the previous meeting, say $D_{t-1}(0)$, whose realization is not observed, and continue averaging using the observations prior to that meeting as well to construct the average $\overline{D}_{C,t_0-1,n_0}$ possibly applying down-weighting for observations further in time from t_0 , i.e., use $\dots, D_{t-1-1}, D_{t-1+1}, D_{t-1+2} \dots, D_{t_0-2}, D_{t_0-1}$. Similarly, observations in $\mathbf{P}$ close to t_0 provide information about what $\mathbb{E}(D_{t_0}(1))$ would have been, though here the successive observations are separated chronologically by the observations in the control sample $\mathbf{C}$.

The economic intuition behind this continuity assumption is that the latent determinants of treatment take-up evolve gradually over time. Firms do not suddenly stop responding to short-term interest rates, and households do not suddenly become insensitive to borrowing costs. Consequently, the propensity to receive treatment is expected to change smoothly.

We now present the formal result for identification of the compliers. The following two assumptions can be justified in large samples when the mean (potential) treatment assignments in both the control and policy samples vary smoothly over time. Under an infill asymptotic embedding where the original observations indexed by $t = 1, \dots, T$ are mapped into the unit interval $[0, 1]$ via $u = t/T$, if $\lim_{T \rightarrow \infty} \mathbb{E}(D_{Tu}(z))$ is continuous in u under a fixed instrument value $z \in \mathbf{Z}$, the following assumptions hold. This type of continuity accommodates general forms of smoothly time-varying means but not abrupt breaks in mean.¹⁴

Assumption 2.7. (i) For any $t \in \mathbf{C}$, $\overline{D}_{C,t,n} \xrightarrow{\mathbb{P}} \mathbb{E}(D_t)$ as $n \rightarrow \infty$ with $n/|\mathbf{C}| \rightarrow 0$. (ii) For $t \in \mathbf{P}$ $\mathbb{E}(D_{t-1}(0)) = \mathbb{E}(D_t(0))$.

Assumption 2.8. (i) For any $t \in \mathbf{P}$ $\overline{D}_{P,t,n} \xrightarrow{\mathbb{P}} \mathbb{E}(D_t)$ as $n \rightarrow \infty$ with $n/|\mathbf{P}| \rightarrow 0$. (ii) For $t \in \mathbf{C}$ $\mathbb{E}(D_t(1)) = \mathbb{E}(D_{s^*(t)}(1))$ where $s^*(t) = \operatorname{argmin}_{s \in \mathbf{P}} |t - s|$ is the closest index in $\mathbf{P}$ to t .

Assumptions 2.7-2.8 impose smoothness restrictions ensuring that the mean of the potential treatment assignment evolves gradually over time under both the no-announcement and announcement trajectories. Under an infill asymptotic embedding, these conditions guarantee that local averages of the observed treatment variable consistently recover the underlying conditional means. Assumption 2.7(i) requires a law of large numbers to apply to the rolling-window sample average of D_t at the points of continuity of $\mathbb{E}(D_t)$. It is a minimal technical assumption. Assumption 2.7(ii) strengthens part (i) a bit by requiring that for $t \in \mathbf{P}$ the potential treatment assignment under the trajectory $Z_t = 0$ has a locally constant mean. Assumption 2.8(i) adapts Assumption 2.7(i) to the observations in $\mathbf{P}$. This

¹⁴However, breaks in the mean of the assignment process can be estimated under some conditions as we explain below. Then, time-smoothness is required to hold only in regimes defined by successive break dates.

is a stronger assumption since two successive observations in the policy sample are separated by several observations in the control sample. Assumption 2.8(ii) requires that $\mathbb{E}(D_t(1))$ for $t \in \mathbf{C}$ is equal to the mean of the potential treatment assignment at the closest date in the policy sample $s^*(t)$. This is a first moment constancy assumption on the potential treatment assignment under the trajectory $Z_t = 1$. Assumption 2.7 is used to identify the compliers in $\mathbf{P}$, while Assumption 2.8 is used to identify the compliers in $\mathbf{C}$.

Assumptions 2.7-2.8 are implied by stationarity and local stationarity conditions commonly imposed in the analysis of economic time series. Together, these assumptions formalize the idea that observations close in time provide reliable information about the counterfactual evolution of the treatment process. This, in turn, enables identification of compliers through local comparisons between control and policy subsamples.

Theorem 2.1. *Let Assumptions 2.6-2.8 hold and $n_0, n_1 \rightarrow \infty$ with $n_0/|\mathbf{C}|, n_1/|\mathbf{P}| \rightarrow 0$. Then:*

- (i) *$t \in \mathbf{P}$ is a complier if and only if $\overline{D}_{P,t,n_1} - \overline{D}_{C,t-1,n_0} \xrightarrow{\mathbb{P}} c$ where $c > 0$.*
- (ii) *$t \in \mathbf{C}$ is a complier if and only if $\overline{D}_{P,s^*(t),n_1} - \overline{D}_{C,t,n_0} \xrightarrow{\mathbb{P}} \tilde{c}$ where $\tilde{c} > 0$.*

Theorem 2.1 shows that the compliers can be identified individually through local comparisons of realized treatment assignments between the policy sample $\mathbf{P}$ and the control sample $\mathbf{C}$. For a policy-sample date $t \in \mathbf{P}$, a complier corresponds to the treatment assignment at t exceeding, in probability, the counterfactual treatment level inferred from nearby control-sample observations. This is captured by a strictly positive limiting gap between a local average taken within $\mathbf{P}$ and the corresponding local average within $\mathbf{C}$. For non-intervention dates $t \in \mathbf{C}$, a complier is characterized by the existence of a nearby policy-sample date whose local treatment assignment exceeds the local control-sample average at t . Economically, the theorem formalizes that a complier is a date where the policy shock genuinely shifts the path of the treatment variable relative to what would have occurred absent the shock. Because the assumptions guarantee smooth evolution of the counterfactual treatment path, local averages in the control sample provide valid proxies for the counterfactual $D_t(z)$. Thus, identifying compliers reduces to detecting persistent, sign-consistent deviations of the observed treatment path from its locally inferred counterfactual trajectory.

To the best of our knowledge, there is no equivalent result in the cross-sectional setting. The assumptions of the theorem are easily satisfied in time series applications. Using Theorem 2.1 is straightforward: one computes the difference between two sample averages and checks whether it is greater than zero. Given the sampling uncertainty associated with the two averages, one can conduct inference using a t -statistic for the null hypothesis $\mathbb{E}(D_{t_0}(1)) -$

$\mathbb{E}(D_{t_0}(0)) = 0$ (t_0 is not a complier) versus the alternative hypothesis that $\mathbb{E}(D_{t_0}(1)) - \mathbb{E}(D_{t_0}(0)) > 0$ (t_0 is a complier).

In the supplement we extend Theorem 2.1 to the case of continuous IV Z_t . This identification argument is quite general and applies well beyond the setting of FOMC announcements. In particular, it covers any narrative identification approach in which Z_t is interpreted as a shock constructed from narrative records (e.g., military spending shocks, tax shocks, oil-supply shocks, natural disasters, etc.). We explicitly demonstrate in the supplement how Theorem 2.1 extends to Ramey’s (2011) identification fiscal multipliers based on military spending news and to Romer and Romer’s (2004) narrative measure of monetary policy shocks. Finally, our framework also applies to cross-sectional settings based on spatial data by replacing continuity over time with continuity over geographical space.

An additional challenge specific to the π -LATE framework is that the set of observations with a first stage $\mathbf{S}_{0,T}$ is also unknown. However, as the following result states, under Assumption 2.4, in the absence of covariates $\tilde{V}_t$, $\mathbf{S}_{0,T}$ is equal to the (identified) set of compliers —i.e., observations for which the first-stage holds individually.

Proposition 2.2. *Suppose Z_t is binary and let Assumptions 2.3 without conditioning on $\tilde{V}_t$, 2.4 and 2.6 hold. Then, the set of compliers coincide with $\mathbf{S}_{0,T}$.*

Knowledge of the compliers sub-population (and hence of the non-compliers sub-population) can be used to test the exclusion restriction (cf. Assumption 2.2) by comparing the mean outcomes of groups of non-compliers across different values of the instrument. Intuitively, for a non-complier the treatment status is not affected by changes in the instrument. One can divide any large subset of non-compliers into two groups according to their assignment status. If one can reject the hypothesis that the average outcomes in these two groups is the same, then the exclusion restriction cannot hold.

Under Assumption 2.4 with $D_t(1) \geq D_t(0)$, the set of non-compliers is $\mathcal{NC} = \{t \in \{1, \dots, T\} : D_t(1) = D_t(0) = D_t\}$. Let $\mathcal{NC}^s$ be any non-empty subset of $\mathcal{NC}$ such that $\mathcal{NC}_{\mathbf{P}}^s = \mathcal{NC}^s \cap \mathbf{P} \neq \emptyset$ and $\mathcal{NC}_{\mathbf{C}}^s = \mathcal{NC}^s \cap \mathbf{C} \neq \emptyset$. We can test the exclusion restriction in Assumption 2.2 under the following assumption on the subsets $\mathcal{NC}_{\mathbf{P}}^s$ and $\mathcal{NC}_{\mathbf{C}}^s$.

Assumption 2.9. (i) $\mathbb{E}[Y_t^*(D_t, z)|t \in \mathcal{NC}_{\mathbf{P}}^s] = \mathbb{E}[Y_r^*(D_r, z)|r \in \mathcal{NC}_{\mathbf{C}}^s]$ for all $t, r \geq 1$, and $z \in \mathbf{Z}$. (ii) For $\mathbf{R} = \mathbf{C}$ or $\mathbf{P}$, $|\mathcal{NC}_{\mathbf{R}}^s|^{-1} \sum_{t \in \mathcal{NC}_{\mathbf{R}}^s} Y_t \xrightarrow{\mathbb{P}} \mathbb{E}[Y_t|t \in \mathcal{NC}_{\mathbf{R}}^s]$ as $|\mathcal{NC}_{\mathbf{R}}^s| \rightarrow \infty$.

Condition (i) states that the potential outcome for non-compliers is mean-stationary and the mean is the same across control and policy subsamples. Condition (ii) states that

a law of large numbers holds for non-compliers observations in both the control and policy subsamples. As long as the policy sample does not tend to contain systematic different values of the policy variable D_t among non-compliers than the control sample, these are relatively mild conditions.¹⁵

Proposition 2.3. *Suppose Z_t is binary and let Assumptions 2.4 and 2.9 hold. If Assumption 2.2 holds, then as $|\mathcal{NC}_{\mathbf{P}}^s|, |\mathcal{NC}_{\mathbf{C}}^s| \rightarrow \infty$,*

$$|\mathcal{NC}_{\mathbf{P}}^s|^{-1} \sum_{t \in \mathcal{NC}_{\mathbf{P}}^s} Y_t - |\mathcal{NC}_{\mathbf{C}}^s|^{-1} \sum_{t \in \mathcal{NC}_{\mathbf{C}}^s} Y_t \xrightarrow{\mathbb{P}} 0.$$

Using Proposition 2.3 to test Assumption 2.2 is simple: since non-compliers can be identified individually using Theorem 2.1, one can immediately compute the sample averages specified in Proposition 2.3 and conduct inference using a t -statistic for the null hypothesis that the population mean of $t \in \mathcal{NC}_{\mathbf{P}}^s$ is equal to that of $t \in \mathcal{NC}_{\mathbf{C}}^s$. The researcher has the ability to choose the subset of non-compliers $\mathcal{NC}^s$ when implementing this test. The simplest choice is to set $\mathcal{NC}^s = \mathcal{NC}$, however, the researcher also has the ability to direct the power of the test toward particular types of non-compliers they may suspect of being more likely to violate the exclusion restriction.

3 High-Frequency Identification of Monetary Policy Effects

To study the effects of monetary policy on real variables, a large literature has relied on high-frequency identification. This exploits the fact that at the time of an FOMC meeting a large amount of economic news is revealed. Here we discuss Rigobon's (2003) heteroskedasticity identification approach which uses a 1-day window [see, e.g., Nakamura and Steinsson (2018)], and can be reformulated as IV-based identification. In Section 3.1 we explain when the resulting reduced-form estimands have a causal meaning within the potential outcome framework of Section 2. In Section 3.2 we discuss the weak identification problem of current approaches and show how the π -LATE framework can be used to strengthen identification.

¹⁵The assumption is stated for the case $h = 0$. However, it can be immediately generalized to any $h > 0$.

3.1 Heteroskedasticity-Based Identification

Consider the following system of equations:

$$\tilde{Y}_t = \beta_0 \tilde{D}_t + \eta_t, \quad \text{and} \quad \tilde{D}_t = a \tilde{Y}_t + e_t, \quad (3.1)$$

where $\tilde{Y}_t$ is the (demeaned) daily change in an outcome variable, (e.g., an asset price or a bond yield) and $\tilde{D}_t$ is the (demeaned) daily change in the unexpected component of a short-term interest rate or policy news (e.g., Δi_t as discussed after Proposition 2.1), η_t is a shock to $\tilde{Y}_t$, e_t is the monetary policy shock and a and β_0 are scalar parameters. The errors η_t and e_t have no serial correlation and are mutually uncorrelated. The parameter of interest is β_0 which represents the causal effect of monetary policy on the outcome variable. The model in (3.1) could arise from a bivariate VAR. In fact, one could add a vector X_t of exogenous variables to the model in (3.1). However, to focus on the main intuition, we follow Nakamura and Steinsson (2018) and we omit X_t and lagged terms of $\tilde{Y}_t$ and $\tilde{D}_t$. See Casini and McCloskey (2025) for a detailed discussion of why the lags can be omitted in this setting.

The model in (3.1) is a special case of the generalized framework studied in Section 2. It is useful because it directly motivates a particular IV estimand. However, we study the causal interpretation of this estimand in the general case for which the linear model with stable parameters is not the correct specification.

Heteroskedasticity-based identification requires that the variance of the monetary shock increases in the days of FOMC announcements, while the variance of other shocks is unchanged. Let T_P denote the number of days containing an FOMC announcement (policy sample), and let T_C denote the number of days that do not contain an FOMC announcement (control sample). Let $\sigma_{e,P}^2 = T_P^{-1} \sum_{t \in P} \mathbb{E}(e_t^2)$ and $\sigma_{e,C}^2 = T_C^{-1} \sum_{t \in C} \mathbb{E}(e_t^2)$ be the average variance of the monetary policy shock in the policy and control samples. Define $\sigma_{\eta,P}^2$ and $\sigma_{\eta,C}^2$ similarly. Then, the identification conditions are

$$\sigma_{e,P} > \sigma_{e,C} \quad \text{and} \quad \sigma_{\eta,P} = \sigma_{\eta,C}. \quad (3.2)$$

The condition $\sigma_{e,P} > \sigma_{e,C}$ is the relevance condition while $\sigma_{\eta,P} = \sigma_{\eta,C}$ is the exclusion restriction. In the supplement we show that the resulting Wald estimand is

$$\beta_{\pi,t,0}^* = \frac{\mathbb{E}(\tilde{D}_t \tilde{Y}_t | Z_t = 1) - \mathbb{E}(\tilde{D}_t \tilde{Y}_t | Z_t = 0)}{\mathbb{E}(\tilde{D}_t^2 | Z_t = 1) - \mathbb{E}(\tilde{D}_t^2 | Z_t = 0)}, \quad (3.3)$$

which corresponds to the Wald estimand (2.4) for $h = 0$, $Y_t = \tilde{D}_t \tilde{Y}_t$, $D_t = \tilde{D}_t^2$ and no conditioning variable $\tilde{V}_t$. The following corollary of Proposition 2.1 presents the causal meaning of $\beta_{\pi,t,0}^*$ under the general setting of Section 2.

Corollary 3.1. (*LATE in heteroskedasticity-based identification*) *Let Assumptions 2.1-2.5 hold for $Y_t = \tilde{D}_t \tilde{Y}_t$ and $D_t = \tilde{D}_t^2$ with $\tilde{D}_t(1)^2 \geq \tilde{D}_t(0)^2$. For $t \in \mathbf{S}_{0,T}$, we have*

$$\beta_{\pi,t,0}^* = \frac{\int_{\mathbf{D}} \mathbb{E} \left(\frac{\partial(\tilde{Y}_{t,0}^*(\tilde{d}))}{\partial(\tilde{d}^2)} \middle| \tilde{D}_t(1)^2 \geq \tilde{d}^2 \geq \tilde{D}_t(0)^2 \right) \mathbb{P} \left(\tilde{D}_t(1)^2 \geq \tilde{d}^2 \geq \tilde{D}_t(0)^2 \right) d(\tilde{d}^2)}{\int_{\mathbf{D}} \mathbb{P} \left(\tilde{D}_t(1)^2 \geq \tilde{d}^2 \geq \tilde{D}_t(0)^2 \right) d(\tilde{d}^2)}. \quad (3.4)$$

Corollary 3.1 shows that the Wald estimand in (3.3) has a causal meaning because it is the ratio of a reduced-form generalized impulse response of $\tilde{D}_t \tilde{Y}_t$ to a first-stage generalized impulse response of $\tilde{D}_t^2$. More specifically, $\beta_{\pi,t,0}^*$ identifies a weighted average of the derivative of the product between the potential outcome and policy variable for compliers. Hence, contrary to popular belief, the causal interpretation of the heteroskedasticity-based estimator (i.e., Rigobon’s estimator) is not the same as that of a standard IV estimator—though it remains local in nature as it averages over compliers.¹⁶ We continue to refer to it as LATE with the understanding that it is a LATE for $\tilde{D}_t \tilde{Y}_t$, not $\tilde{Y}_t$ itself.

Here the compliers are the observations for which the announcement induces a higher volatility of the policy $\tilde{D}_t$. In contrast, the non-compliers are characterized by idiosyncratic or general equilibrium factors that dominate the news specific to the announcement. That is, regimes where $\tilde{D}_t^2$ remains low regardless of the presence of an announcement correspond to “never-takers,” while regimes where $\tilde{D}_t^2$ remains high even in the absence of an announcement correspond to “always-takers.” Noting that $D_t = \tilde{D}_t^2$ in this context, we can apply Theorem 2.1 to identify the compliers individually. We do so in the empirical application in Section 7.

It is important to consider how the interpretation of the causal effect identified by $\beta_{\pi,t,0}^*$ in Corollary 3.1 varies with the functional relationship between $\tilde{Y}_t$ and $\tilde{D}_t$. Let us begin with the linear case with stable parameters as in (3.1). From (3.3), simple algebra shows that $\beta_{\pi,t,0}^*$ reduces to β_0 when the denominator of (3.4) is nonzero, which means that Rigobon’s estimator identifies the causal effect of the policy (i.e., the slope coefficient in (3.1)). This result does not generally extend to the case where β_0 is time-varying or the first-stage is zero. At most one could identify a π -LATE. We will return to this in Section 3.2.

¹⁶We also note that the interpretation becomes even more complex over the region when $\tilde{D}_t$ changes sign. The interpretation of the marginal causal effect becomes an average over the two signs of $\tilde{D}_t$.

Let us turn to analyzing the consequences of nonlinearities. When D_t and the shock η_t are additively separable (i.e., $Y_t = \varphi_D(D_t) + \varphi_\eta(\eta_t)$ for some nonlinear functions $\varphi_D(\cdot)$ and $\varphi_\eta(\cdot)$), [Kolesár and Plagborg-Møller \(2025\)](#) show that the estimand resulting from a regression of Y_t on D_t using $Z_t = (W_t - \mathbb{E}(W_t))D_t$ as an instrument for which $\text{Cov}(D_t^2, W_t) \neq 0$ identifies a weighted average of marginal effects of the policy shock e_t with weights that are not guaranteed to be positive. As a result, the researcher may infer an incorrect sign for the marginal effects. Thus, this estimand is not weakly causal [cf. [Blandhol, Bonney, Mogstad, and Torgovitsky \(2025\)](#)].¹⁷ The authors also note that for the case $Y_t = e_t\varphi_\eta(\eta_t)$ with $\mathbb{E}[\varphi_\eta(\eta_t)] = 0$ and $e_t \perp \eta_t$ the estimand is nonzero while the true causal effect of the policy shock is zero since $\mathbb{E}[Y_t | e_t] = 0$.

Corollary 3.1 provides even more negative news about the effect of nonlinearities for heteroskedasticity-based identification than that shown by [Kolesár and Plagborg-Møller \(2025\)](#): in a general nonparametric model, Corollary 3.1 implies that Rigobon’s Wald estimand $\beta_{\pi,t,0}^*$, which is in general different from the IV estimand examined by [Kolesár and Plagborg-Møller \(2025\)](#), does not necessarily equal a weighted average of marginal effects. The intuition is that the instrument affects $\text{Var}(D_t)$ and not $\mathbb{E}(D_t)$, so variation in the instrument induces exogenous variation in D_t^2 , which has a causal effect on $D_t Y_t$ not just Y_t . In short, it is generally difficult to interpret $\beta_{\pi,t,0}^*$ when the true model is nonlinear. Thus, we concur with the recommendation of [Kolesár and Plagborg-Møller \(2025\)](#) that the linearity assumption should be checked carefully when using heteroskedasticity-based identification. This is likely even more important in the context of SVARs and local projections than in the current event study setting since the former aggregates data over a month or a quarter while the latter uses relatively higher frequency data (e.g., a 30-minute or 1-day change in policy and outcome variables around an announcement), where linearity may be a more credible assumption since a nonlinear function can be locally well approximated by a linear one.

3.2 Weak or Lack of Identification and the Usefulness of π -LATE

The key identification condition that the volatility of the policy variable is higher during FOMC announcement days appears reasonable in principle, since each announcement day is likely to be associated with substantial monetary news. However, the volatility of monetary policy variables can be high for other reasons. There are multi-year periods during which

¹⁷[Kitagawa, Wang, and Xu \(2025\)](#) show how a reasonable economic interpretation can potentially be restored when the weights are negative.

the volatility of several macroeconomic variables is elevated. In this case, general equilibrium factors dominate the news specific to the announcement. For example, during the 2007-09 financial crisis and the Covid-19 pandemic, volatility was high across many macroeconomic and financial variables. These facts pose serious challenges for identification, as the first-stage condition may not hold for all t . To see this, examine the denominator in (3.3). If the first-stage does not hold for all t , we may have

$$T_P^{-1} \sum_{t \in \mathbf{P}} \text{Var}(\tilde{D}_t) - T_C^{-1} \sum_{t \in \mathbf{C}} \text{Var}(\tilde{D}_t) \approx 0, \quad (3.5)$$

which would render the estimate of the average treatment effect highly imprecise.

Using an F -test for weak identification, Lewis (2022) shows that the monetary policy effects based on a 1-day window in Nakamura and Steinsson (2018) appear to be weakly-identified. We show that this arises from significant time variation in the volatility of the policy variable within both policy and control samples. Figure 2 plots $\tilde{D}_t$ (2-Year Treasury yields) for the control and policy samples. The policy sample includes all regularly scheduled FOMC meeting days from 1/1/2000 to 3/19/2014. The control sample includes all Tuesdays and Wednesdays that are not FOMC meeting days between 1/1/2000 and 12/31/2012.

There appear to be multiple volatility regimes. Using the structural break test from Casini and Perron (2024), which allows for stable or smoothly varying volatility under the null and abrupt breaks under the alternative, we detect three breaks in the control sample. The first break (April 24, 2007) marks the start of the 2007-09 financial crisis. The second (July 28, 2009) captures the crisis period itself, characterized by the highest volatility. Afterward, volatility returns to pre-crisis levels until the third break (February 2, 2011), which aligns with the zero lower bound (ZLB) period and the start of unconventional monetary policy. The final regime shows the lowest volatility, reflecting initial policy effects and stabilization.¹⁸

These findings show significant time variation in $\text{Var}(\tilde{D}_t)$. In the second regime, control-sample volatility is close to the policy-sample average, contributing to the weak identification in (3.5). Nakamura and Steinsson (2018) find their estimates imprecise and not economically meaningful for some of the interest rates they use as outcome variables. Lewis (2022) reports a first-stage F -statistic of 8.11—well below the 23 critical value—suggesting weak identification.

¹⁸We do not test for breaks in the policy sample due to small size ($T_P = 74$), treating it as a single regime.

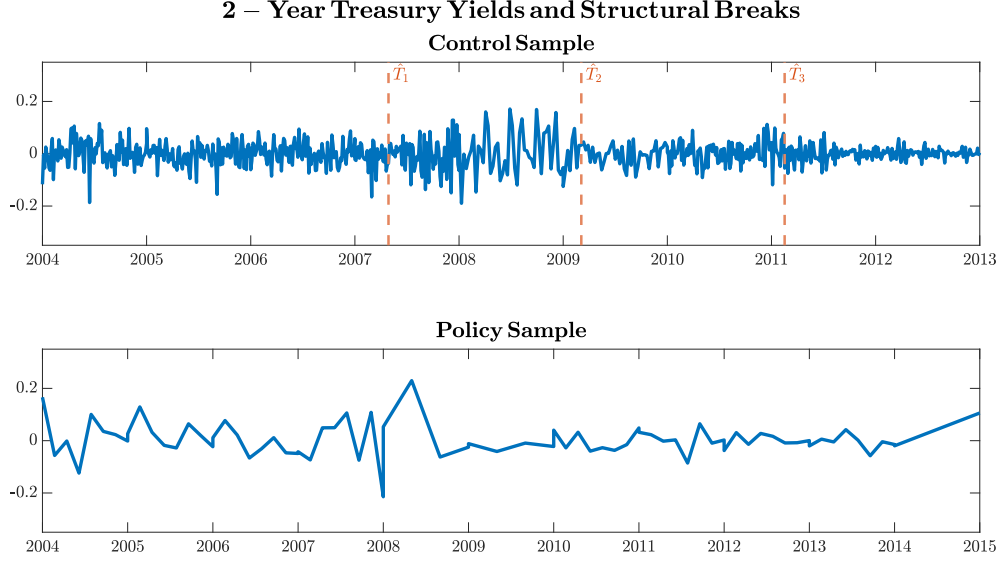


Figure 2: Plot of 2-years Treasury yields in control (top panel) and policy sample (bottom panel). Vertical broken lines are the estimated break dates using [Casini and Perron’s \(2024\)](#) test.

We propose to focus on π -LATE. The fraction π_0 of the sample (i.e., all $t \in \mathbf{S}_{0,T}$) that has a first-stage corresponds to the regimes in the control sample where $\text{Var}(\tilde{D}_t)$ is low (relative to its average level). For example, it is likely that the regime $[\hat{T}_1 + 1, \hat{T}_2]$ does not belong to $\mathbf{S}_{0,T}$ since $\text{Var}(\tilde{D}_t)$ within this regime appears close to the average volatility in the policy sample. By construction, it is easier to identify π -LATE than full sample LATE. The usefulness of π -LATE depends on the magnitude of π_0 : a small π_0 implies that identification is achievable only in a small portion of the population, whereas a large π_0 indicates that the identified π -LATE is representative of a substantial part of the population.¹⁹

The π -LATE parameter is the same as the LATE parameter in Section 3.1 but instead of supposing that a first-stage exists, only uses observations with a nonzero first-stage. Let $T_{P,S}$ denote the number of days in $\mathbf{S}_{0,T}$ that contain an FOMC announcement, and let $T_{C,S}$ the number of days in $\mathbf{S}_{0,T}$ that do not contain an FOMC announcement. This means

¹⁹It is possible that in practice the π_0 fraction of the sample contains a mixture of strong and weak identification. We discuss weak identification in the context of π -LATE formally in Section 6.

$T_{P,S} + T_{C,S} = \pi_0 T$.²⁰ Let $\mathbf{P}_S = \mathbf{P} \cap \mathbf{S}_{0,T}$ and $\mathbf{C}_S = \mathbf{C} \cap \mathbf{S}_{0,T}$. The Wald estimand is

$$\tilde{\beta}_{\pi,t,0}^* = \frac{\mathbb{E}(\tilde{D}_t \tilde{Y}_t | t \in \mathbf{P}_S) - \mathbb{E}(\tilde{D}_t \tilde{Y}_t | t \in \mathbf{C}_S)}{\mathbb{E}(\tilde{D}_t^2 | t \in \mathbf{P}_S) - \mathbb{E}(\tilde{D}_t^2 | t \in \mathbf{C}_S)}, \quad (3.6)$$

to which Corollary 3.1 immediately applies without the (now redundant) qualifier “for $t \in \mathbf{S}_{0,T}$.” π -LATE is the average treatment effect for the sub-population for which a first-stage holds: observations for which $\tilde{D}_t^2$ is induced to be higher by the announcement (i.e., the sub-population of compliers in $\mathbf{S}_{0,T}$).

If the treatment effect is constant across the population [e.g., as in (3.1)], then the π -LATE for the sub-population $\mathbf{S}_{0,T}$ is equal to both the LATE and ATE in the full population. To determine which treatment effect is identified, we must determine which parts of the sample belong to $\mathbf{S}_{0,T}$. We discuss this in Sections 4-5.

4 Testing for Full Population Identification Failure

In this section, we introduce a test of the null hypothesis that no subpopulation exists for which a LATE can be identified, even weakly. In other words, the test assesses whether identifying a sub-population LATE is possible at all. However, we strongly caution against using this as a pretest before estimation or inference, as doing so may introduce pretest bias and invalidates standard inference unless the inference method is modified to account for the pretest [see, e.g., Andrews (2018)]. Instead, the test should be viewed as a diagnostic tool for evaluating whether there is evidence of identifiable sub-population LATEs in a given application. We apply it for this purpose to several existing studies that appear to face identification challenges. Notably, such a pretest is unnecessary for conducting identification-robust inference on sub-population LATEs, which we discuss in Section 6.

In accord with the analysis of Section 2, consider an IV regression model with a single endogenous variable and multiple instruments. In matrix format, the structural equation is

$$Y = D\beta + X\gamma_1 + u, \quad t = 1, \dots, T, \quad (4.1)$$

where Y is a $T \times 1$ vector of outcome variables, D is $T \times 1$ vector of endogenous variables,

²⁰For notational simplicity we assume that $\pi_0 T$ is an integer so that we avoid using the notation $\lfloor \pi_0 T \rfloor$, where $\lfloor \cdot \rfloor$ denotes the largest smaller integer function.

X is a $T \times p$ matrix of p exogenous regressors, u is a $T \times 1$ vector of error terms, and $\beta \in \mathbb{R}$ and $\gamma_1 \in \mathbb{R}^p$ are unknown parameters. The reduced-form equation is

$$D_t = Z_t' \theta \mathbf{1}\{t \in \mathbf{S}_{0,T}\} + X_t' \gamma_2 + e_t, \quad (4.2)$$

where Z_t is a $q \times 1$ vector of instruments, e_t is an error term, and $\theta \in \mathbb{R}^q$ and $\gamma_2 \in \mathbb{R}^p$ are unknown parameters. For $t \notin \mathbf{S}_{0,T}$, the instrument Z_t is irrelevant. For $t \in \mathbf{S}_{0,T}$, the instrument Z_t is relevant if $\theta \neq 0$. We assume that $|\mathbf{S}_{0,T}| = \pi_0 T$ for some $\pi_0 \in (0, 1]$, noting that this is without loss of generality since it does not rule out complete identification failure which occurs when $\theta = 0$ for any $\pi_0 \in (0, 1]$. The hypothesis testing problem is

$$H_{\theta,0} : \theta = 0 \quad \text{versus} \quad H_{\theta,1} : \theta \neq 0.$$

We discuss both the cases for which the sub-population $\mathbf{S}_{0,T}$ is known and unknown. For the sake of the exposition, we focus on homogeneous θ in $\mathbf{S}_{0,T}$.²¹

Consider the $(\pi T \times T)$ selection matrix S_T that selects the πT rows of a matrix corresponding to the indices in $\mathbf{S}_T$. That is, for an arbitrary $T \times k$ matrix A , $S_T A$ is the $(\pi T \times k)$ matrix whose elements are the rows of A that correspond to the indices in $\mathbf{S}_T$. For example, if $\mathbf{S}_T = \{1, \dots, 0.25T, 0.75T + 1, \dots, T\}$,

$$S_T A = \left[A^{(1,:)}' : \dots : A^{(0.25T,:)}' : A^{(0.75T+1,:)}' : \dots : A^{(T,:)}' \right]',$$

where $A^{(r,:)}$ denotes the r^{th} row of the matrix A . Using the standard projection matrix notation, $P_A = A(A'A)^{-1}A'$ and $M_A = I - P_A$, let $\tilde{A}(S_T) = M_{S_T X} S_T A$ for any arbitrary $T \times k$ matrix A . The following F test statistic is useful for testing whether $\theta = 0$ in the regression (4.2) when the sub-population $\mathbf{S}_{0,T}$ is known:

$$F_T(\mathbf{S}_T) = \frac{\tilde{D}(S_T)' \tilde{Z}(S_T) \hat{J}(S_T)^{-1} \tilde{Z}(S_T)' \tilde{D}(S_T)}{q(\pi T - p - q)},$$

for $\mathbf{S}_T = \mathbf{S}_{0,T}$ and $Z = [Z_1 : \dots : Z_T]'$ and $\hat{J}(S_T)$ a consistent estimate of the long-run variance $\lim_{T \rightarrow \infty} (T\pi)^{-1} \text{Var}(\tilde{Z}(S_T)' S_T e)$ with $e = [e_1 : \dots : e_T]'$. HAC or DK-HAC estimators can be used to estimate the long-run variance [cf. Andrews (1991), Casini (2023) and Newey and

²¹We could allow for $\theta_t \neq 0$ for $t \in \mathbf{S}_{0,T}$ at the expense of additional notation and longer proofs, though the key insights would not change. Actually, the computational procedures we develop to implement our methods allow $\theta_t \neq 0$ for $t \in \mathbf{S}_{0,T}$.

West (1987)].

For the case of an unknown sub-population, we follow the structural break literature and search for maximal identification strength over all sub-populations of minimal size $\pi_L T$ that can be partitioned into m distinct smaller sub-populations, where $\pi_L > 0$ and $1 \leq m \leq m_+$ for some upper bound on the number of regimes $m_+ > 0$:

$$F_T^* = \sup_{\pi \in [\pi_L, 1]} \max_{1 \leq m \leq m_+} \sup_{\mathbf{S}_T \in \Xi_{\epsilon, \pi, m, T}} F_T(\mathbf{S}_T),$$

where $\Xi_{\epsilon, \pi, m, T}$ denotes the set of all possible partitions of a fraction π of $\{1, \dots, T\}$ that involve m regimes $((\lambda_{L,1}T, \lambda_{R,1}T), \dots, (\lambda_{L,m}T, \lambda_{R,m}T))$ for $\lambda_{L,i}, \lambda_{R,i} \in [0, 1]$ such that (i) $\lambda_{L,i} < \lambda_{R,i}$ for all i , (ii) $\lambda_{R,i} < \lambda_{L,i+1}$ for $i = 1, \dots, m-1$, (iii) $|\lambda_{R,i} - \lambda_{L,i}| \geq \epsilon$ for all i and some (small) $\epsilon > 0$ and (iv) $\sum_{i=1}^m (\lambda_{R,i} - \lambda_{L,i}) = \pi$. Conditions (i) and (ii) correspond to $T\lambda_{L,i}$ ($T\lambda_{R,i}$) denoting the start (end) date of regime i within the sub-population $\mathbf{S}_T$ while condition (iii) implies that each regime involves a non-negligible fraction of the sample. The statistic F_T^* thus implicitly searches for maximal identification strength over all possible sub-populations of size $\pi_L T$ and larger with less than m_+ distinct regimes that are at least an ϵ fraction of the overall sample size.

The tuning parameters π_L and ϵ determine the types of sub-populations for which the test can detect identification: smaller values of π_L allow detection in smaller sub-populations, while smaller values of ϵ enable detection in sub-populations with shorter regimes. The choice of these lower bounds should be guided by the empirical context, reflecting the smallest sub-population and regime sizes for which LATE inference remains meaningful in the application.²² In our simulations and empirical applications we set $\pi_L = 0.6$ and $\epsilon = 0.05$.

For X'_t the t^{th} row of X , let $w_t = (X'_t, Z'_t)'$ and $W_r(\cdot)$ denote a r -vector of independent Wiener processes on $[0, 1]$. We derive the asymptotic null distributions of $F_T(\mathbf{S}_T)$ and F_T^* under the following standard high-level assumptions that permit both heteroskedastic and serially correlated errors. Sufficient conditions for them can be found in the supplement.

Assumption 4.1. $T^{-1} \sum_{t=1}^{\lfloor Ts \rfloor} w_t w'_t \xrightarrow{\mathbb{P}} sQ$, uniformly in $s \in [0, 1]$ for some p.d. matrix Q .

Assumption 4.2. $T^{-1/2} \sum_{t=1}^{\lfloor Ts \rfloor} w_t e_t \Rightarrow \Omega_{we}^{1/2} W_{p+q}(s)$ for some p.d. variance matrix Ω_{we} .

Assumption 4.3. $\hat{J}(S_T)$ is p.d. for all T , $\mathbf{S}_T \in \Xi_{\epsilon, \pi, m, T}$ and $\hat{J}(S_T) \xrightarrow{\mathbb{P}} \lim_{T \rightarrow \infty} (T\pi)^{-1} \text{Var}(e' S'_T \tilde{Z}(S_T))$ uniformly in $\mathbf{S}_T \in \Xi_{\epsilon, \pi, m, T}$.

²²In the structural break literature, common recommendations for ϵ are 0.05, 0.10 and 0.15. See Casini and Perron (2019) for a review.

Theorem 4.1. *Let Assumptions 4.1-4.3 hold. Under $H_{\theta,0}$,*

$$F_T(\mathbf{S}_T) \Rightarrow F(\mathbf{S}) \quad \text{if} \quad \mathbf{S}_T \in \Xi_{\epsilon,\pi,m,T}, \quad \text{and} \quad F_T^* \Rightarrow \sup_{\pi \in [\pi_L, 1]} \max_{1 \leq m \leq m_+} \sup_{\mathbf{S} \in \Xi_{\epsilon,\pi,m}} F(\mathbf{S}),$$

where $\mathbf{S} = \lim_{T \rightarrow \infty} T^{-1} \mathbf{S}_T$, $\Xi_{\epsilon,\pi,m} = \lim_{T \rightarrow \infty} T^{-1} \Xi_{\epsilon,\pi,m,T}$ and

$$F(\mathbf{S}) = \frac{1}{q\pi} \left\| \sum_{i=1}^m (W_q(\lambda_{R,i}) - W_q(\lambda_{L,i})) \right\|^2.$$

When $\pi = 1$ ($\pi_L = 1$ and $m_+ = 1$), $F_T(\mathbf{S}_T)$ (F_T^*) reduces to the usual first-stage F -statistic for $\theta = 0$ in (4.2). For $\pi \in (0, 1)$ ($\pi_L \in (0, 1)$), the consistency of tests against $H_{\theta,1}$ using $F_T(\mathbf{S}_{0,T})$ (F_T^*) follows from similar arguments as for the $\pi_0 = 1$ case. The asymptotic null distributions of both $F(\mathbf{S})$ and F_T^* are free of nuisance parameters. The critical values are obtained via simulations and reported in Table 4 for up to $m_+ = 6$ and up to $q = 6$.

5 Estimation of LATE and Identified Sub-Populations

We discuss estimation of the LATE parameter β in (4.1) in both the cases of a known and unknown sub-population $\mathbf{S}_{0,T}$, as well as estimation of $\mathbf{S}_{0,T}$ itself in the latter case. When $\mathbf{S}_{0,T}$ is known, estimation of β is an application of IV estimation for which $Z_t \mathbf{1}\{t \in \mathbf{S}_{0,T}\}$ is treated as the vector of instruments. Let this estimator be denoted as $\hat{\beta}(\mathbf{S}_{0,T})$.

On the other hand, when the sub-population $\mathbf{S}_{0,T}$ is unknown, we must estimate it first. Although $\mathbf{S}_{0,T}$ can be estimated consistently in the special case of a binary instrument under the conditions of Proposition 2.2 and Theorem 2.1, it can also be estimated more generally. We discuss two methods. The first is more computationally straightforward but the second is more efficient because it uses the information in both structural and reduced-form equations (4.1)-(4.2). We follow the structural change literature and assume that π_0 and m_0 are known, i.e., the practitioner has previously used the tests from Section 4 to determine π_0 and m_0 .

We begin with the first estimator. Consider the $T \times T$ matrix C_T that selects the πT rows of a matrix corresponding to the indices in $\mathbf{S}_T$ while setting the remaining $(1 - \pi)T$ rows to zero. For example, for a $T \times k$ matrix A , if $\mathbf{S}_T = \{1, \dots, 0.25T, 0.75T + 1, \dots, T\}$,

$$C_T A = \left[A^{(1,:)}' : \dots : A^{(0.25T,:)}' : 0_{k \times 1} : \dots : 0_{k \times 1} : A^{(0.75T+1,:)}' : \dots : A^{(T,:)}' \right]'$$

Let $\bar{A}(C_T) = M_X C_T A$ so that for a given $\mathbf{S}_T$, the OLS estimators of θ and γ_2 in (4.2) can be ex-

pressed as $\hat{\theta}_{OLS}(\mathbf{S}_T) = (\bar{Z}(C_T)' \bar{Z}(C_T))^{-1} \bar{Z}(C_T)' D$ and $\hat{\gamma}_{2,OLS}(\mathbf{S}_T) = (X' M_{C_T Z} X)^{-1} X' M_{C_T Z} D$. Our first estimator of $\mathbf{S}_{0,T}$ minimizes the sum of squared residuals of the reduced-form:

$$\hat{\mathbf{S}}_{T,OLS} = \underset{\mathbf{S}_T \in \Xi_{\epsilon, \pi_0, m_0, T}}{\operatorname{argmin}} \left(D - C_T Z \hat{\theta}_{OLS}(\mathbf{S}_T) - X \hat{\gamma}_{2,OLS}(\mathbf{S}_T) \right)' \left(D - C_T Z \hat{\theta}_{OLS}(\mathbf{S}_T) - X \hat{\gamma}_{2,OLS}(\mathbf{S}_T) \right).$$

Correspondingly, we estimate β with $\hat{\beta}(\hat{\mathbf{S}}_{T,OLS})$. In the supplement, we present the consistency results for $\hat{\mathbf{S}}_{T,OLS}$ and $\hat{\beta}(\hat{\mathbf{S}}_{T,OLS})$, and we consider a second estimator of $\mathbf{S}_{0,T}$, namely a GLS criterion that minimizes an efficiently weighted combination of the sum of squared residuals of both the structural (4.1) and reduced-form equation (4.2).

In model (4.1) the LATE parameter β is constant, so π -LATE is the full population LATE and $\hat{\beta}(\hat{\mathbf{S}}_{T,OLS})$ is consistent for the LATE parameter β . This can be a precise estimate even when a first-stage F test detects full sample weak identification because it uses the most-strongly identified subsample. When the model (4.1) is misspecified, so that LATEs may be nonlinear and time-varying, $\hat{\beta}(\hat{\mathbf{S}}_{T,OLS})$ is still consistent for a weighted average of the LATEs in the $\mathbf{S}_{0,T}$ subsample if the $\mathbf{S}_{0,T}$ subsample exhibits strong identification.

The estimator $\hat{\mathbf{S}}_{T,OLS}$ and the test statistic F_T^* solve an optimization problem over many partitions. This is computationally more complex than problems in the structural breaks literature, as it involves optimizing both over sample partitions and identification strength. We address this challenge by proposing an efficient algorithm based on dynamic programming, extending the approach of [Bai and Perron \(2003\)](#) to our setting.²³

6 Identification-Robust Inference

We consider tests on β in (4.1) that are robust to weak identification in both the cases for which the sub-population $\mathbf{S}_{0,T}$ is known and unknown. The hypothesis testing problem is $H_0 : \beta = \beta_0$ versus $H_1 : \beta \neq \beta_0$. Here we present results for the case of unknown sub-population $\mathbf{S}_{0,T}$ and weak instruments. We defer to the supplement the formal treatment of the case of known $\mathbf{S}_{0,T}$ and strong instruments. Rewrite (4.1)-(4.2) as

$$y = \bar{Z}(C_{0,T}) \theta a'_\beta + X \eta + v, \quad y = [Y : D], \quad v = [v_1 : e], \quad a_\beta = (\beta, 1)', \quad \eta = [\gamma : \phi], \quad (6.1)$$

²³While [Antoine and Boldea \(2018\)](#) consider the case of a single break, and [Magnusson and Mavroeidis \(2014\)](#) study a related context, neither provide a computational solution—referring to the problem as “computationally demanding.”

where $v_1 = u + \beta e$, $\gamma = \gamma_1 + \phi\beta$ and $\phi = \gamma_2 + (X'X)^{-1}X'C_{0,T}Z\theta$. When $\mathbf{S}_{0,T}$ is known, it is straightforward to use existing tests in the identification-robust linear IVs literature to test H_0 [cf. [Anderson and Rubin \(1949\)](#), [Andrews, Moreira, and Stock \(2006\)](#), [Kleibergen \(2002\)](#) and [Moreira \(2003\)](#)]. However, Proposition [S.B.1](#) in the supplement shows that $Z'M_X y$ is not a sufficient statistic for $(\beta, \theta)'$ but $\bar{Z}(C_{0,T})'y$ is, implying that existing tests suffer a loss in efficiency because they treat Z rather than $C_{0,T}Z$ as the matrix of IVs. Efficient tests are therefore functions of $\bar{Z}(C_{0,T})'y$. [Magnusson and Mavroeidis \(2014\)](#) consider a model similar to [\(6.1\)](#). Our model specifies that θ is nonzero in the sub-population $\mathbf{S}_{0,T}$ and is zero in $\mathbf{S}_{0,T}^c$ where $\mathbf{S}_{0,T}^c$ is the complement of $\mathbf{S}_{0,T}$. [Magnusson and Mavroeidis \(2014\)](#) allow the first-stage coefficient θ_t to be generally time-varying for some of their tests. Their tests are based on the full sample of observations whereas our tests are based on a lower-dimensional statistic since we do not use the sub-population $\mathbf{S}_{0,T}^c$. This allows us to obtain gains in efficiency.

When $\mathbf{S}_{0,T}$ is known we can apply the results of [Andrews, Moreira, and Stock \(2006\)](#) to form identification-robust tests of H_0 vs H_1 that are functions of $\bar{Z}(C_{0,T})'y$ and are robust to both heteroskedasticity and autocorrelation (HAR) in the reduced-form errors $\{v_t\}$. Suppose $\hat{\Sigma}_{N_1}(\mathbf{S}_{0,T})$, $\hat{\Sigma}_{N_1,N_2}(\mathbf{S}_{0,T})$ and $\hat{\Sigma}_{N_2}(\mathbf{S}_{0,T})$ are consistent estimators of $\Sigma_{N_1}(\mathbf{S}_0)$, $\Sigma_{N_1,N_2}(\mathbf{S}_0)$ and $\Sigma_{N_2}(\mathbf{S}_0)$ under H_0 , where these latter quantities are defined by

$$\Sigma_{v\bar{Z}}(\mathbf{S}_0) = \begin{bmatrix} \Sigma_{N_1}(\mathbf{S}_0) & \Sigma_{N_1N_2}(\mathbf{S}_0)' \\ \Sigma_{N_1N_2}(\mathbf{S}_0) & \Sigma_{N_2}^*(\mathbf{S}_0) \end{bmatrix}, \quad (6.2)$$

$$\Sigma_{N_2}(\mathbf{S}_0) = \Sigma_{N_2}^*(\mathbf{S}_0) - \Sigma_{N_1N_2}(\mathbf{S}_0) \Sigma_{N_1}^{-1}(\mathbf{S}_0) \Sigma_{N_1N_2}(\mathbf{S}_0)'$$

for $\Sigma_{v\bar{Z}}(\mathbf{S}_0) = \Sigma_{v\bar{Z}}(\mathbf{S}_0, \mathbf{S}_0)$, with

$$\Sigma_{v\bar{Z}}(\mathbf{S}, \mathbf{S}') = \lim_{T \rightarrow \infty} \text{Cov} \left(T^{-1/2} \sum_{t=1}^T \begin{bmatrix} v_t' b_0 \bar{Z}_t(C_T) \\ v_t' \Sigma_v^{-1} a_{0,\beta} \bar{Z}_t(C_T) \end{bmatrix}, T^{-1/2} \sum_{t=1}^T \begin{bmatrix} v_t' b_0 \bar{Z}_t(C_T') \\ v_t' \Sigma_v^{-1} a_{0,\beta} \bar{Z}_t(C_T') \end{bmatrix} \right)$$

for $\mathbf{S} = \lim_{T \rightarrow \infty} T^{-1} \mathbf{S}_T$, $\mathbf{S}' = \lim_{T \rightarrow \infty} T^{-1} \mathbf{S}_T'$, $b_0 = (1, -\beta_0)'$ and $a_{0,\beta} = (\beta_0, 1)'$, where v_t and $\bar{Z}_t(C_T)$ are the transposes of the t th rows of v and $\bar{Z}(C_T)$.²⁴ Let $\hat{\Sigma}_v(\mathbf{S}_{0,T}) = (T - q - p)^{-1} \hat{v}(\mathbf{S}_{0,T})' \hat{v}(\mathbf{S}_{0,T})$ with $\hat{v}(\mathbf{S}_{0,T}) = y - P_{\bar{Z}(C_{0,T})} y - P_X y$. Define

$$N_{1,T}(\mathbf{S}_{0,T}) = \hat{\Sigma}_{N_1}^{-1/2}(\mathbf{S}_{0,T}) T^{-1/2} \bar{Z}(C_{0,T})' y b_0 \quad \text{and} \quad (6.3)$$

$$N_{2,T}(\mathbf{S}_{0,T}) = \hat{\Sigma}_{N_2}^{-1/2}(\mathbf{S}_{0,T}) \left(T^{-1/2} \bar{Z}(C_{0,T})' y \hat{\Sigma}_v^{-1}(\mathbf{S}_{0,T}) a_{0,\beta} - \hat{\Sigma}_{N_1N_2}(\mathbf{S}_{0,T}) \hat{\Sigma}_{N_1}^{-1/2}(\mathbf{S}_{0,T}) N_{1,T}(\mathbf{S}_{0,T}) \right).$$

²⁴See the supplement for details on how to construct these estimators and for consistency results.

Consider the following HAR versions of the Anderson-Rubin (AR), Lagrange multiplier (LM) and likelihood ratio statistics based on the sufficient statistic $\bar{Z}(C_{0,T})'y$:

$$\begin{aligned} AR_T(\mathbf{S}_{0,T}) &= M_{1,T}(\mathbf{S}_{0,T}), & LM_T(\mathbf{S}_{0,T}) &= \frac{M_{1,2,T}(\mathbf{S}_{0,T})^2}{M_{2,T}(\mathbf{S}_{0,T})}, \\ LR_T(\mathbf{S}_{0,T}) &= \frac{1}{2} \left(M_{1,T}(\mathbf{S}_{0,T}) - M_{2,T}(\mathbf{S}_{0,T}) + \sqrt{(M_{1,T}(\mathbf{S}_{0,T}) - M_{2,T}(\mathbf{S}_{0,T}))^2 + 4M_{1,2,T}(\mathbf{S}_{0,T})^2} \right), \end{aligned} \quad (6.4)$$

where $M_{1,T}(\mathbf{S}_{0,T}) = N_{1,T}(\mathbf{S}_{0,T})' N_{1,T}(\mathbf{S}_{0,T})$, $M_{1,2,T}(\mathbf{S}_{0,T}) = N_{1,T}(\mathbf{S}_{0,T})' N_{2,T}(\mathbf{S}_{0,T})$ and $M_{2,T}(\mathbf{S}_{0,T}) = N_{2,T}(\mathbf{S}_{0,T})' N_{2,T}(\mathbf{S}_{0,T})$. The conditional likelihood ratio (CLR) test of level α rejects H_0 when $LR_T(\mathbf{S}_{0,T}) > \kappa_\alpha(N_{2,T}(\mathbf{S}_{0,T}))$, where the critical value function $\kappa_\alpha(\cdot)$ is defined such that $\kappa_\alpha(n_2)$ is the $1 - \alpha$ quantile of the large-sample conditional distribution of $LR_T(\mathbf{S}_{0,T})$ under H_0 , given $N_{2,T}(\mathbf{S}_{0,T}) = n_2$:

$$\frac{1}{2} \left(\mathcal{Z}'_q \mathcal{Z}_q - n'_2 n_2 + \sqrt{(\mathcal{Z}'_q \mathcal{Z}_q - n'_2 n_2)^2 + 4(\mathcal{Z}'_q n_2)^2} \right),$$

where $\mathcal{Z}_q \sim \mathcal{N}(0, I_q)$. The critical value function $\kappa_\alpha(\cdot)$ is approximated in [Moreira \(2003\)](#). The LM and AR tests reject H_0 when $LM_T > \chi_1^2(1 - \alpha)$ and $AR_T > \chi_q^2(1 - \alpha)$, where $\chi_q^2(1 - \alpha)$ denotes the $1 - \alpha$ quantile of a chi-squared distribution with q degrees of freedom.

When $\mathbf{S}_{0,T}$ is known the results of [Andrews et al. \(2006\)](#) imply that the CLR, LM and AR tests have limiting null rejection probabilities equal to α under weak IV asymptotics, $\theta = c/T^{1/2}$ for some nonstochastic $c \in \mathbb{R}^q$, under a weakening of Assumptions [6.1-6.4](#) below for which these assumptions need only hold pointwise in $\mathbf{S}_T$. These tests are asymptotically similar and therefore have asymptotically correct size in the presence of weak IVs.

For the case of an unknown sub-population, the identification-robust tests in the extant literature no longer apply because the set of instruments $C_{0,T}Z$ is unknown and must be estimated. In this section, we show how to form HAR CLR, LM and AR tests with correct asymptotic null rejection probabilities under both weak and strong IV asymptotics. To estimate the true sub-population $\mathbf{S}_{0,T}$ when constructing these tests let

$$\hat{\mathbf{S}}_T = \arg \max_{\mathbf{S}_T \in \mathcal{S}} M_{2,T}(\mathbf{S}_T), \quad \text{where} \quad \mathcal{S} = \bigcup_{1 \leq m \leq m_+} \bigcup_{\pi \in (\epsilon, 1]} \Xi_{\epsilon, \pi, m, T}. \quad (6.5)$$

Proposition [S.B.2](#) in the supplement shows that the process $\{\bar{Z}(C_T)'y\}_{\mathbf{S}_T \in \mathcal{S}}$ is sufficient for (β, θ') in a canonical Gaussian setting analogous to that in [Andrews et al. \(2006\)](#) so that there is no loss in efficiency from using the unknown sub-population AR, LM and LR statistics,

$LR_T(\hat{\mathbf{S}}_T)$, $LM_T(\hat{\mathbf{S}}_T)$ and $AR_T(\hat{\mathbf{S}}_T)$, which are only functions of the process $\{\bar{Z}(C_T)'y\}_{\mathbf{S}_T \in \mathcal{S}}$.

We establish the asymptotic validity of the HAR CLR, LM and AR tests in the unknown sub-population setting under a weak set of high-level sufficient conditions on the IVs, exogenous variables and errors. Define $w(\mathbf{S}_T) = [C_T Z : X]$.

Assumption 6.1. *Let $\mathbf{S} = T^{-1}\mathbf{S}_T$, $\mathbf{S}' = T^{-1}\mathbf{S}'_T$, and $\mathcal{S}_\infty = \{\mathbf{S} : \mathbf{S} = \lim_{T \rightarrow \infty} \mathbf{S}_T \text{ for some sequence } \mathbf{S}_T \in \mathcal{S}\}$. There exists a matrix-valued function $Q : \mathcal{S}_\infty \times \mathcal{S}_\infty \rightarrow \mathbb{R}^{(q+p) \times (q+p)}$ such that $\sup_{\mathbf{S}_T, \mathbf{S}'_T \in \mathcal{S}} \|T^{-1}w(\mathbf{S}_T)'w(\mathbf{S}'_T) - Q(\mathbf{S}, \mathbf{S}')\| \xrightarrow{\mathbb{P}} 0$. Moreover, the diagonal limits are uniformly positive definite: $\inf_{\mathbf{S} \in \mathcal{S}_\infty} \lambda_{\min}(Q(\mathbf{S}, \mathbf{S}')) > 0$.*

Assumption 6.2. *$T^{-1}v'v \xrightarrow{\mathbb{P}} \Sigma_v$ for some 2×2 p.d. matrix Σ_v .*

Assumption 6.3. *For $\mathbf{S}_T, \mathbf{S}'_T \in \mathcal{S}$ and $\mathbf{S} = \lim_{T \rightarrow \infty} T^{-1}\mathbf{S}_T$, $\mathbf{S}' = \lim_{T \rightarrow \infty} T^{-1}\mathbf{S}'_T$, $T^{-1/2}\text{vec}(w(\mathbf{S}_T)'v) \Rightarrow \mathcal{G}(\mathbf{S})$, where $\mathcal{G}(\cdot)$ is a mean-zero Gaussian process indexed by $\mathbf{S} \subseteq (0, 1]$ with $2(q+p) \times 2(q+p)$ covariance function $\Psi(\mathbf{S}, \mathbf{S}') = \lim_{T \rightarrow \infty} T^{-1}\text{Cov}(\text{vec}(w(\mathbf{S}_T)'v), \text{vec}(w(\mathbf{S}'_T)'v))$.*

In Assumption 6.3, $\text{vec}(\cdot)$ denotes the vec operator. The quantities $Q(\cdot)$, Σ_v , and $\Psi(\cdot)$ are assumed to be unknown. Assumptions 6.1-6.2 hold under suitable conditions by a (uniform) law of large numbers. Assumption 6.3 holds under suitable conditions by a functional central limit theorem. Assumptions 6.1-6.3 are consistent with non-normal, heteroskedastic, autocorrelated errors and IVs and regressors that may be random or non-random.²⁵

We assume that we can consistently estimate $\Sigma_{v\bar{Z}}(\mathbf{S}) \equiv \Sigma_{v\bar{Z}}(\mathbf{S}, \mathbf{S})$ uniformly in $\mathbf{S}_T$.

Assumption 6.4. *We have an estimator $\hat{\Sigma}_{v\bar{Z}}(\mathbf{S}_T)$ such that $\hat{\Sigma}_{v\bar{Z}}(\mathbf{S}_T) \xrightarrow{\mathbb{P}} \Sigma_{v\bar{Z}}(\mathbf{S})$ uniformly in $\mathbf{S}_T \in \mathcal{S}$ for $\mathbf{S} = \lim_{T \rightarrow \infty} T^{-1}\mathbf{S}_T$.*

Note that this assumption immediately implies the uniform consistency of $\hat{\Sigma}_{N_2}(\mathbf{S}_T) = \hat{\Sigma}_{N_2}^*(\mathbf{S}_T) - \hat{\Sigma}_{N_1 N_2}(\mathbf{S}_T) \hat{\Sigma}_{N_1}^{-1}(\mathbf{S}_T) \hat{\Sigma}_{N_1 N_2}(\mathbf{S}_T)'$ as well. Consistent estimators of $\Sigma_{v\bar{Z}}$ are HAC and DK-HAC estimators.²⁶

Finally, we impose a second-order stationarity condition for $v'_t b_0 \bar{Z}_t(C_T)$ and $v'_t \Sigma_v^{-1} a_{0,\beta} \bar{Z}_t(C_T)$.

Assumption 6.5. *Let $\pi(\mathbf{S})$ equal the Lebesgue measure of $\mathbf{S} \subseteq (0, 1]$. Assume that $\Sigma_{v\bar{Z}}(\mathbf{S}, \mathbf{S}') = \pi(\mathbf{S} \cap \mathbf{S}') \Sigma_{v\bar{Z}}$ where $\mathbf{S}, \mathbf{S}' \subseteq (0, 1]$ and $\Sigma_{v\bar{Z}}$ is p.d.*

²⁵In the supplement we provide primitive sufficient conditions for Assumptions 6.1-6.3.

²⁶In the supplement we provide weak sufficient conditions, even allowing for certain forms of nonstationarity, that ensure this assumption holds.

Assumption 6.5 is implied by a uniform law of large numbers and functional central limit theorem for partial sum processes under second-order stationarity. Under weak IV asymptotics, $T^{-1}\hat{\mathbf{S}}_T$ is not consistent for $\mathbf{S}_0$. Assumption 6.5 is needed in order to show that $N_{1,T}(\cdot)$ and $N_{2,T}(\cdot)$ are asymptotically independent processes. Under strong IV asymptotics we can dispense with Assumption 6.5 because $T^{-1}\hat{\mathbf{S}}_T \xrightarrow{\mathbb{P}} \mathbf{S}_0$ and the limit of the processes $N_{1,T}(\cdot)$ and $N_{2,T}(\cdot)$ have zero covariance when evaluated at a fixed $\mathbf{S}_0$.

Define the LR, LM and AR statistics in this context according to (6.4), replacing $\mathbf{S}_{0,T}$ with $\hat{\mathbf{S}}_T$. We now establish the correct asymptotic null rejection probabilities of the sub-population-estimated plug-in HAR CLR, LM and AR tests under weak identification.

Theorem 6.1. *Let Assumptions 6.1-6.5 hold and suppose $\theta = c/T^{1/2}$ for some nonstochastic $c \in \mathbb{R}^q$. We have: (i) $AR_T(\hat{\mathbf{S}}_T) \xrightarrow{d} \chi_q^2$ under H_0 ; (ii) $LM_T(\hat{\mathbf{S}}_T) \xrightarrow{d} \chi_1^2$ under H_0 ; (iii) $\mathbb{P}_{\beta_0}(LR_T(\hat{\mathbf{S}}_T) > \kappa_\alpha(N_{2,T}(\hat{\mathbf{S}}_T))) \rightarrow \alpha$ where $\mathbb{P}_{\beta_0}(\cdot)$ is the probability computed under H_0 .*

The key to establishing these asymptotic validity results is to show that each of the above statements hold conditional on the realization of $N_{2,T}(\cdot)$. This can be readily established from the facts that the stochastic processes $N_{1,T}(\cdot)$ and $N_{2,T}(\cdot)$ are asymptotically independent by construction, $\hat{\mathbf{S}}_T$ is a function of $N_{2,T}(\cdot)$ and $N_{1,T}(\mathbf{S}_T) \Rightarrow \mathcal{N}(0, I_q)$ under H_0 .²⁷

7 Empirical Evidence on LATE of Monetary Policy

We illustrate our methods by revisiting the identification of monetary policy effects in the framework of Nakamura and Steinsson (2018), introduced in Section 3. They use a bivariate model (3.1) to estimate the causal effect of $\tilde{D}_t$ on $\tilde{Y}_t$, employing both event-study and heteroskedasticity-based identification approaches. The dependent variable $\tilde{Y}_t$ is either the nominal or real 2-Year instantaneous Treasury forward rate and the policy variable is either the daily change in nominal 2-Year Treasury yields or the 30-minute change in a “policy news” series constructed as the first principal component of the unanticipated 30-minute changes in five selected interest rates. Heteroskedasticity-based identification assumes the variance of

²⁷In addition to identification-robust tests of H_0 vs H_1 , since the causal interpretation of β depends upon the sub-population $\mathbf{S}_{0,T}$, practitioners may wish to simultaneously report the result of these tests along with a corresponding estimate of the sub-population. More specifically, failure to reject H_0 should be interpreted as failure to reject that the estimand is equal to $\hat{\beta}_0$, where the estimand is interpreted as a weighted average of the LATEs for the estimated sub-population $\hat{\mathbf{S}}_T$. Given that the tests of H_0 remain asymptotically valid conditional on the realization of $N_{2,T}(\cdot)$ and the fact that $\hat{\mathbf{S}}_T$ is a function of $N_{2,T}(\cdot)$, the tests remain asymptotically valid when interpreted conditional on the value of $\hat{\mathbf{S}}_T$.

the monetary shock rises on FOMC announcement days, while the variance of other shocks remains constant [cf. eq. (3.2)]. FOMC dates define the policy sample $\mathbf{P}$, and analogous non-FOMC dates define the control sample $\mathbf{C}$. Nakamura and Steinsson’s instrument for $\widetilde{D}_t^2$ is defined as $Z_t = \mathbf{1}\{t \in \mathbf{P}\}$, corresponding to the model in Section 3. We focus on the same period: January 1, 2004, to March 19, 2014.

Lewis (2022) recently analyzes this problem by developing a first-stage F -test for weak identification. He finds that weak identification is not rejected when $\widetilde{D}_t$ is the 1-day change in nominal 2-Year Treasury yields, but is strongly rejected when $\widetilde{D}_t$ is the 30-minute policy news series. This supports Nakamura and Steinsson’s (2018) observation that the daily policy variable may suffer from weaker identification. Unlike Nakamura and Steinsson (2018), Lewis (2022) estimates the model using GMM and does not impose the assumption that the non-monetary policy shock η_t has equal variance across the treatment and control samples.

Section 7.1 reports results of our test for full sample identification failure. Section 7.2 presents causal effect estimates based on the most strongly-identified subsample. Section 7.3 provides identification-robust inference results, and Section 7.4 estimates compliers at the individual level and tests the exclusion restriction.

7.1 Testing for Identification Failure

We present the results of our test for identification failure over all sub-populations from Section 4 in Table 1 considering values of π_L from 0.6 to 1. For the 30-minute policy news variable, the F_T^* statistic is very large and identification failure is rejected. This supports the finding in Lewis (2022) and intuition in Nakamura and Steinsson (2018) that the 30-minute policy news variable leads to stronger identification in the full sample. In contrast, for the 1-day change in nominal Treasury yields, identification failure cannot be strongly rejected in the full sample: the F_T^* statistic at $\pi_L = 1$ (i.e., full sample) is only slightly larger than the 1% critical value. The F_T^* statistic increases substantially as π_L decreases and it is very far from the critical values. This is clear evidence that identification is much stronger over subsamples. At $\pi_L = 0.9$ it reaches 33.87, clearly rejecting identification failure in the π -subsample (with $\pi = 0.9$ or 0.95) over which the supremum of $F_T(\mathbf{S}_T)$ is computed. The F_T^* statistic increases monotonically with smaller π_L due to the increasing number of partitions considered. For example, at $\pi_L = 0.8$, F_T^* is 54.78—nearly seven times the full sample value. Overall, the results indicate that strong identification may hold when using a 1-day window, but only within subsamples comprising at most 90% of the data. The weak identification

reported by Lewis (2022) using a 1-day window around FOMC announcements likely does not stem solely from volatility returning to normal after announcements. Rather, a small subsample (10–20% of the data) exhibits weak or failed identification, contributing to the weaker identification exhibited in the full sample.

Table 1: Tests for Identification Failure over all Sub-Populations

$\tilde{D}_t \backslash \pi_L$	F_T^* statistic and critical values				
	0.6	0.7	0.8	0.9	1
30-minute “policy news”	$10^4 \times 95.36$	$10^4 \times 56.50$	$10^4 \times 32.45$	$10^4 \times 18.75$	$10^4 \times 7.42$
1-day nominal Treasury yields	155.69	88.22	54.78	33.88	8.09
1% critical values	17.89	15.58	13.11	10.66	6.46
5% critical values	12.74	10.70	8.88	7.02	3.85

F_T^* statistics for first-stage identification failure. $\tilde{D}_t$ is either the 30-minute policy news series or 1-day change in nominal Treasury yields. π_L is the minimum fraction of the sample over which the supremum of the $F(\mathbf{S}_T)$ is computed. Maximum number of breaks is set to $m_+ = 5$.

7.2 Estimation in Strongly-Identified Subsample

We turn to estimation of π_0 and $\mathbf{S}_{0,T}$ using the methods from Section 5, and then to estimating the LATE of monetary policy based on the strongly-identified subsample, $\hat{\beta}(\hat{\mathbf{S}}_{T,OLS})$, or simply, $\hat{\pi}$ -sample, where $\hat{\pi} = |\hat{\mathbf{S}}_{T,OLS}|/T$.²⁸ We focus on $\hat{\beta}(\hat{\mathbf{S}}_{T,OLS})$; results using $\hat{\beta}(\hat{\mathbf{S}}_{T,FGLS})$ are similar. Figure 3 plots the 1-day changes in 2-Year yields for the control and policy samples and highlights the regimes included in the strongly-identified subsample $\hat{\mathbf{S}}_{T,OLS}$. The estimate $\hat{\pi} = 0.8$ implies that in 80% of the sample, the first-stage is strong and identification holds. In the control sample, the excluded periods include the first seven months of 2005 and the regime surrounding the financial crisis (2007–2009). As shown in the figure, volatility during the crisis period is much higher than in the rest of the control group and higher than the average volatility in the treatment group. This subsample appears to drive the apparent full sample weak identification. Since our method searches for maximum identification strength, it correctly excludes this period when computing π -LATE.²⁹ The interpretation is that in both excluded regimes—especially during the financial crisis—market uncertainty was elevated even on non-FOMC days, violating the identification assumption.

²⁸Recall that when the instrument is binary, the Wald and IV estimands coincide. Hence, the results and intuition developed in Sections 2–3 directly apply to the TSLS estimand analyzed here.

²⁹The other excluded period (January to July 2005) does not display obviously high volatility but shows some persistence, with a short-duration cluster below the mean toward the end.

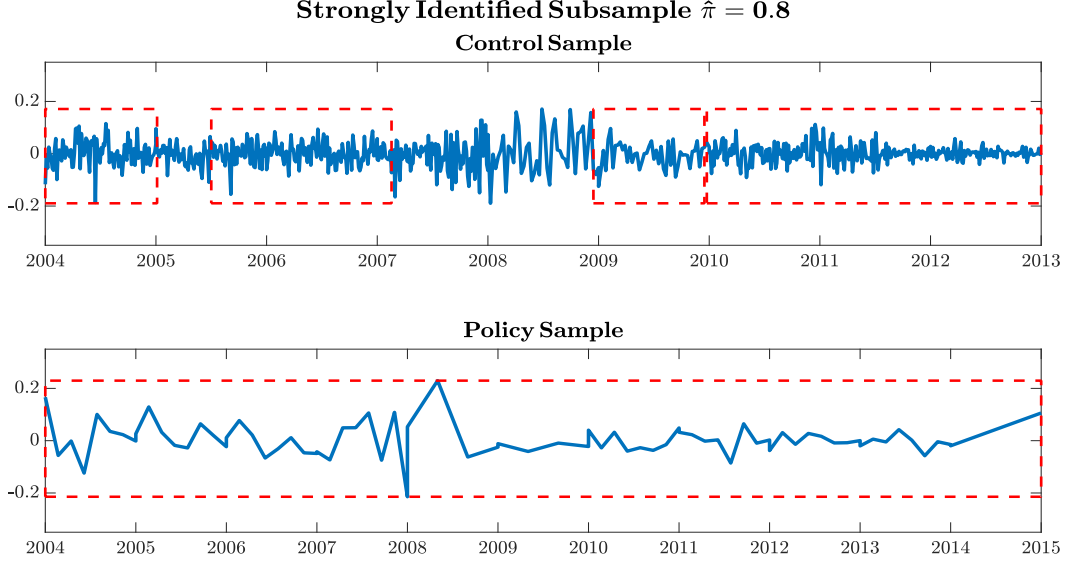


Figure 3: Plot of $\tilde{D}_t$ (2-Year Treasury yields) in the control sample (top panel) and policy sample (bottom panel). The red rectangles indicate subsamples included in the strongly-identified subsample $\hat{\mathbf{S}}_{T,OLS}$ where $\hat{\pi} = 0.8$.

We now estimate the causal effect of monetary policy using the $\hat{\pi}$ -sample, where by construction the LATE is most strongly-identified. We compare these results with full sample estimates obtained using two-stage least squares (TSLS) and GMM, following Nakamura and Steinsson (2018) and Lewis (2022), respectively. Table 2 presents the results. Starting with the full sample estimates: when the policy variable is the 30-minute policy news series, TSLS and GMM yield very similar point estimates for both nominal and real forward rates, and both are statistically significant using standard and robust confidence intervals.³⁰

As noted by Lewis (2022), the assumption that non-monetary shocks have equal variance across treatment and control groups does not bias the TSLS estimates, as they closely match the GMM ones. One explanation is that the GMM estimate of a (capturing reverse causality from forward rates to policy news) is both near zero and statistically significant (not reported). Since potential bias from this assumption is proportional to $a(\sigma_{\eta,P}^2 - \sigma_{\eta,C}^2)$, and a is close to zero, the resulting bias is negligible even if the variances $\sigma_{\eta,P}^2$ and $\sigma_{\eta,C}^2$ differ.

Turning to the case where the policy variable is the 1-day change in 2-Year Treasury yields, the TSLS and GMM estimates differ markedly from each other and from those based on

³⁰The robust confidence intervals for the GMM estimates are based on the subset K -test in Lewis (2022).

the 30-minute policy news series. Notably, the GMM estimate of β is negative for nominal forwards and positive for real forwards, but in neither case is it statistically significant—whether using standard or robust confidence intervals.

Table 2: Estimation of β

dep. var.	30-minute Policy News		1-day 2-Year Yield	
	Nominal	Real	Nominal	Real
Full Sample				
TSLS				
β	1.10**	0.96***	1.14***	0.97***
standard CI	[0.17, 2.02]	[0.41, 1.51]	[0.83, 1.45]	[0.40, 1.565]
GMM				
β	1.07**	0.94***	-0.27	1.31
standard CI	[0.17, 1.98]	[0.36, 1.51]	[-4.90, 4.36]	[-3.74, 6.35]
robust CI	[0.27, 3.25]	[0.44, 2.38]	[-77.27, 0.94]	[-253.70, 1.92]
π -sample based on $\hat{\mathbf{S}}_{T,OLS}$ with $\hat{\pi} = 0.8$				
TSLS				
β	1.11**	0.97***	1.13***	0.92***
standard CI	[0.19, 2.02]	[0.42, 1.51]	[0.92, 1.30]	[0.56, 1.28]
GMM				
β	1.07**	0.94***	0.65*	0.86**
standard CI	[0.17, 1.96]	[0.38, 1.50]	[-0.02, 1.31]	[0.29, 1.43]

TSLS and GMM estimates of β . The GMM estimates allow for changes also in the variance of η_t across regimes. The dependent variable is the 1-day change in either nominal or real 2-Year instantaneous Treasury forward rate. The policy variable is either the 30-minute changes in the “policy news” variable or 1-day changes in the 2-Year nominal Treasury yield. The standard 95% confidence interval is based on the standard normal critical values. For the GMM estimates, the robust 95% confidence interval is based on the subset K -test in [Lewis \(2022\)](#). Asterisks indicate statistical significance at the 10%, 5%, or 1% level based on standard intervals.

As discussed by [Lewis \(2022\)](#), these estimates are difficult to interpret in economically meaningful terms. He also shows that the GMM estimates of a are nonzero and proposed a second dimension of policy news to account for the findings. However, the opposing signs of β across nominal and real forwards complicate this interpretation. Ultimately, he concludes that these results are inconsistent with [Nakamura and Steinsson’s \(2018\)](#) “background noise” view of the non-monetary shock η_t which assumes that its volatility remains unchanged between FOMC and non-FOMC days.

We contribute to this discussion by presenting TSLS and GMM estimates based on the most strongly-identified $\hat{\pi}$ -sample. We focus first on standard confidence intervals and defer weak identification-robust inference to [Table 3](#). The bottom panel of [Table 2](#) shows that, for the 30-minute policy news variable, the TSLS and GMM estimates, including their statistical

significance, are virtually unchanged. As expected—given the apparent strong identification in the full sample—results are broadly similar when using the $\hat{\pi}$ -sample.³¹

Finally, we turn to the $\hat{\pi}$ -sample estimates using the 1-day window for the policy. The GMM estimates differ sharply from those in the full sample: for both nominal and real forwards, they now have the same sign and are statistically significant. This suggests that the opposite signs reported by Lewis (2022) likely stemmed from weak identification, rendering those estimates unreliable.³² Notably, the GMM estimates are now similar in magnitude to those based on the 30-minute policy variable, supporting a more meaningful interpretation.³³

Overall, this analysis highlights the advantage of using the most strongly-identified $\hat{\pi}$ -sample. Given weak identification in the full sample when using 1-day Treasury yields as the policy variable, the corresponding estimates should be discarded. In contrast, evidence from the $\hat{\pi}$ -sample shows that TSLS and GMM produce similar, positive estimates for β , consistent with monetary policy affecting real forward rates, as predicted by New Keynesian models, and supporting the existence of a forward guidance channel.³⁴

7.3 Weak Identification-Robust Inference

We apply the weak identification-robust tests proposed in Section 6 and compare them to existing full sample tests LM_T and LR_T .³⁵ We test the null hypothesis $H_0 : \beta = 0$ against $H_1 : \beta \neq 0$. Results are shown in Table 3. When the policy variable is the 30-minute policy news, identification is strong in the full sample. Accordingly, both the proposed and existing tests yield similar results: all tests reject at the 5% level for both nominal and real forwards (not reported for brevity). Nakamura and Steinsson (2018) showed that the effect of policy news peaks at the 2-Year maturity and declines with longer maturities. Consistent with this,

³¹The confidence intervals in the $\hat{\pi}$ -sample are even slightly tighter.

³²While the TSLS estimates are nearly unchanged from the full sample, this should not be taken as evidence of their reliability. Under weak IVs, their similarity to the $\hat{\pi}$ -sample results may simply be coincidental.

³³We also verified that the GMM estimate of a is 0.70 for nominal forwards and -0.91 for real forwards. It is intuitive that the estimate of a is close to zero when using a 30-minute window but significantly different from zero with a 1-day window. In the narrow 30-minute window around an FOMC announcement, reverse causality from $\tilde{Y}_t$ to $\tilde{D}_t$ is limited, as monetary news is more pronounced than other shocks—though some endogeneity may still arise from omitted factors affecting both. In contrast, over a full day, asset price movements can influence short-term interest rates, making reverse causality more likely.

³⁴However, the results do not yet support a second meaningful dimension of news, as proposed by Lewis (2022), since the sign of the GMM estimate of a is unstable across nominal and real forwards. Regarding Nakamura and Steinsson’s (2018) “background noise” interpretation of non-monetary shocks, we find no clear evidence against it: in the $\hat{\pi}$ -sample, identification appears strong, and TSLS and GMM estimates consistently share the same sign and similar magnitudes.

³⁵We do not report the AR_T test since for $q = 1$ it is equivalent to the LM_T test.

Table 3: Identification-Robust Inference on β

2-Year Forwards	30-minute Policy News						1-day change in 2-Year Yields					
	Nominal			Real			Nominal			Real		
α	0.10	0.05	0.01	0.10	0.05	0.01	0.10	0.05	0.01	0.10	0.05	0.01
LM_T	✓	✓	×	✓	✓	✓	×	×	×	✓	✓	×
CLR_T	✓	✓	×	✓	✓	✓	×	×	×	✓	✓	×
$LM_T(\hat{\mathbf{S}}_T)$	✓	✓	×	✓	✓	✓	✓	✓	×	✓	✓	✓
$CLR_T(\hat{\mathbf{S}}_T)$	✓	✓	×	✓	✓	×	✓	✓	✓	✓	✓	✓

Weak identification-robust tests on β . The dependent variable $\tilde{Y}_t$ is the 2-Year forward rates. $\tilde{D}_t$ is either the 30-minute policy news series or the 1-day nominal Treasury yields. Significance levels are $\alpha = 0.10, 0.05, 0.01$. A ✓ indicates rejection H_0 ; a × non-rejection.

we find weaker statistical significance for the 5-Year. In line with theoretical predictions, the long-run impact of monetary policy shocks on real interest rates (i.e., the 10 Year forwards) approaches zero (not reported).

Let us instead consider the 1-day change in 2-Year yields as the policy variable. The existing LM_T and LR_T tests do not reject the null at any standard significance level for nominal forwards, and at the 1% level for real forwards. In contrast, our proposed tests based on $\hat{\mathbf{S}}_T$ show much stronger rejections, aligning with the results using the 30-minute policy news series, which indicate a positive causal effect on 2-Year forwards.

7.4 Identification and Estimation of Compliers, and Exclusion Restriction

We now identify compliers individually by applying Theorem 2.1. Under heteroskedasticity-based identification, the sample rolling window averages in Theorem 2.1 correspond to rolling window variances, i.e., $\bar{D}_{P,t,n_1}$ and $\bar{D}_{C,t,n_0}$ are equal to $\bar{\sigma}_{P,t,n_1}^2$ and $\bar{\sigma}_{C,t,n_0}^2$ in this context, where $\bar{\sigma}_{P,t,n_1}^2$ and $\bar{\sigma}_{C,t,n_0}^2$ are defined analogously but using $\tilde{D}_t^2$ for D_t .³⁶ We use two-sided rolling windows with $n_0 = 101$ and $n_1 = 15$.³⁷ For each t_0 we test the null hypothesis $H_0 : \mathbb{E}(D_{t_0}^2(1)) - \mathbb{E}(D_{t_0}^2(0)) = 0$ (t_0 is a non-complier) versus the one tailed alternative $H_1 : \mathbb{E}(D_{t_0}^2(1)) - \mathbb{E}(D_{t_0}^2(0)) > 0$ (t_0 is a complier). We use the t -statistic

$$t_{t_0} = \begin{cases} \frac{\sqrt{n_0}(\bar{\sigma}_{P,t_0,n_1}^2 - \bar{\sigma}_{C,t_0-1,n_0}^2)}{\sqrt{J_{\text{HAC},t_0-1}(1+n_0/n_1)}} & t_0 \in \mathbf{P} \\ \frac{\sqrt{n_0}(\bar{\sigma}_{P,s^*(t_0),n_1}^2 - \bar{\sigma}_{C,t_0,n_0}^2)}{\sqrt{J_{\text{HAC},t_0}(1+n_0/n_1)}} & t_0 \in \mathbf{C}, \end{cases}$$

³⁶More specifically, we have $\bar{\sigma}_{C,t_0-1,n_0}^2 = \frac{1}{n_0} \sum_{s \in N_0(t_0)} \tilde{D}_s^2$, and $\bar{\sigma}_{P,t_0,n_1}^2 = \frac{1}{n_1} \sum_{s \in N_1(t_0)} \tilde{D}_s^2$.

³⁷We also used one-sided backward rolling windows, and the results are essentially unchanged.

where J_{HAC,t_0} is the Newey-West estimator with $\lfloor n_0^{1/3} \rfloor$ lags applied to $\tilde{D}_s^2 - n_0^{-1} \sum_{k \in N_0(t_0+1)} \tilde{D}_k^2$.

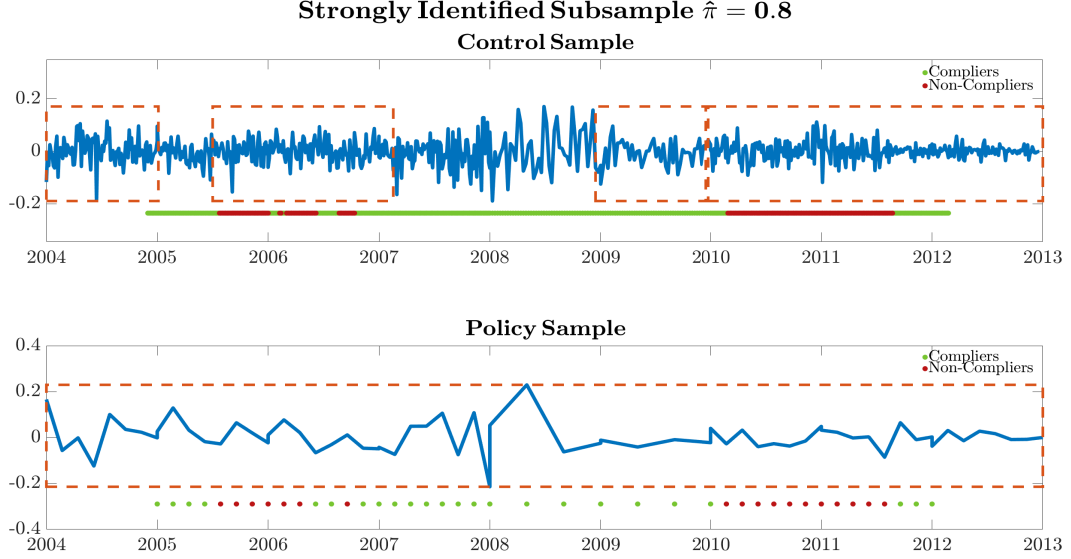


Figure 4: Plot of $\tilde{D}_t$ (2-Years Treasury yields) in the control sample (top panel) and policy sample (bottom panel). The orange rectangles indicate subsamples included in the strongly-identified set $\hat{\mathbf{S}}_{T,OLS}$ where $\hat{\pi} = 0.8$. Green filled circles indicate compliers; red filled circles indicate non-compliers. Time points without colored markers correspond to cases where rolling sample variances could not be computed due to proximity to the start or end of the sample.

The results at the 0.10 significance level are shown in Figure 4. Approximately 62% of observations are classified as compliers. The non-compliers are mostly concentrated in the period from 2010 to mid-2011, which corresponds to the early phase of the zero lower bound (ZLB) period following the 2008–09 recession. During this time, the Fed relied primarily on qualitative forward guidance—e.g., stating that economic conditions were “likely to warrant exceptionally low levels of the federal funds rate for some time.” In August 2011, the Fed shifted to more explicit, calendar-based guidance, stating that such conditions were “likely to warrant exceptionally low levels of the federal funds rate at least through mid-2013.” Thus, the non-complier period aligns with the phase of the ZLB when forward guidance was less aggressive as the policy announcements by then only imply a near zero-rate horizon for the following three to four quarters, significantly shorter than what the ZLB constraint would actually have implied. Non-compliers also appear intermittently between mid-2005 and mid-2006. One possible explanation is that this period coincided with the latter stages of the

Fed’s highly predictable tightening cycle, during which policy rates were increased by 25 basis points at successive meetings and forward guidance emphasized a “measured pace” of tightening. Because markets largely anticipated these policy actions, FOMC announcements may have conveyed relatively little new information, resulting in weaker local compliance. However, these episodes are less persistent than those observed during the early ZLB period.³⁸

Finally, we use Theorem 2.1 and Proposition 2.3 to test the exclusion restriction (cf. Assumption 2.2). Here the exclusion restriction means that the volatility of the non-monetary policy shocks remain constant across FOMC and non-FOMC days. We consider the whole set of non-compliers $\mathcal{NC}$. We test the null hypothesis that the exclusion restriction holds by using the following t -statistic:

$$t_{\text{exclusion}} = \frac{\sqrt{|\mathcal{NC}_C|} (\bar{Y}_{\mathcal{NC}_P^s} - \bar{Y}_{\mathcal{NC}_C^s})}{\sqrt{J_{\text{HAC}, \mathcal{NC}^s}(1 + |\mathcal{NC}_C|/|\mathcal{NC}_P|)}}$$

where $\bar{Y}_{\mathcal{NC}_P} = \frac{1}{|\mathcal{NC}_P|} \sum_{t \in \mathcal{NC}_P} Y_t$, $\bar{Y}_{\mathcal{NC}_C} = \frac{1}{|\mathcal{NC}_C|} \sum_{t \in \mathcal{NC}_C} Y_t$, $Y_t = \tilde{D}_t \tilde{Y}_t$ and $J_{\text{HAC}, \mathcal{NC}^s}$ is the Newey-West estimator applied to $\bar{Y}_{\mathcal{NC}_P} - \bar{Y}_{\mathcal{NC}_C}$. For the real (nominal) 2-Year forward rate, we find $t_{\text{exclusion}} = 0.01$ ($t_{\text{exclusion}} = 0.60$), and thus fail to reject the exclusion restriction. This supports the “background noise” interpretation of Nakamura and Steinsson (2018) as non-monetary shocks appear to have constant variance across FOMC and non-FOMC days.

8 Conclusions

This paper discusses identification, estimation and inference on dynamic LATE. We show that compliers can be identified individually and the exclusion restriction can be tested using a t -test. While weak identification is common in the full sample in practice, strong identification often appears to hold in a sizable subsample. We propose a method to isolate this strongly-identified subsample, enabling consistent estimation and inference.

³⁸The set of compliers does not coincide with the set of observations in the strongly-identified $\hat{\pi}$ -sample. However, this does not necessarily imply a violation of monotonicity (cf. Proposition 2.2). First, the complier status is determined via a t -test whereas the strongly-identified $\hat{\pi}$ -sample is determined via estimation. Second, the complier status is determined by the rolling window variances at each t , whereas inclusion in the $\hat{\pi}$ -sample depends on how these variances contribute to the average volatility in the control sample relative to that in the policy sample. In other words, inclusion in the $\hat{\pi}$ -sample is a global criterion, influenced by volatility elsewhere in the sample, while complier status reflects whether the first stage holds individually.

Supplemental Materials: The online supplement [cf. [Casini et al. \(2026a\)](#)] includes Monte Carlo simulations, proofs of the results of Sections 2-4 and 6. The non-online supplement [cf. [Casini et al. \(2026b\)](#)] contains the theoretical results and corresponding proofs for the estimators in Section 5 and additional results.

References

- ABADIE, A. (2003): “Semiparametric Instrumental Variable Estimation of Treatment Response Models,” *Journal of Econometrics*, 113(2), 231–263.
- ANDERSON, T. W., AND H. RUBIN (1949): “Estimation of the Parameters of a Single Equation in a Complete System of Stochastic Equations,” *Annals of Mathematical Statistics*, 20(1), 46–63.
- ANDREWS, D. W. K. (1991): “Heteroskedasticity and Autocorrelation Consistent Covariance Matrix Estimation,” *Econometrica*, 59(3), 817–858.
- ANDREWS, D. W. K., M. MOREIRA, AND J. H. STOCK (2006): “Optimal Two-Sided Invariant Similar Tests for Instrumental Variables Regression,” *Econometrica*, 74(3), 715–752.
- ANDREWS, I. (2018): “Valid Two-Step Identification-Robust Confidence Sets for GMM,” *Review of Economics and Statistics*, 100(2), 337–348.
- ANDREWS, I., J. H. STOCK, AND L. SUN (2019): “Weak Instruments in IV Regression: Theory and Practice,” *Annual Review of Economics*, 11, 727–753.
- ANGRIST, J. D., K. GRADY, AND G. W. IMBENS (2000): “The Interpretation of Instrumental Variables Estimators in Simultaneous Equations Models with an Application to the Demand for Fish,” *Review of Economic Studies*, 67(3), 499–527.
- ANGRIST, J. D., G. W. IMBENS, AND D. B. RUBIN (1996): “Identification of Causal Effects Using Instrumental Variables,” *Journal of the American Statistical Association*, 91(434), 444–455.
- ANGRIST, J. D., AND G. M. KUERSTEINER (2011): “Causal Effects of Monetary Shocks: Semiparametric Conditional Independence Tests with a Multinomial Propensity Score,” *Review of Economics and Statistics*, 93(3), 725–747.
- ANTOINE, B., AND O. BOLDEA (2018): “Efficient Estimation with Time-Varying Information and the New Keynesian Phillips Curve,” *Journal of Econometrics*, 204(2), 268–300.
- BAI, J., AND P. PERRON (2003): “Computation and Analysis of Multiple Structural Changes,” *Journal of Applied Econometrics*, 18, 1–22.
- BARNICHON, B., AND G. MESTERS (2025): “Innovations Meet Narratives—Improving the Power-Credibility Trade-Off in Macro,” *Working paper*.
- BLANDHOL, C., J. BONNEY, M. MOGSTAD, AND A. TORGOVITSKY (2025): “When is TSLS Actually LATE?,” *NBER Working Paper 29709*.
- CASINI, A. (2023): “Theory of Evolutionary Spectra for Heteroskedasticity and Autocorre-

- lation Robust Inference in Possibly Misspecified and Nonstationary Models,” *Journal of Econometrics*, 235(2), 372–392.
- CASINI, A., AND A. MCCLOSKEY (2025): “Identification, Estimation and Inference in High-Frequency Event Study Regressions,” *arXiv preprint arXiv 2406.15667*.
- CASINI, A., A. MCCLOSKEY, L. ROLLA, AND R. PALA (2026a): “Online Supplement to ”Dynamic Local Average Treatment Effects in Time Series”,” *Unpublish manuscript*.
- (2026b): “Supplement Not For Online Publication to ”Dynamic Local Average Treatment Effects in Time Series”,” *Unpublish manuscript*.
- CASINI, A., AND P. PERRON (2019): “Structural Breaks in Time Series,” *Oxford Research Encyclopedia of Economics and Finance*, Oxford University Press.
- (2024): “Change-Point Analysis of Time Series with Evolutionary Spectra,” *Journal of Econometrics*, 242(2), 105811.
- GERTLER, M., AND P. KARADI (2015): “Monetary Policy Surprises, Credit Costs, and Economic Activity,” *American Economic Journal: Macroeconomics*, 7(1), 44–76.
- HECKMAN, J. J. (1996): “Identification of Causal Effects Using Instrumental Variables: Comment,” *Journal of the American Statistical Association*, 91(434), 459–462.
- HECKMAN, J. J., AND E. VYTLACIL (2005): “Structural Equations, Treatment Effects, and Econometric Policy Evaluation,” *Econometrica*, 73(3), 669–738.
- (2007): “Econometric Evaluation of Social Programs, Part I: Causal Models, Structural Models and Econometric Policy Evaluation,” in *Handbook of Econometrics*, ed. by J. J. Heckman, and E. E. Leamer, vol. 6B, pp. 4779–4874. Elsevier.
- IMBENS, G. W. (2010): “Better LATE Than Nothing: Some Comments on Deaton (2009) and Heckman and Urzua (2009),” *Journal of Economic Literature*, 48(2), 399–423.
- IMBENS, G. W., AND J. D. ANGRIST (1994): “Identification and Estimation of Local Average Treatment Effects,” *Econometrica*, 62(2), 467–475.
- IMBENS, G. W., AND D. B. RUBIN (1997): “Estimating Outcome Distributions for Compliers in Instrumental Variables Models,” *Review of Economic Studies*, 64(4), 555–574.
- KÄNZIG, D. R. (2021): “The Macroeconomic Effects of Oil Supply News: Evidence from OPEC Announcements,” *American Economic Review*, 111(4), 1092–1125.
- KILIAN, L. (2009): “Not All Oil Price Shocks Are Alike: Disentangling Demand and Supply Shocks in the Crude Oil Market,” *American Economic Review*, 99(3), 1053–1069.
- KITAGAWA, T. (2015): “A Test for Instrument Validity,” *Econometrica*, 83(5), 2043–2063.
- KITAGAWA, T., W. WANG, AND M. XU (2025): “Nonlinearity in Dynamic Causal Effects: Making the Bad into the Good, and the Good into the Great?,” *Journal of Business Economics and Statistics*, forthcoming.
- KLEIBERGEN, F. (2002): “Pivotal Statistics for Testing Structural Parameters in Instrumental Variables Regression,” *Econometrica*, 70(5), 1781–1803.
- KOLESÁR, M., AND M. PLAGBORG-MØLLER (2025): “Dynamic Causal Effects in a Nonlinear World: the Good, the Bad, and the Ugly,” *arXiv preprint arXiv 2411.10415*.
- LEWIS, D. J. (2022): “Robust Inference in Models Identified via Heteroskedasticity,” *Review of Economics and Statistics*, 104(3), 510–524.
- MAGNUSSON, L. M., AND S. MAVROEIDIS (2014): “Identification Using Stability Restriction

- tions,” *Econometrica*, 82(5), 1799–1851.
- MERTENS, K., AND M. O. RAVN (2013): “The Dynamic Effects of Personal and Corporate Income Tax Changes in the United States,” *American Economic Review*, 103(4), 1212–1247.
- MONTIEL OLEA, J. L., AND C. PFLUEGER (2013): “A Robust Test for Weak Instruments,” *Journal of Business and Economic Statistics*, 31(3), 358–369.
- MONTIEL OLEA, J. L., J. H. STOCK, AND M. W. WATSON (2021): “Inference in SVARs Identified with External Instruments,” *Journal of Econometrics*, 225(1), 74–87.
- MOREIRA, M. (2003): “A Conditional Likelihood Ratio Test for Structural Models,” *Econometrica*, 71(4), 1027–1048.
- NAKAMURA, E., AND J. STEINSSON (2018): “High Frequency Identification of Monetary Non-Neutrality: The Information Effect,” *Quarterly Journal of Economics*, 133(3), 1283–1330.
- NEWKEY, W. K., AND K. D. WEST (1987): “A Simple Positive Semidefinite, Heteroskedastic and Autocorrelation Consistent Covariance Matrix,” *Econometrica*, 55(3), 703–708.
- PLAGBORG-MØLLER, M., AND C. K. WOLF (2022): “Instrumental Variable Identification of Dynamic Variance Decompositions,” *Journal of Political Economy*, 130(8), 2164–2202.
- RAMBACHAN, A., AND N. SHEPHARD (2021): “When Do Common Time Series Estimands Have Nonparametric Causal Meaning?,” *Unpublished Manuscript, Harvard University*.
- RAMEY, V. A. (2011): “Identifying Government Spending Shocks: It’s All in the Timing,” *Quarterly Journal of Economics*, 126(1), 1–50.
- RIGOBON, R. (2003): “Identification Through Heteroskedasticity,” *Review of Economics and Statistics*, 85(4), 777–792.
- RIGOBON, R., AND B. SACK (2003): “Measuring the Reaction of Monetary Policy to the Stock Market,” *The Quarterly Journal of Economics*, 118(2), 639–669.
- ROMER, C. D., AND D. H. ROMER (1989): “Does Monetary Policy Matter? A New Test in the Spirit of Friedman and Schwartz,” *NBER Macroeconomics Annual*, 4, 121–184.
- (2004): “A New Measure of Monetary Shocks: Derivation and Implications,” *American Economic Review*, 94(4), 1055–1084.
- (2010): “The Macroeconomic Effects of Tax Changes: Estimates Based on a New Measure of Fiscal Shocks,” *American Economic Review*.
- RUBIN, D. (1974): “Estimating Causal Effects of Treatments in Randomized and Non-Randomized Studies,” *Journal of Educational Psychology*, 66(5), 688–701.
- SOJITRA, R. B., AND V. SYRGKANIS (2025): “Dynamic Local Average Treatment Effects,” *arXiv preprint arXiv:2405.01463*.
- STAIGER, D., AND J. H. STOCK (1997): “Instrumental Variables Regression with Weak Instruments,” *Econometrica*, 65(3), 557–586.
- STOCK, J. H., AND M. W. WATSON (2018): “Identification and Estimation of Dynamic Causal Effects in Macroeconomics Using External Instruments,” *Economic Journal*, 128(610), 917–948.

Online Supplement to “Dynamic Local Average Treatment Effects in Time Series”

ALESSANDRO CASINI

University of Rome Tor Vergata

ADAM McCLOSKEY

University of Colorado at Boulder

LUCA ROLLA

University of Rome Tor Vergata

RAIMONDO PALA

University of Rome Tor Vergata

22nd July 2026

Abstract

This supplemental material is for online publication. It contains additional results on identification-robust inference, Monte Carlo simulations and proofs of the results.

S.A Critical Values of F_T^*

Table 4: Critical values of F_T^*

$\alpha = 0.10$						
$\pi_L \backslash q$	1	2	3	4	5	10
0.50	11.56	7.68	6.04	5.11	4.51	3.15
0.60	10.22	6.96	5.53	4.72	4.19	3.00
0.70	8.54	6.10	4.92	4.33	3.87	2.82
0.80	6.89	5.26	4.32	3.77	3.51	2.59
0.90	5.46	4.16	3.47	3.15	2.92	2.29
1.00	2.68	2.32	2.07	1.94	1.85	1.61
$\alpha = 0.05$						
$\pi_L \backslash q$	1	2	3	4	5	10
0.50	14.03	9.02	6.94	5.78	5.03	3.44
0.60	12.74	8.21	6.38	5.34	4.70	3.29
0.70	10.70	7.27	5.75	4.97	4.38	3.10
0.80	8.88	6.36	5.10	4.39	3.96	2.86
0.90	7.02	5.19	4.15	3.71	3.41	2.56
1.00	3.85	3.05	2.58	2.37	2.23	1.84
$\alpha = 0.01$						
$\pi_L \backslash q$	1	2	3	4	5	10
0.50	19.83	11.65	8.88	7.16	6.22	4.01
0.60	17.89	10.88	8.21	6.74	5.75	3.89
0.70	15.58	9.78	7.62	6.37	5.50	3.67
0.80	13.11	8.72	6.69	5.61	4.93	3.45
0.90	10.66	7.28	5.67	4.80	4.38	3.10
1.00	6.46	4.53	3.72	3.30	3.00	2.35

S.B Additional Results on Identification-Robust Inference

We present the sufficiency results referenced in Section 6 as well as other results.

S.B.1 Known Sub-Population

When $\mathbf{S}_{0,T}$ is known, it is straightforward to use existing tests in the identification-robust linear IVs literature to test H_0 with known optimality properties under certain conditions. However, one must be careful in defining the appropriate statistics when applying existing tests in this setting in order to maintain efficiency.

With fixed regressors X and Z and reduced-form errors v that are i.i.d. across rows with each row being bivariate normally distributed with a mean of zero and a known nonsingular covariance matrix Σ_v , careful application of the results of [Andrews, Moreira, and Stock \(2006\)](#) imply that $\bar{Z}(C_{0,T})'y$ is a sufficient statistic for $(\beta, \theta)'$. This implies that there can be no loss in efficiency from focusing on tests that are functions of only $\bar{Z}(C_{0,T})'y$. On the other hand, the following proposition implies a loss in efficiency from tests that are functions of only $Z'M_X y$, which a casual user may be tempted to use when constructing identification-robust tests for β .

Assumption S.B.1. $\text{rank}((I - P_{M_X Z}) M_X C_{0,T} Z) > 0$.

Assumption [S.B.1](#) requires that the residualized instrument variation in $\mathbf{S}_{0,T}$ is not fully spanned by the full-sample residualized instrument. This rules out degenerate designs in which $M_X C_{0,T} Z$ and $M_X Z$ have the same relevant span.

Proposition S.B.1. *For the model in [\(6.1\)](#) with fixed regressors X and Z and reduced-form errors v that are i.i.d. across rows with each row being bivariate normal with a zero mean and known p.d. covariance matrix Σ_v . If $\pi_0 < 1$ and Assumption [S.B.1](#) holds, then $Z'M_X y$ is not sufficient for $(\beta, \theta)'$.*

This result implies that existing tests (i.e., CLR, LM and AR tests) and extensions thereof are not efficient when Z is treated as the matrix of IVs rather than $C_{0,T} Z$.

The results of [Andrews, Moreira, and Stock \(2006\)](#) imply that the CLR, LM and AR tests have limiting null rejection probabilities equal to α under weak instrument asymptotics. In addition, results in [Andrews, Moreira, and Stock \(2006\)](#) imply asymptotic near-optimality properties of the CLR test under a stronger set of assumptions that may not hold in the presence of serial correlation in $\{v_t\}$.

S.B.2 Unknown Sub-Population

In [Section 6](#) we propose CLR, LM and AR statistics where we plug-in the estimate for the unknown sub-population: $CLR_T(\hat{\mathbf{S}}_T)$, $LM_T(\hat{\mathbf{S}}_T)$ and $AR_T(\hat{\mathbf{S}}_T)$. To motivate their use in the unknown sub-population setting, we establish the analog of the sufficiency result of [Andrews, Moreira, and Stock \(2006\)](#) in this setting.

Proposition S.B.2. *For the model in [\(6.1\)](#) with fixed regressors X and Z and reduced form errors v that are i.i.d. across rows with each row being bivariate normal with a zero mean,*

known *p.d.* covariance matrix Σ_v and an unknown sub-population $\mathbf{S}_{0,T} \in \mathcal{S}$, the Gaussian process $\{\bar{Z}(C_T)'y\}_{\mathbf{S}_T \in \mathcal{S}}$ is a sufficient statistic for $(\beta, \theta)'$.

Since the AR, LM and CLR statistics are only functions of $\{\bar{Z}(C_T)'y\}_{\mathbf{S}_T \in \mathcal{S}}$, this result implies that these tests entail no loss in efficiency relative to tests using the entire data.

We finally show that under weak IV asymptotics, $\hat{\mathbf{S}}_T$ is the maximum likelihood estimator of $\mathbf{S}_T$ under H_0 , so that the sub-population estimate we use when constructing and interpreting the identification-robust tests is efficient.

Proposition S.B.3. *For the model in (6.1) with fixed regressors X and Z and reduced form errors v that are i.i.d. across rows with each row being bivariate normal with a zero mean, *p.d.* covariance matrix Σ_v and unknown sub-population $\mathbf{S}_T \in \mathcal{S}$, if Assumptions 6.1 and 6.3-6.4 hold and $\theta = c/T^{1/2}$ for some nonstochastic $c \in \mathbb{R}^q$, $\hat{\mathbf{S}}_T$ is asymptotically equivalent to the maximum likelihood estimator of $\mathbf{S}_T$ under H_0 .*

The result in Proposition S.B.3 continues to hold under strong IVs, i.e., $\theta \neq 0$ is fixed. See Casini, McCloskey, Rolla, and Pala (2026b).

Magnusson and Mavroeidis (2014) propose tests for β that are robust to changes in θ whether through persistent time variation or breaks. For the latter case, they assumed the number of breaks is known, whereas our approach does not require this prior knowledge. For the weak IVs case, Magnusson and Mavroeidis (2014) consider a single break and build split-sample tests based on the sufficient statistic $\{Z(\tau)'y\}_{\tau \in [0,1]}$ where $Z(\tau) = [[\{Z_t'\}_{t=1}^{\lfloor \tau T \rfloor} \mathbf{0}'] : [\mathbf{0}' \{Z_t'\}_{t=\lfloor \tau T \rfloor + 1}^T]'$. This high-dimensional statistic effectively uses the full data sequence evaluated at all potential split points. In contrast, our framework focuses on subsamples where θ is nonzero, allowing us to construct a lower-dimensional sufficient statistic. In other words, their statistic is not minimal sufficient [cf. Lehmann and Romano (2005)], whereas ours is—making our tests more efficient in this setting.

S.C Monte Carlo Simulations

S.C.1 Finite-Sample Size and Power of F_T^*

We study the finite-sample rejection frequencies of the F_T^* test using a simulation experiment calibrated to real data from the analysis in Nakamura and Steinsson (2018) introduced in Section 3. We use the same sequences of policy dates and control dates, $\mathbf{P}$ and $\mathbf{C}$, to construct

the instrument as

$$Z_t = \begin{cases} \frac{T}{T_P}, & t \in \mathbf{P} \\ -\frac{T}{T_C}, & t \in \mathbf{C}. \end{cases}$$

Under heteroskedasticity-based identification the first-stage equation can be equivalently written as $D_t = \theta Z_t D_t + e_t$ for some θ and e_t [cf. [Rigobon and Sack \(2003\)](#) and [Lewis \(2022\)](#)]. Hence, we generate D_t according to the following data-generating process (DGP):

$$D_t = \begin{cases} \frac{e_t}{1-\theta_1 Z_t}, & t \leq \lfloor T/4 \rfloor \\ \frac{e_t}{1-\theta_2 Z_t}, & \lfloor T/4 \rfloor + 1 \leq t \leq \lfloor T/4 \rfloor + \lfloor (1-\pi_0) T \rfloor \\ \frac{e_t}{1-\theta_3 Z_t}, & \lfloor T/4 \rfloor + \lfloor (1-\pi_0) T \rfloor + 1 \leq t \leq T, \end{cases} \quad (\text{S.C.1})$$

where $\pi_0 = 0.4, 0.6, 0.8$ and $T = 400$. We set T_P and T_C equal to the number of policy and control dates that occur in the first T observations in [Nakamura and Steinsson's \(2018\)](#) sample. We specify $e_t = \rho_e e_{t-1} + v_{e,t}$, where $\rho_e \in \{0, 0.25, 0.5, 0.75\}$, $v_{e,t} \sim \text{i.i.d. } \mathcal{N}(0, \sigma_v^2)$ and σ_v^2 is set equal to the sample variance of the policy variable (2-years nominal Treasury yields). We set $\theta_1 = \theta_2 = \theta_3 = 0$ under the null hypothesis. Under the alternative, $\theta_1 = \theta_3 > 0$ and $\theta_2 = 0$.

We also consider the following DGP:³⁹

$$Y_t = \beta D_t + \gamma_1 X_t + u_t, \quad (\text{S.C.2})$$

where $X_t \sim \text{i.i.d. } \mathcal{N}(1, 1)$ for all t and

$$D_t = \begin{cases} \theta_1 Z_t + \gamma_2 X_t + e_t, & t \leq \lfloor T/4 \rfloor \\ \theta_2 Z_t + \gamma_2 X_t + e_t, & \lfloor T/4 \rfloor + 1 \leq t \leq \lfloor T/4 \rfloor + \lfloor (1-\pi_0) T \rfloor \\ \theta_3 Z_t + \gamma_2 X_t + e_t, & \lfloor T/4 \rfloor + \lfloor (1-\pi_0) T \rfloor + 1 \leq t \leq T, \end{cases} \quad (\text{S.C.3})$$

$Z_t \sim \text{i.i.d. } \mathcal{N}(1, 1)$, and u_t and e_t are i.i.d. jointly normal with mean zero and covariance

$$\Sigma_{ue} = \begin{bmatrix} 1 & \rho \\ \rho & 1 \end{bmatrix}, \quad (\text{S.C.4})$$

³⁹This is also used to compare the performance of the estimators of β discussed in [Section 5](#) and of the identification-robust tests discussed in [Section 6](#).

with $\rho \in \{0.25, 0.50, 0.75\}$ and $\gamma_1 = \gamma_2 \in \{0, 1\}$. Under the null hypothesis we set $\theta_1 = \theta_2 = \theta_3 = 0$. Under the alternative hypothesis we set $\theta_1 = \theta_3 = dT^{-1/2}$ with $d \in \{4, 10, 16\}$ and $\theta_2 = 0$. We also consider two additional specifications for θ_2 . In the first, $\theta_2 = -0.5dT^{-1/2}$, so θ_2 has the opposite sign to θ_1 and θ_3 . In the second, $\theta_2 = -0.5$. In both cases the instrument is relevant throughout the sample (no identification failure). However, in the second regime the first-stage effect tends to offset instrument relevance in the other regimes because of the sign reversal. We set $\pi_0 \in \{0.6, 0.8\}$ and $T = 200$. We also consider a variant of the DGP with serially correlated data. We assume $u_t = \rho_u u_{t-1} + v_{u,t}$ and $e_t = \rho_e e_{t-1} + v_{e,t}$ with $v_{u,t}$ and $v_{e,t}$ being jointly normal with mean zero and covariance Σ_{ue} as in (S.C.4) with $\rho \in \{0.25, 0.50, 0.75\}$.

Throughout the simulation study, F_T^* is implemented with $\pi_L = 0.6$, $\epsilon = 0.05$, $m_+ = 5$. For both F_T^* and the full sample F_T we use the Newey-West estimator with bandwidth equal to the popular rule $\lfloor T^{-1/3} \rfloor$ for $\hat{J}(\cdot)$.⁴⁰ The significance level is 5% and the number of simulations is 5,000. Figure 5 plots the rejection rates of F_T^* and the full sample F_T for the calibrated DGP in (S.C.1) with $\pi_0 = 0.6$. Both F_T and F_T^* yield accurate rejection rates, providing evidence for the reliability of our empirical results in Section 7. F_T^* is much more powerful than F_T by about across all values of ρ_e .

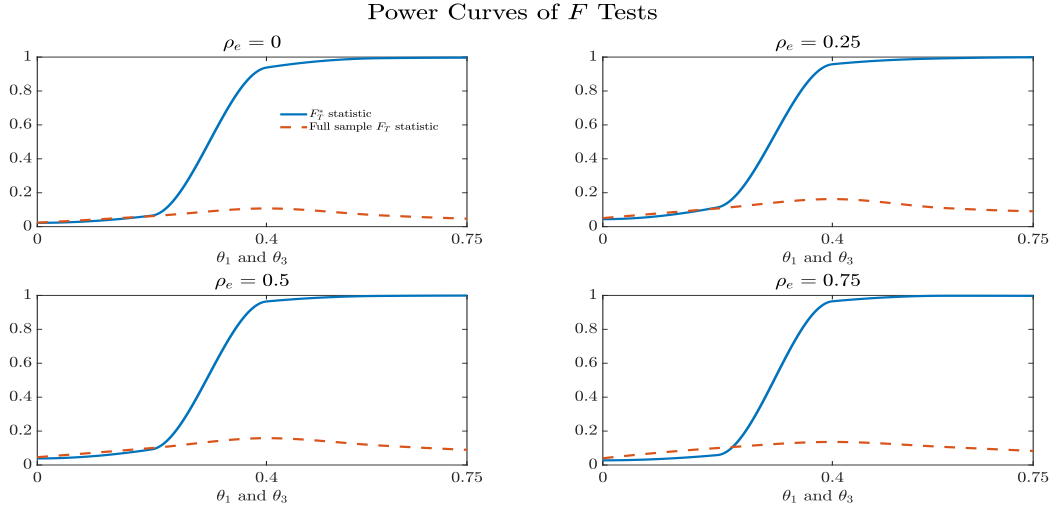


Figure 5: Power curves of F_T^* and F_T for the DGP calibrated to Nakamura and Steinsson's (2018) sample with $T = 400$.

⁴⁰We also consider data-dependent bandwidths. The results are similar and not reported.

Moving on to the DGP specified in (S.C.2)-(S.C.3), Table 5 presents the size and size-adjusted power of the F tests for the case $\theta_2 = -0.5$ under the alternative. Both statistics show accurate null rejection rates. The size-adjusted power of F_T^* is much higher than that of the full sample F_T for all values of d . Power gains are larger when $\pi_0 = 0.6$ than when 0.8 since in the former case the strongly-identified subsample is smaller, making the full sample F_T rely more heavily on subsamples that suffer from identification failure.

Table 5: Size and size-adjusted power of F tests under alternative hypothesis with $\theta_2 = -0.5$

$\rho = 0.25, \pi_0 = 0.6$	$\theta_1 = \theta_2 = \theta_3 = 0$ (null)	$d = 4$	$d = 8$	$d = 12$	$d = 16$	$d = 20$	$d = 24$
full sample F_T	0.061	0.023	0.123	0.478	0.797	0.933	0.981
F_T^*	0.058	0.128	0.426	0.838	0.971	0.996	1.000
$\rho = 0.75, \pi_0 = 0.6$	$\theta_1 = \theta_2 = \theta_3 = 0$ (null)	$d = 4$	$d = 8$	$d = 12$	$d = 16$	$d = 20$	$d = 24$
full sample F_T	0.063	0.034	0.107	0.394	0.712	0.879	0.957
F_T^*	0.035	0.142	0.389	0.774	0.943	0.989	0.998
$\rho = 0.25, \pi_0 = 0.8$	$\theta_1 = \theta_2 = \theta_3 = 0$ (null)	$d = 2$	$d = 6$	$d = 10$	$d = 14$	$d = 18$	$d = 22$
full sample F_T	0.061	0.038	0.642	0.969	0.998	1.000	1.000
F_T^*	0.058	0.087	0.853	0.992	1.000	1.000	1.000
$\rho = 0.75, \pi_0 = 0.8$	$\theta_1 = \theta_2 = \theta_3 = 0$ (null)	$d = 2$	$d = 6$	$d = 10$	$d = 14$	$d = 18$	$d = 22$
full sample F_T	0.063	0.049	0.468	0.907	0.993	0.999	1.000
F_T^*	0.035	0.093	0.696	0.980	0.999	1.000	1.000

Figure 6 plots the size-adjusted power of the F tests for the specification $\theta_2 = -0.5dT^{-1/2}$.⁴¹ In this specification, under the alternative the instrument is relevant throughout the sample. However, because θ_2 has opposite sign to θ_1 and θ_3 the contribution of the second regime tends to offset those of the first and third. Consistent with this, the plots show that F_T^* is substantially more powerful than the full sample F_T . The resulting power gains exceed those obtained when θ_2 shares the same sign as θ_1 and θ_3 .

S.C.2 Finite-sample Size and Power of Weak Identification-Robust Tests

We study the finite-sample rejection frequencies of the proposed tests and compare their performance to the tests analyzed by Andrews, Moreira, and Stock (2006) and Magnusson and Mavroeidis (2014). The analysis is based on the same DGP described in (S.C.2)-(S.C.3)

⁴¹Note that the size of the tests is that reported in Table 5 since the DGPs for the two specifications are the same under the null.

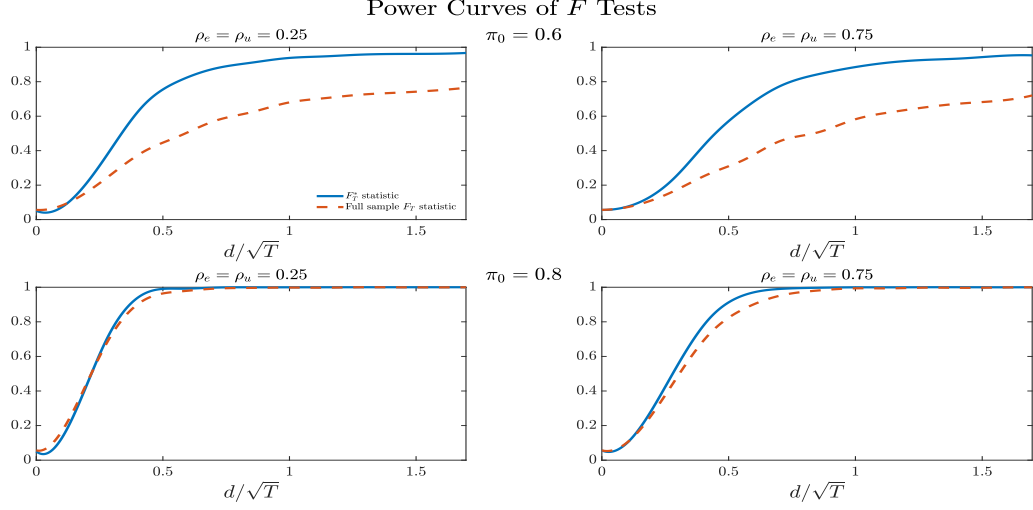


Figure 6: Power curves of F_T^* and full sample F_T for $\theta_2 = -0.5dT^{-1/2}$ under the alternative hypothesis with $T = 200$.

with parameters $\rho \in \{0.25, 0.5, 0.75\}$, $\gamma_1 = \gamma_2 \in \{0, 1\}$, $\theta_1 = \theta_3 = dT^{-1/2}$ with $d \in \{4, 8, 10, 16, 24\}$ and $\theta_2 = 0$. We consider values of $\pi_0 \in \{0.4, 0.6, 0.8\}$ and $T = \{200, 400\}$. Under the null hypothesis we set $\beta = 0$.

We compare the performance of our proposed test statistics $AR_T(\hat{\mathbf{S}}_T)$, $LM_T(\hat{\mathbf{S}}_T)$ and $CLR_T(\hat{\mathbf{S}}_T)$ with their full sample counterparts and with the test statistics Split-S, Split-CLR, qLL-S, ave-S and exp-S analyzed by [Andrews, Moreira, and Stock \(2006\)](#) and with the tests statistics Split-S, Split-CLR, qLL-S, ave-S and exp-S proposed by [Magnusson and Mavroeidis \(2014\)](#). For all tests, we consider heteroskedasticity and autocorrelation-robust versions using the Newey-West estimator with bandwidth equal to the popular rule $1/\lfloor T^{1/3} \rfloor$.⁴² For qLL-S, the tuning parameters c and $\tilde{c}$ are set to 10, following the recommendation in [Magnusson and Mavroeidis \(2014\)](#). The significance level is fixed at 5%, and the number of Monte Carlo replications is set to 10,000 throughout the analysis.

Tables 6-7 report the null rejections frequencies of the various test statistics. In Table 6 the sample size is set to $T = 400$, $\rho = 0.25$ and the errors are assumed to be i.i.d.⁴³ The results show that qLL-S and exp-S tests systematically yield rejection rates below the nominal significance level, while the Split-CLR is systematically oversized. LM_T and $LM_T(\hat{\mathbf{S}}_T)$ show accurate null rejection frequencies for all values of d , whereas CLR_T and $CLR_T(\hat{\mathbf{S}}_T)$ tend to

⁴²We also experiment with data-dependent bandwidths, though results are similar and therefore omitted.

⁴³When the number of instruments is one, the AR tests are not reported, as they are numerically equivalent to the LM tests. For results with serially correlated errors, see the supplement. The conclusions are similar to i.i.d. errors.

produce slightly oversized rejection rates when the instrument is weak (i.e., $d = 4$ and $d = 8$). We note that $LM_T(\hat{\mathbf{S}}_T)$ and $CLR_T(\hat{\mathbf{S}}_T)$ have often more accurate rejection rates than their full sample counterparts. These findings are consistent across values of π_0 . In Table 7 we consider a smaller sample size of $T = 200$ and $\rho \in \{0.25, 0.50, 0.75\}$. The qualitative patterns remain similar. As ρ increases $LM_T(\hat{\mathbf{S}}_T)$ and $CLR_T(\hat{\mathbf{S}}_T)$ become slightly more oversized than their full sample counterparts for $d = 4$, although the opposite occurs for larger d . Thus, the proposed tests demonstrate good size control, even in small samples.

In the supplement, we examine the impact of serial correlation on the null rejection rates. Under strong serial dependence ($\rho_e = \rho_u = 0.75$) all tests exhibit rejection rates that exceed the nominal significance level. Specifically, $LM_T(\hat{\mathbf{S}}_T)$ and $CLR_T(\hat{\mathbf{S}}_T)$ are a bit more oversized than LM_T and CLR_T but similar to qLL-S. Under weak serial dependence $\rho_e = \rho_u = 0.25$, the proposed tests $LM_T(\hat{\mathbf{S}}_T)$ and $CLR_T(\hat{\mathbf{S}}_T)$ are only slightly more oversized than their full sample counterparts, LM_T and CLR_T .

Finally, we turn to the comparison of size-adjusted power, as reported in Figures 7-8. Neither test achieves unit power when the instrument is weak. The results indicate that the proposed tests consistently achieve the highest size-adjusted power across all specifications considered. The power gains are substantial, averaging around approximately 20–30% and reaching 40–50% in the most favorable cases. The latter coincide with the specifications $\theta_2 = -0.5dT^{-1/2}$ and $\theta_2 = -0.5$. Consistent with the theoretical predictions, the gains are more pronounced for smaller values of π_0 , provided that π_0 is not too small. Overall, these finite-sample results support our theoretical relative efficiency results.

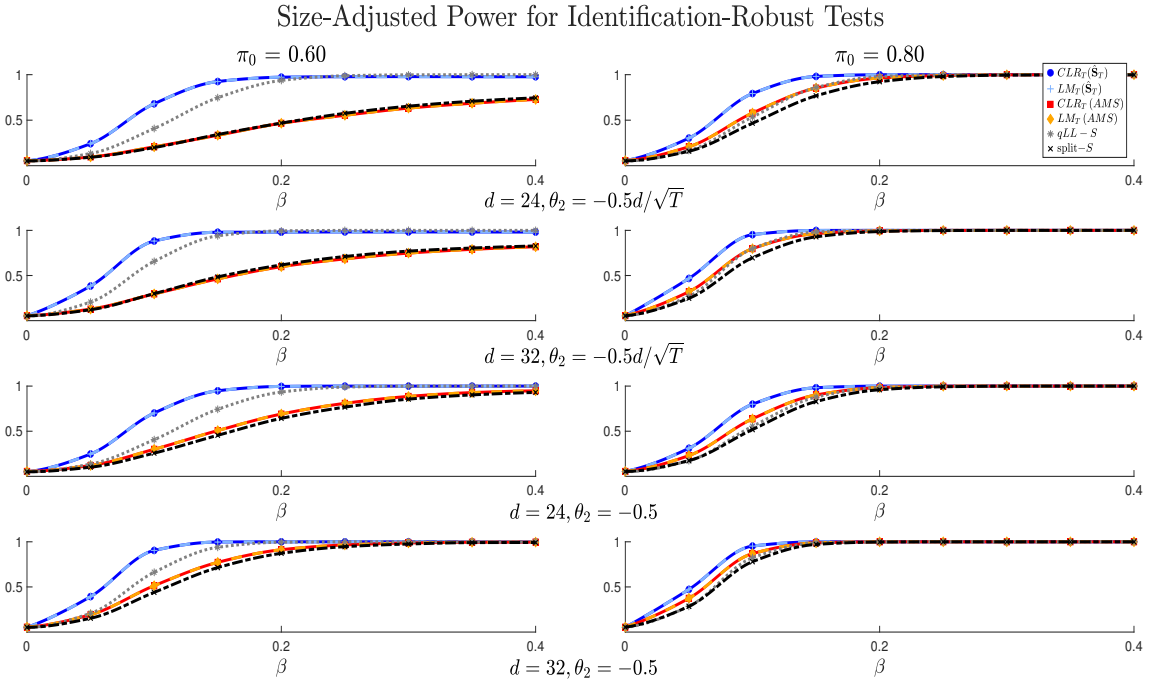
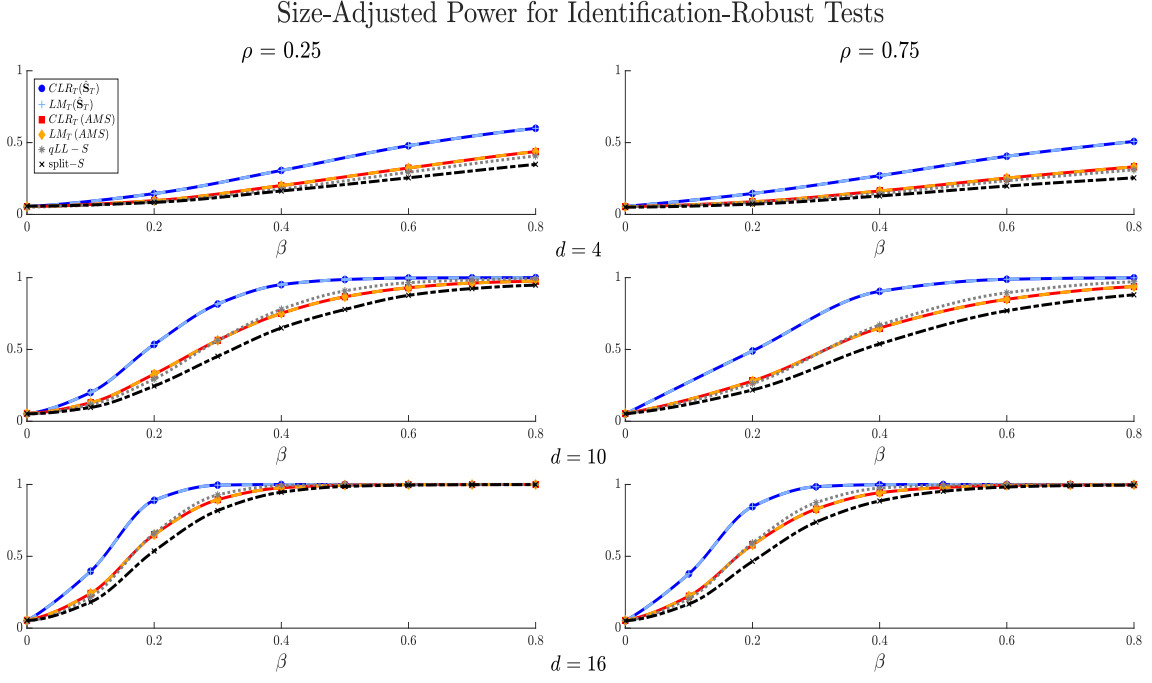
Table 6: Finite-Sample Null Rejection Frequencies of Tests

$T = 400$	$\gamma_1 = \gamma_2 = 0$ and $\rho_e = \rho_u = 0$				$\gamma_1 = \gamma_2 = 1$ and $\rho_e = \rho_u = 0$			
$\rho = 0.25, \pi_0 = 0.6$	$d = 4$	$d = 8$	$d = 12$	$d = 16$	$d = 4$	$d = 8$	$d = 12$	$d = 16$
LM_T	0.057	0.057	0.053	0.056	0.059	0.060	0.058	0.056
CLR_T	0.078	0.065	0.060	0.062	0.083	0.074	0.067	0.065
$LM_T(\hat{\mathbf{S}}_T)$	0.061	0.052	0.049	0.050	0.065	0.055	0.054	0.048
$CLR_T(\hat{\mathbf{S}}_T)$	0.072	0.054	0.051	0.050	0.079	0.074	0.055	0.049
split – S	0.039	0.038	0.036	0.036	0.042	0.042	0.038	0.034
split – CLR	0.115	0.122	0.117	0.115	0.110	0.130	0.122	0.123
qqL – S	0.027	0.031	0.028	0.056	0.031	0.029	0.031	0.026
ave – S	0.044	0.039	0.043	0.038	0.043	0.043	0.042	0.038
exp – S	0.019	0.020	0.017	0.019	0.017	0.022	0.021	0.019
$\rho = 0.25, \pi_0 = 0.4$	$d = 4$	$d = 8$	$d = 12$	$d = 16$	$d = 4$	$d = 8$	$d = 12$	$d = 16$
LM_T	0.058	0.058	0.058	0.058	0.059	0.058	0.056	0.060
CLR_T	0.087	0.079	0.073	0.072	0.083	0.081	0.072	0.073
$LM_T(\hat{\mathbf{S}}_T)$	0.061	0.060	0.065	0.064	0.065	0.065	0.062	0.066
$CLR_T(\hat{\mathbf{S}}_T)$	0.081	0.071	0.073	0.071	0.079	0.076	0.070	0.072
split – S	0.036	0.036	0.038	0.037	0.042	0.039	0.039	0.038
split – CLR	0.104	0.119	0.120	0.121	0.110	0.119	0.119	0.122
qqL – S	0.027	0.028	0.029	0.031	0.028	0.030	0.032	0.032
ave – S	0.040	0.043	0.047	0.043	0.044	0.043	0.037	0.044
exp – S	0.015	0.017	0.017	0.018	0.017	0.020	0.020	0.019
$\rho = 0.25, \pi_0 = 0.8$	$d = 4$	$d = 8$	$d = 12$	$d = 16$	$d = 4$	$d = 8$	$d = 12$	$d = 16$
LM_T	0.055	0.056	0.057	0.058	0.056	0.060	0.058	0.060
CLR_T	0.069	0.061	0.060	0.060	0.076	0.067	0.062	0.062
$LM_T(\hat{\mathbf{S}}_T)$	0.060	0.061	0.053	0.049	0.061	0.059	0.056	0.052
$CLR_T(\hat{\mathbf{S}}_T)$	0.068	0.063	0.054	0.049	0.071	0.062	0.058	0.053
split – S	0.038	0.041	0.047	0.045	0.043	0.047	0.044	0.041
split – CLR	0.117	0.128	0.135	0.135	0.122	0.133	0.133	0.128
qqL – S	0.028	0.032	0.029	0.030	0.025	0.030	0.032	0.031
ave – S	0.045	0.044	0.047	0.040	0.042	0.046	0.044	0.041
exp – S	0.028	0.019	0.017	0.020	0.018	0.022	0.019	0.018

Model M1 and M2. The null hypothesis is $H_0 : \beta = 0$.

Table 7: Finite-Sample Null Rejection Frequencies of Tests

$\gamma_1 = \gamma_2 = 0$ and $\rho_e = \rho_u = 0$									
$T = 200, \pi_0 = 0.6$	$\rho = 0.25$			$\rho = 0.50$			$\rho = 0.75$		
	$d = 4$	$d = 10$	$d = 16$	$d = 4$	$d = 10$	$d = 16$	$d = 4$	$d = 10$	$d = 16$
LM_T	0.061	0.061	0.061	0.062	0.062	0.059	0.062	0.062	0.062
CLR_T	0.084	0.073	0.070	0.085	0.073	0.063	0.080	0.071	0.070
$LM_T(\hat{\mathbf{S}}_T)$	0.068	0.054	0.053	0.075	0.057	0.059	0.082	0.057	0.054
$CLR_T(\hat{\mathbf{S}}_T)$	0.081	0.057	0.054	0.086	0.059	0.059	0.090	0.058	0.054
split – S	0.034	0.035	0.034	0.035	0.034	0.035	0.036	0.035	0.034
split – CLR	0.105	0.113	0.111	0.106	0.115	0.115	0.110	0.115	0.115
qqL – S	0.015	0.017	0.017	0.019	0.019	0.014	0.017	0.018	0.018
ave – S	0.035	0.038	0.035	0.038	0.036	0.039	0.036	0.039	0.042
exp – S	0.012	0.012	0.012	0.012	0.012	0.014	0.012	0.012	0.012
$\gamma_1 = \gamma_2 = 1$ and $\rho_e = \rho_u = 0$									
$T = 200, \pi_0 = 0.6$	$\rho = 0.25$			$\rho = 0.50$			$\rho = 0.75$		
	$d = 4$	$d = 10$	$d = 16$	$d = 4$	$d = 10$	$d = 16$	$d = 4$	$d = 10$	$d = 16$
LM_T	0.066	0.066	0.066	0.064	0.064	0.064	0.069	0.069	0.069
CLR_T	0.089	0.077	0.075	0.089	0.075	0.073	0.089	0.080	0.078
$LM_T(\hat{\mathbf{S}}_T)$	0.070	0.060	0.056	0.082	0.061	0.056	0.095	0.058	0.056
$CLR_T(\hat{\mathbf{S}}_T)$	0.085	0.063	0.057	0.094	0.064	0.056	0.105	0.060	0.057
split – S	0.040	0.040	0.039	0.041	0.040	0.038	0.040	0.037	0.038
split – CLR	0.110	0.123	0.121	0.112	0.112	0.119	0.118	0.119	0.120
qqL – S	0.020	0.021	0.021	0.024	0.019	0.024	0.025	0.023	0.024
ave – S	0.041	0.040	0.039	0.042	0.036	0.037	0.040	0.041	0.038
exp – S	0.014	0.015	0.015	0.015	0.016	0.015	0.014	0.016	0.014



S.D Mathematical Proofs

S.D.1 Proofs of the Results of Sections 2-3

S.D.1.1 Preliminary Lemmas

Lemma S.D.1. *Let Assumption 2.4 hold, $g_t(\cdot)$ be locally absolutely continuous on $\mathbf{D} \subseteq \mathbb{R}$ and $\mathbb{E}[\int_{\mathbf{D}} |\partial g_t(d)/\partial d| dd | \tilde{V}_t] < \infty$. For $t \in \mathbf{S}_{0,T}$, $\tilde{v}$ in the support of $\tilde{V}_t$, and $z, z' \in \mathbf{Z}$, we have*

$$\begin{aligned} & \mathbb{E} \left(g_t(D_t(z')) - g_t(D_t(z)) \mid \tilde{V}_t = \tilde{v} \right) \\ &= \int_{\mathbf{D}} \mathbb{E} \left(\frac{\partial}{\partial d} g_t(d) \mid D_t(z) \leq d \leq D_t(z'), \tilde{V}_t = \tilde{v} \right) \mathbb{P} \left(D_t(z) \leq d \leq D_t(z') \mid \tilde{V}_t = \tilde{v} \right) dd. \end{aligned}$$

Proof of Lemma S.D.1. Suppose that Assumption 2.4 holds with $D_t(z') \geq D_t(z)$. We have

$$\begin{aligned} & \mathbb{E} \left(g_t(D_t(z')) - g_t(D_t(z)) \mid \tilde{V}_t = \tilde{v} \right) \\ &= \mathbb{E} \left(\int_{D_t(z)}^{D_t(z')} \frac{\partial g_t}{\partial d}(d) dd \mid \tilde{V}_t = \tilde{v} \right) \\ &= \mathbb{E} \left(\int_{\mathbf{D}} \frac{\partial g_t}{\partial d}(d) \mathbf{1}_{\{D_t(z) \leq d \leq D_t(z')\}} dd \mid \tilde{V}_t = \tilde{v} \right) \\ &= \int_{\mathbf{D}} \mathbb{E} \left(\frac{\partial g_t}{\partial d}(d) \mathbf{1}_{\{D_t(z) \leq d \leq D_t(z')\}} \mid \tilde{V}_t = \tilde{v} \right) dd \\ &= \int_{\mathbf{D}} \mathbb{E} \left(\frac{\partial g_t}{\partial d}(d) \mid D_t(z) \leq d \leq D_t(z'), \tilde{V}_t = \tilde{v} \right) \mathbb{P} \left(D_t(z) \leq d \leq D_t(z') \mid \tilde{V}_t = \tilde{v} \right) dd, \end{aligned}$$

where the first equality follows from local absolute continuity and the fundamental theorem of calculus and the third equality follows from Fubini's theorem and integrability. $\square$

S.D.1.2 Proof of Proposition 2.1

Consider first the denominator of $\beta_{\pi,t,h}(\tilde{v})$. We appeal to Assumption 2.1 and apply Lemma S.D.1 with $g_t(d) = d$ to obtain for $t \in \mathbf{S}_{0,T}$,

$$\mathbb{E} \left(D_t \mid Z_t = z', \tilde{V}_t = \tilde{v} \right) - \mathbb{E} \left(D_t \mid Z_t = z, \tilde{V}_t = \tilde{v} \right) = \int_{\mathbf{D}} \mathbb{P} \left(D_t(z) \leq d \leq D_t(z') \mid \tilde{V}_t = \tilde{v} \right) dd.$$

Noting that $Y_{t,h}$ is a function of $D_t = d$, we appeal to Assumptions 2.1-2.2 and 2.5(i) and apply Lemma S.D.1 with $g_t(d) = Y_{t,h}^*(d)$ to the numerator of $\beta_{\pi,t,h}(\tilde{v})$ to obtain,

$$\begin{aligned} & \mathbb{E}(Y_{t+h} | Z_t = z', \tilde{V}_t = \tilde{v}) - \mathbb{E}(Y_{t+h} | Z_t = z, \tilde{V}_t = \tilde{v}) \\ &= \int_{\mathbf{D}} \mathbb{E}\left(\frac{\partial Y_{t,h}^*}{\partial d}(d) | D_t(z) \leq d \leq D_t(z'), \tilde{V}_t = \tilde{v}\right) \times \mathbb{P}(D_t(z) \leq d \leq D_t(z') | \tilde{V}_t = \tilde{v}) dd. \end{aligned}$$

Using the derived expressions for the numerator and denominator of $\beta_{\pi,t,h}(\tilde{v})$ yields the expression given in the proposition, which is well-defined by Assumption 2.3. Note that by definition $w_t(d|\tilde{v}) \geq 0$ and $\int_{\mathbf{D}} w_t(d|\tilde{v}) dd = 1$. $\square$

S.D.1.3 Proof of Theorem 2.1

Lemma S.D.2. *Let Assumption 2.6 hold. For each t , $D_t(1) > D_t(0)$ with probability one if and only if $\mathbb{E}(D_t(1)) > \mathbb{E}(D_t(0))$.*

Proof of Lemma S.D.2. First, note that $\mathbb{P}(D_t(1) > D_t(0)) = 1$ immediately implies $\mathbb{E}(D_t(1)) > \mathbb{E}(D_t(0))$. To see the reverse implication, suppose $\mathbb{E}(D_t(1)) > \mathbb{E}(D_t(0))$. Then it must be the case that $\mathbb{P}(D_t(1) > D_t(0)) > 0$. But then Assumption 2.6 immediately implies $\mathbb{P}(D_t(0) < D_t(1)) = 1$. $\square$

Consider first the policy sample. Given $t \in \mathbf{P}$ we have $Z_t = 1$ and $D_t = D_t(1)$. By Lemma S.D.2 $t \in \mathbf{P}$ is a complier if and only if $\mathbb{E}(D_t(1)) > \mathbb{E}(D_t(0))$. By Assumption 2.7(ii) we have $\mathbb{E}(D_t(0)) - \mathbb{E}(D_{t-1}(0)) = 0$. The latter implies $t \in \mathbf{P}$ is a complier if and only if

$$\begin{aligned} \mathbb{E}(D_t(1)) - \mathbb{E}(D_{t-1}(0)) &= \mathbb{E}(D_t(1)) - \mathbb{E}(D_t(0)) \\ &> \mathbb{E}(D_t(0)) - \mathbb{E}(D_t(0)) = 0. \end{aligned}$$

By Assumption 2.7(i), $\overline{D}_{C,t-1,n_0} \xrightarrow{\mathbb{P}} \mathbb{E}(D_{t-1}(0))$ as $n_0 \rightarrow \infty$. By Assumption 2.8(i), $\overline{D}_{P,t,n_1} \xrightarrow{\mathbb{P}} \mathbb{E}(D_t(1))$ as $n_1 \rightarrow \infty$. Thus, $\overline{D}_{P,t,n_1} - \overline{D}_{C,t-1,n_0} \xrightarrow{\mathbb{P}} c$ as $n_0, n_1 \rightarrow \infty$, where $c > 0$ if and only if $t \in \mathbf{P}$ is a complier.

Now consider the control sample. Since $t \in \mathbf{C}$ we have $Z_t = 0$ and $D_t = D_t(0)$. Using Assumption 2.8(ii) and Lemma S.D.2 we have $t \in \mathbf{C}$ is a complier if and only if

$$\mathbb{E}(D_{s^*(t)}(1)) - \mathbb{E}(D_t(0)) = \mathbb{E}(D_t(1)) - \mathbb{E}(D_t(0))$$

$$> \mathbb{E}(D_t(0)) - \mathbb{E}(D_t(0)) = 0.$$

By Assumption 2.7(i), $\overline{D}_{C,t,n_0} \xrightarrow{\mathbb{P}} \mathbb{E}(D_t(0))$ as $n_0 \rightarrow \infty$. By Assumption 2.8(i), $\overline{D}_{P,s^*(t),n_1} \xrightarrow{\mathbb{P}} \mathbb{E}(D_{s^*(t)}(1))$ as $n_1 \rightarrow \infty$. Thus, $\overline{D}_{P,s^*(t),n_1} - \overline{D}_{C,t,n_0} \xrightarrow{\mathbb{P}} \tilde{c}$ as $n_0, n_1 \rightarrow \infty$, where $\tilde{c} > 0$ if and only if $t \in \mathbf{C}$ is a complier. $\square$

S.D.1.4 Proof of Proposition 2.2

Under Assumption 2.4 with $D_t(1) \geq D_t(0)$, non-compliers are characterized by $\mathbb{P}(D_t(1) = D_t(0)) = 1$ using Assumption 2.6 so that any non-complier cannot belong to $\mathbf{S}_{0,T}$. On the other hand, if t is a complier, $\mathbb{E}[D_t(1)] \neq \mathbb{E}[D_t(0)]$ since $\mathbb{P}(D_t(1) > D_t(0)) = 1$ by the definition of a complier so that $t \in \mathbf{S}_{0,T}$. $\square$

S.D.1.5 Proof Proposition 2.3

For $t \in \mathcal{NC}_{\mathbf{C}}^s$, $Z_t = 0$ so that $Y_t = Y_t^*(D_t(0), 0)$ and Assumption 2.9(ii) implies

$$|\mathcal{NC}_{\mathbf{C}}^s|^{-1} \sum_{t \in \mathcal{NC}_{\mathbf{C}}^s} Y_t \xrightarrow{\mathbb{P}} \mathbb{E}[Y_t | t \in \mathcal{NC}_{\mathbf{C}}^s] = \mathbb{E}[Y_t^*(D_t(0), 0) | t \in \mathcal{NC}_{\mathbf{C}}^s] = \mathbb{E}[Y_t^*(D_t, 0) | t \in \mathcal{NC}_{\mathbf{C}}^s],$$

as $|\mathcal{NC}_{\mathbf{C}}^s| \rightarrow \infty$ since t is a non-complier. Similarly, Assumption 2.9(ii) implies

$$|\mathcal{NC}_{\mathbf{P}}^s|^{-1} \sum_{t \in \mathcal{NC}_{\mathbf{P}}^s} Y_t \xrightarrow{\mathbb{P}} \mathbb{E}[Y_t^*(D_t, 1) | t \in \mathcal{NC}_{\mathbf{P}}^s]$$

as $|\mathcal{NC}_{\mathbf{P}}^s| \rightarrow \infty$. But Assumption 2.9(i) implies $\mathbb{E}[Y_t^*(D_t, 0) | t \in \mathcal{NC}_{\mathbf{C}}^s] = \mathbb{E}[Y_t^*(D_t, 0) | t \in \mathcal{NC}_{\mathbf{P}}^s]$, which is in turn equal to $\mathbb{E}[Y_t^*(D_t, 1) | t \in \mathcal{NC}_{\mathbf{P}}^s]$ under Assumption 2.2. $\square$

S.D.2 Proofs of the Results of Sections 4, 6 and S.B

S.D.2.1 Proof of Theorem 4.1

Suppose that $\mathbf{S}_T \in \Xi_{\epsilon, \pi, m, T}$. Note that $\tilde{Z}(S_T) = M_{S_T X} S_T Z$ and

$$\tilde{D}(S_T) = M_{S_T X} S_T D = M_{S_T X} S_T (X\gamma_2 + e) = M_{S_T X} S_T e$$

under $H_{\theta,0}$. Thus,

$$\widetilde{D}(S_T)' \widetilde{Z}(S_T) = (S_T e)' M_{S_T X} M_{S_T X} S_T Z = (S_T e)' M_{S_T X} S_T Z = (S_T e)' \widetilde{Z}(S_T).$$

By Assumptions 4.1-4.2 we have

$$T^{-1/2} \widetilde{Z}(S_T)' \widetilde{D}(S_T) = T^{-1/2} \sum_{t \in \mathbf{S}_T} \widetilde{Z}_t e_t \Rightarrow B(\lambda_{L,1}, \lambda_{R,m}),$$

where $B(\lambda_{L,1}, \lambda_{R,m})$ is q -dimensional Brownian motion with covariance $J(S) = \lim_{T \rightarrow \infty} (\pi T)^{-1} \text{Var}((S_T Z)' M_{S_T X} S_T e)$ and $S = \lim_{T \rightarrow \infty} S_T$. Since $\widehat{J}(S_T) \xrightarrow{\mathbb{P}} J(S)$ uniformly by Assumption 4.3, we have

$$F_T(\mathbf{S}_T) \Rightarrow \frac{B(\lambda_{L,1}, \lambda_{R,m})' J(S)^{-1} B(\lambda_{L,1}, \lambda_{R,m})}{q\pi} = \frac{1}{q\pi} \left\| \sum_{i=1}^m (W_q(\lambda_{R,i}) - W_q(\lambda_{L,i})) \right\|^2.$$

The result then follows from the continuous mapping theorem. $\square$

S.D.2.2 Proof of Theorem 6.1

Lemma S.D.3. *Let Assumptions 6.1-6.3 hold and suppose $\theta = c/T^{1/2}$ for some nonstochastic $c \in \mathbb{R}^q$. We have $\widehat{\Sigma}_v(\mathbf{S}_T) \xrightarrow{\mathbb{P}} \Sigma_v$ uniformly in $\mathbf{S}_T \in \mathcal{S}$.*

Proof of Lemma S.D.3. Recall that C_T is the selection matrix that corresponds to $\mathbf{S}_T$. Using $y = \overline{Z}(C_{0,T}) \theta a'_\beta + X\eta + v$, $P_X \overline{Z}(C_T) = 0$ and $P_{\overline{Z}(C_T)} X = 0$, we have

$$\begin{aligned} \widehat{v}(\mathbf{S}_T) &= y - P_{\overline{Z}(C_T)} y - P_X y = y - P_{\overline{Z}(C_T)} y - X\eta - P_X v \\ &= \overline{Z}(C_{0,T}) \theta a'_\beta + v - P_{\overline{Z}(C_T)} y - P_X v \\ &= \overline{Z}(C_{0,T}) \theta a'_\beta + v - P_{\overline{Z}(C_T)} \overline{Z}(C_{0,T}) \theta a'_\beta - P_{\overline{Z}(C_T)} v - P_X v \\ &= M_{\overline{Z}(C_T)} \overline{Z}(C_{0,T}) \theta a'_\beta + v - P_{\overline{Z}(C_T)} v - P_X v. \end{aligned} \tag{S.D.1}$$

Then, using $\overline{Z}(C_T)' X = 0$, we have

$$\begin{aligned} \widehat{v}(\mathbf{S}_T)' \widehat{v}(\mathbf{S}_T) &= \left(M_{\overline{Z}(C_T)} \overline{Z}(C_{0,T}) \theta a'_\beta + v - P_{\overline{Z}(C_T)} v - P_X v \right)' \\ &\quad \times \left(M_{\overline{Z}(C_T)} \overline{Z}(C_{0,T}) \theta a'_\beta + v - P_{\overline{Z}(C_T)} v - P_X v \right) \end{aligned} \tag{S.D.2}$$

$$\begin{aligned}
&= a_\beta \theta' \bar{Z}(C_{0,T})' M_{\bar{Z}(C_T)} \bar{Z}(C_{0,T}) \theta a'_\beta + a_\beta \theta' \bar{Z}(C_{0,T})' M_{\bar{Z}(C_T)} v \\
&\quad + v' M_{\bar{Z}(C_T)} \bar{Z}(C_{0,T}) \theta a'_\beta + v' v - v' P_{\bar{Z}(C_T)} v - v' P_X v.
\end{aligned}$$

Using Assumptions 6.1, 6.3 and $\theta = cT^{-1/2}$ we have

$$\begin{aligned}
&T^{-1} v' M_{\bar{Z}(C_T)} \bar{Z}(C_{0,T}) \theta a'_\beta \tag{S.D.3} \\
&= T^{-1} v' \bar{Z}(C_{0,T}) \frac{c}{\sqrt{T}} a'_\beta - T^{-1} v' \bar{Z}(C_T) \left(\bar{Z}(C_T)' \bar{Z}(C_T) \right)^{-1} \bar{Z}(C_T)' \bar{Z}(C_{0,T}) \frac{c}{\sqrt{T}} a'_\beta \\
&= T^{-1/2} O_{\mathbb{P}}(1) \frac{c}{\sqrt{T}} a'_\beta - T^{-1/2} \left(T^{-1/2} v' \bar{Z}(C_T) \right) \left(T^{-1} \bar{Z}(C_T)' \bar{Z}(C_T) \right)^{-1} T^{-1} \bar{Z}(C_T)' \bar{Z}(C_{0,T}) \frac{c}{\sqrt{T}} a'_\beta \\
&= O_{\mathbb{P}}(T^{-1}), \quad \text{and}
\end{aligned}$$

$$\begin{aligned}
&T^{-1} a_\beta \theta' \bar{Z}(C_{0,T})' M_{\bar{Z}(C_T)} \bar{Z}(C_{0,T}) \theta a'_\beta \tag{S.D.4} \\
&= T^{-1} a_\beta \theta' \bar{Z}(C_{0,T})' \bar{Z}(C_{0,T}) \theta a'_\beta - T^{-1} a_\beta \theta' \bar{Z}(C_{0,T})' \bar{Z}(C_T) \left(\bar{Z}(C_T)' \bar{Z}(C_T) \right)^{-1} \bar{Z}(C_T)' \bar{Z}(C_{0,T}) \theta a'_\beta \\
&= a_\beta \frac{c'}{\sqrt{T}} T^{-1} \bar{Z}(C_{0,T})' \bar{Z}(C_{0,T}) \frac{c}{\sqrt{T}} a'_\beta \\
&\quad - a_\beta \frac{c'}{\sqrt{T}} T^{-1} \bar{Z}(C_{0,T})' \bar{Z}(C_T) \left(T^{-1} \bar{Z}(C_T)' \bar{Z}(C_T) \right)^{-1} T^{-1} \bar{Z}(C_T)' \bar{Z}(C_{0,T}) \frac{c}{\sqrt{T}} a'_\beta = O_{\mathbb{P}}(T^{-1})
\end{aligned}$$

so that

$$T^{-1} \hat{v}'(\mathbf{S}_T) \hat{v}(\mathbf{S}_T) - \Sigma_v = \left(T^{-1} v' v - \Sigma_v \right) - T^{-1} v' P_{\bar{Z}(C_T)} v - T^{-1} v' P_X v + O_{\mathbb{P}}(T^{-1}). \tag{S.D.5}$$

By Assumption 6.2, the first term on the right-hand side of (S.D.5) converges in probability to zero. The second term satisfies

$$T^{-1} v' P_{\bar{Z}(C_T)} v = T^{-1} \left\| P_{\bar{Z}(C_T)} v \right\|^2 = T^{-1} O_{\mathbb{P}}(1) = o_{\mathbb{P}}(1),$$

under bounded second moments, fixed q , and weak dependence conditions. All convergence results above hold uniformly in $\mathbf{S}_T$. The argument for the third term of (S.D.5) is analogous.

□

Lemma S.D.3 shows that $\hat{\Sigma}_v(\mathbf{S}_T)$ is consistent for Σ_v for all $\mathbf{S}_T \in \mathcal{S}$. The convergence in the lemma occurs uniformly over all true parameters β, c, γ, ϕ and over $\mathbf{S}_T \in \mathcal{S}$. Thus, estimation based on any partition $\mathbf{S}_T \in \mathcal{S}$ leads to residuals $\{\hat{v}(\mathbf{S}_T)\}$ that can be used to

construct a consistent estimate $\hat{\Sigma}_v(\mathbf{S}_T)$ for Σ_v under weak instruments.

For some partitions $\mathbf{S}, \mathbf{S}' \subseteq (0, 1]$, partition $Q(\mathbf{S}, \mathbf{S}')$ conformably with

$$w(\mathbf{S}_T)'w(\mathbf{S}'_T) = \begin{bmatrix} Z(C_T)'Z(C'_T) & Z(C_T)'X \\ X'Z(C'_T) & X'X \end{bmatrix},$$

where $Z(C_T) = C_T Z$, so that

$$Q(\mathbf{S}, \mathbf{S}') = \begin{bmatrix} Q_{11}(\mathbf{S}, \mathbf{S}') & Q_{12}(\mathbf{S}) \\ Q_{21}(\mathbf{S}') & Q_{22} \end{bmatrix},$$

and let $Q(\mathbf{S}) = Q(\mathbf{S}, \mathbf{S})$. In addition, note that

$$\begin{aligned} \Sigma_{N_1}(\mathbf{S}, \mathbf{S}') &= J(\mathbf{S}) B_0 \Psi(\mathbf{S}, \mathbf{S}') B_0' J(\mathbf{S}')', \quad \Sigma_{N_1 N_2}(\mathbf{S}, \mathbf{S}') = J(\mathbf{S}) A_0 \Psi(\mathbf{S}, \mathbf{S}') B_0' J(\mathbf{S}')', \\ \Sigma_{N_2}^*(\mathbf{S}, \mathbf{S}') &= J(\mathbf{S}) A_0 \Psi(\mathbf{S}, \mathbf{S}') A_0' J(\mathbf{S}')' \end{aligned} \quad (\text{S.D.6})$$

for $J(\mathbf{S}) = [I_q : -Q_{12}(\mathbf{S}) Q_{22}^{-1}]$, $B_0 = (b_0' \otimes I_{q+p})$ and $A_0 = (\Sigma_v^{-1} a_{0,\beta})' \otimes I_{q+p}$.⁴⁴

Finally, let $N_{1,\infty}(\cdot)$ and $N_{2,\infty}(\cdot)$ be independent q -dimensional Gaussian processes indexed by $\mathbf{S} \subseteq (0, 1]$ with

$$\begin{aligned} N_{1,\infty}(\mathbf{S}) &\sim \mathcal{N}(\Sigma_{N_1}^{-1/2}(\mathbf{S}) \Sigma_{\bar{Z}}(\mathbf{S}, \mathbf{S}_0) c a_\beta' b_0, I_q), \\ N_{2,\infty}(\mathbf{S}) &\sim \mathcal{N}(\Sigma_{N_2}^{-1/2}(\mathbf{S}) (\Sigma_{\bar{Z}}(\mathbf{S}, \mathbf{S}_0) c a_\beta' \Sigma_v^{-1} a_{0,\beta} - \Sigma_{N_1 N_2}(\mathbf{S}) \Sigma_{N_1}^{-1}(\mathbf{S}) \Sigma_{\bar{Z}}(\mathbf{S}, \mathbf{S}_0) c a_\beta' b_0), I_q), \end{aligned} \quad (\text{S.D.7})$$

where $\Sigma_{\bar{Z}}(\mathbf{S}, \mathbf{S}') = Q_{11}(\mathbf{S}, \mathbf{S}') - Q_{12}(\mathbf{S}) Q_{22}^{-1} Q_{21}(\mathbf{S}')$.

Lemma S.D.4. *Let Assumptions 6.1-6.5 hold and suppose $\theta = c/T^{1/2}$ for some nonstochastic $c \in \mathbb{R}^q$. Then, for $\mathbf{S}_T \in \mathcal{S}$ and $\mathbf{S} = \lim_{T \rightarrow \infty} T^{-1} \mathbf{S}_T$, we have $(N_{1,T}(\mathbf{S}_T), N_{2,T}(\mathbf{S}_T)) \Rightarrow (N_{1,\infty}(\mathbf{S}), N_{2,\infty}(\mathbf{S}))$.*

Proof of Lemma S.D.4. By Assumption 6.1,

$$T^{-1} \bar{Z}(C_T)' \bar{Z}(C'_T) = T^{-1} Z(C_T)' Z(C'_T) - T^{-1} Z(C_T)' P_X Z(C'_T) \xrightarrow{\mathbb{P}} \Sigma_{\bar{Z}}(\mathbf{S}, \mathbf{S}'), \quad (\text{S.D.8})$$

uniformly in $\mathbf{S}_T, \mathbf{S}'_T \in \mathcal{S}$. By Assumptions 6.1 and 6.3, we have uniformly in $\mathbf{S}_T \in \mathcal{S}$,

$$T^{-1/2} \bar{Z}(C_T)' v b_0 = T^{-1/2} (Z(C_T) - P_X Z(C_T))' v b_0 \quad (\text{S.D.9})$$

⁴⁴See (S.D.9) and (S.D.12) for details.

$$\begin{aligned}
&= T^{-1/2} \left(Z(C_T) - X Q_{22}^{-1} Q_{21}(\mathbf{S}_T) \right)' v b_0 + o_{\mathbb{P}}(1) \\
&= \left[I_q : -Q_{12}(\mathbf{S}_T) Q_{22}^{-1} \right] T^{-1/2} w(\mathbf{S}_T)' v b_0 + o_{\mathbb{P}}(1) \\
&= \left[I_q : -Q_{12}(\mathbf{S}_T) Q_{22}^{-1} \right] (b'_0 \otimes I_{q+p}) T^{-1/2} \text{vec}(w(\mathbf{S}_T)' v), \\
&\Rightarrow J(\mathbf{S}) B_0 \mathcal{G}(\mathbf{S}).
\end{aligned}$$

Using (S.D.6), (S.D.8), (S.D.9) and Assumption 6.4,

$$\begin{aligned}
N_{1,T}(\mathbf{S}_T) &= \hat{\Sigma}_{N_1}^{-1/2}(\mathbf{S}_T) T^{-1/2} \bar{Z}(C_T)' \left(T^{-1/2} \bar{Z}(C_{0,T}) c a'_\beta + v \right) b_0 \\
&\Rightarrow \Sigma_{N_1}^{-1/2}(\mathbf{S}) \Sigma_{\bar{Z}}(\mathbf{S}, \mathbf{S}_0) c a'_\beta b_0 + \Sigma_{N_1}^{-1/2}(\mathbf{S}) J(\mathbf{S}) B_0 \mathcal{G}(\mathbf{S}) \sim N_{1,\infty}(\mathbf{S})
\end{aligned} \tag{S.D.10}$$

since $J(\mathbf{S}) B_0 \mathcal{G}(\mathbf{S}) \sim \mathcal{N}(0, \Sigma_{N_1}(\mathbf{S}))$. Similarly, using Lemma S.D.3, Assumptions 6.1, 6.3 and 6.4, (S.D.6) and (S.D.10), we have

$$\begin{aligned}
N_{2,T}(\mathbf{S}_T) &= \hat{\Sigma}_{N_2}^{-1/2}(\mathbf{S}_T) \\
&\quad \times \left(T^{-1/2} \bar{Z}(C_T)' \left(T^{-1/2} \bar{Z}(C_{0,T}) c a'_\beta + v \right) \hat{\Sigma}_v^{-1}(\mathbf{S}_T) a_{0,\beta} - \hat{\Sigma}_{N_1 N_2}(\mathbf{S}_T) \hat{\Sigma}_{N_1}^{-1/2}(\mathbf{S}_T) N_{1,T}(\mathbf{S}_T) \right) \\
&= \hat{\Sigma}_{N_2}^{-1/2}(\mathbf{S}_T) \\
&\quad \times \left(\left(T^{-1} \bar{Z}(C_T)' \bar{Z}(C_{0,T}) c a'_\beta + T^{-1/2} \bar{Z}(C_T)' v \right) \Sigma_v^{-1} a_{0,\beta} - \hat{\Sigma}_{N_1 N_2}(\mathbf{S}_T) \hat{\Sigma}_{N_1}^{-1/2}(\mathbf{S}_T) N_{1,T}(\mathbf{S}_T) \right) + o_{\mathbb{P}}(1) \\
&\Rightarrow \Sigma_{N_2}^{-1/2}(\mathbf{S}) \left(\Sigma_{\bar{Z}}(\mathbf{S}, \mathbf{S}_0) c a'_\beta \Sigma_v^{-1} a_{0,\beta} \right) + \Sigma_{N_2}^{-1/2}(\mathbf{S}) J(\mathbf{S}) A_0 \mathcal{G}(\mathbf{S}) \\
&\quad - \Sigma_{N_2}^{-1/2}(\mathbf{S}) \left(\Sigma_{N_1 N_2}(\mathbf{S}) \Sigma_{N_1}^{-1}(\mathbf{S}) \Sigma_{\bar{Z}}(\mathbf{S}, \mathbf{S}_0) c a'_\beta b_0 \right) \\
&\quad - \Sigma_{N_2}^{-1/2}(\mathbf{S}) \left(\Sigma_{N_1 N_2}(\mathbf{S}) \Sigma_{N_1}^{-1}(\mathbf{S}) J(\mathbf{S}) B_0 \mathcal{G}(\mathbf{S}) \right) \sim N_{2,\infty}(\mathbf{S}),
\end{aligned} \tag{S.D.11}$$

where $T^{-1} \bar{Z}(C_T)' \bar{Z}(C_{0,T}) \xrightarrow{\mathbb{P}} \Sigma_{\bar{Z}}(\mathbf{S}, \mathbf{S}_0)$ uniformly over $\mathbf{S}_T \in \mathcal{S}$ by (S.D.8) and

$$T^{-1/2} \bar{Z}(C_T)' v \Sigma_v^{-1} a_{0,\beta} \Rightarrow \left[I_q : -Q_{12}(\mathbf{S}) Q_{22}^{-1} \right] \left(a'_{0,\beta} \Sigma_v^{-1} \otimes I_{q+p} \right) \mathcal{G}(\mathbf{S}) \tag{S.D.12}$$

in analogy with the arguments that show (S.D.9). The distributional equivalence of the limit in (S.D.11) can be seen after noting

$$\begin{aligned}
&\text{Var} \left(J(\mathbf{S}) A_0 \mathcal{G}(\mathbf{S}) - \Sigma_{N_1 N_2}(\mathbf{S}) \Sigma_{N_1}^{-1}(\mathbf{S}) J(\mathbf{S}) B_0 \mathcal{G}(\mathbf{S}) \right) \\
&= J(\mathbf{S}) A_0 \Psi(\mathbf{S}, \mathbf{S}) A'_0 J(\mathbf{S})' - J(\mathbf{S}) A_0 \Psi(\mathbf{S}, \mathbf{S}) B'_0 J(\mathbf{S})' \Sigma_{N_1}^{-1}(\mathbf{S}) \Sigma_{N_1 N_2}(\mathbf{S})' \\
&\quad - \Sigma_{N_1 N_2}(\mathbf{S}) \Sigma_{N_1}^{-1}(\mathbf{S}) J(\mathbf{S}) B_0 \Psi(\mathbf{S}, \mathbf{S}) A'_0 J(\mathbf{S})'
\end{aligned} \tag{S.D.13}$$

$$\begin{aligned}
& + \Sigma_{N_1 N_2}(\mathbf{S}) \Sigma_{N_1}^{-1}(\mathbf{S}) J(\mathbf{S}) B_0 \Psi(\mathbf{S}, \mathbf{S}) B_0' J(\mathbf{S})' \Sigma_{N_1}^{-1}(\mathbf{S}) \Sigma_{N_1 N_2}(\mathbf{S})' \quad (\text{S.D.14}) \\
& = \Sigma_{N_2}^*(\mathbf{S}) - \Sigma_{N_1 N_2}(\mathbf{S}) \Sigma_{N_1}^{-1}(\mathbf{S}) \Sigma_{N_1 N_2}(\mathbf{S})' - \Sigma_{N_1 N_2}(\mathbf{S}) \Sigma_{N_1}^{-1}(\mathbf{S}) \Sigma_{N_1 N_2}(\mathbf{S})' \\
& \quad + \Sigma_{N_1 N_2}(\mathbf{S}) \Sigma_{N_1}^{-1}(\mathbf{S}) \Sigma_{N_1}(\mathbf{S}) \Sigma_{N_1}^{-1}(\mathbf{S}) \Sigma_{N_1 N_2}(\mathbf{S})' \\
& = \Sigma_{N_2}^*(\mathbf{S}) - \Sigma_{N_1 N_2}(\mathbf{S}) \Sigma_{N_1}^{-1}(\mathbf{S}) \Sigma_{N_1 N_2}(\mathbf{S})' = \Sigma_{N_2}(\mathbf{S}).
\end{aligned}$$

The weak convergence in (S.D.9) occurs jointly with that in (S.D.11) since $N_{1,T}(\cdot)$ and $N_{2,T}(\cdot)$ are functions of the same data. And finally, they are asymptotically independent since they are asymptotically Gaussian and

$$\begin{aligned}
& \text{Cov} \left(\Sigma_{N_1}^{-1/2}(\mathbf{S}) J(\mathbf{S}) B_0 \mathcal{G}(\mathbf{S}), \Sigma_{N_2}^{-1/2}(\mathbf{S}') \left(J(\mathbf{S}') A_0 - \Sigma_{N_1 N_2}(\mathbf{S}') \Sigma_{N_1}^{-1}(\mathbf{S}') J(\mathbf{S}') B_0 \right) \mathcal{G}(\mathbf{S}') \right) \\
& = \Sigma_{N_1}^{-1/2}(\mathbf{S}) J(\mathbf{S}) B_0 \Psi(\mathbf{S}, \mathbf{S}') \left(A_0' J(\mathbf{S}')' - B_0' J(\mathbf{S}')' \Sigma_{N_1}^{-1}(\mathbf{S}') \Sigma_{N_1 N_2}(\mathbf{S}')' \right) \Sigma_{N_2}^{-1/2}(\mathbf{S}') \\
& = \Sigma_{N_1}^{-1/2}(\mathbf{S}) \left(\Sigma_{N_1 N_2}(\mathbf{S}', \mathbf{S})' - \Sigma_{N_1}(\mathbf{S}, \mathbf{S}') \Sigma_{N_1}^{-1}(\mathbf{S}') \Sigma_{N_1 N_2}(\mathbf{S}')' \right) \Sigma_{N_2}^{-1/2}(\mathbf{S}') \\
& = \pi(\mathbf{S} \cap \mathbf{S}') \Sigma_{N_1}^{-1/2}(\mathbf{S}) \left(\Sigma_{N_1 N_2}' - \Sigma_{N_1} \Sigma_{N_1}^{-1} \Sigma_{N_1 N_2}' \right) \Sigma_{N_2}^{-1/2}(\mathbf{S}') = 0
\end{aligned}$$

for any $\mathbf{S}, \mathbf{S}' \subseteq (0, 1]$ by Assumption 6.5. $\square$

Inspection of the proof shows that the results hold uniformly over compact sets of true values of β and c (including the zero vector) and over arbitrary sets of true γ and ϕ values.

Proof of Theorem 6.1. Let

$$\begin{aligned}
M_\infty(\mathbf{S}) &= [N_{1,\infty}(\mathbf{S}) : N_{2,\infty}(\mathbf{S})]' [N_{1,\infty}(\mathbf{S}) : N_{2,\infty}(\mathbf{S})], \quad (\text{S.D.15}) \\
\overline{M}_{1,\infty}(\mathbf{S}) &= \left(N_{1,\infty}(\mathbf{S})' N_{1,\infty}(\mathbf{S}), N_{1,\infty}(\mathbf{S})' N_{2,\infty}(\mathbf{S}) \right)', \\
M_{2,\infty}(\mathbf{S}) &= N_{2,\infty}(\mathbf{S})' N_{2,\infty}(\mathbf{S}), \quad M_{1,2,\infty}(\mathbf{S}) = N_{1,\infty}(\mathbf{S})' N_{2,\infty}(\mathbf{S}), \quad M_{1,\infty}(\mathbf{S}) = N_{1,\infty}(\mathbf{S})' N_{1,\infty}(\mathbf{S}).
\end{aligned}$$

By Lemma S.D.4 $N_{1,T}(\cdot)$ and $N_{2,T}(\cdot)$ are asymptotically independent. $\hat{\mathbf{S}}_T$ depends on $N_{2,T}(\cdot)$ only. Thus, $N_{1,\infty}(\cdot)$ and $\hat{\mathbf{S}}_T$ are asymptotically independent. Using Lemma S.D.4 we yield that under H_0 $N_{1,\infty}(\hat{\mathbf{S}}_T)$ is Gaussian with zero mean and variance I_q . Thus, $M_{1,\infty}(\mathbf{S}_T) \sim M_{1,\infty}$ for all $\mathbf{S}_T \in \mathcal{S}$. Part (i) follows by using the continuous mapping theorem.

For part (ii), note that conditional on $N_{2,\infty}(\cdot)$ Lemma S.D.4 implies that $M_{1,2,\infty}(\hat{\mathbf{S}}_T)$ is Gaussian with zero mean and variance $N_{2,\infty}(\hat{\mathbf{S}}_T)' N_{2,\infty}(\hat{\mathbf{S}}_T)$. The result then follows from the continuous mapping theorem.

We now move to part (iii). By Lemma S.D.4 $N_{1,T}(\cdot)$ and $N_{2,T}(\cdot)$ are asymptotically independent and $\hat{\mathbf{S}}_T$ depends on $N_{2,T}(\cdot)$ only. $\hat{\mathbf{S}}_T$ is inconsistent and converges in distribution

to a random variable $\mathbf{S}_\infty$. Thus, $LR_T(\hat{\mathbf{S}}_T)$ has asymptotically the same distribution as

$$\begin{aligned} CLR_\infty(M_{1,\infty}(\mathbf{S}_\infty), M_{2,\infty}(\mathbf{S}_\infty), \Sigma_v, \beta_0) \\ = \frac{1}{2} \left(M_{1,\infty}(\mathbf{S}_\infty) - M_{2,\infty}(\mathbf{S}_\infty) + \sqrt{(M_{1,\infty}(\mathbf{S}_\infty) - M_{2,\infty}(\mathbf{S}_\infty))^2 + 4M_{1,2,\infty}^2(\mathbf{S}_\infty)} \right). \end{aligned}$$

Conditional on $N_{2,T}(\cdot)$ $\hat{\mathbf{S}}_T$ is fixed. Define $\kappa_{CLR,\alpha}(M_{1,\infty}(\mathbf{S}_\infty), m_2, \Sigma_v, \beta_0)$ to be the $1 - \alpha$ quantile of the null distribution of $CLR_\infty(M_{1,\infty}(\mathbf{S}_\infty), m_2, \Sigma_v, \beta_0)$. Since under H_0 $\kappa_{CLR,\alpha}(\hat{\mathbf{S}}_T)$ has asymptotically the same distribution as $\kappa_{CLR,\alpha}(M_{1,\infty}(\mathbf{S}_\infty), M_{2,\infty}(\mathbf{S}_\infty), \Sigma_v, \beta_0)$, we have that the distribution of $LR_T(\hat{\mathbf{S}}_T) - \kappa_{CLR,\alpha}(\hat{\mathbf{S}}_T)$ is asymptotically the same distribution as

$$CLR_\infty(M_{1,\infty}(\mathbf{S}_\infty), M_{2,\infty}(\mathbf{S}_\infty), \Sigma_v, \beta_0) - \kappa_{CLR,\alpha}(M_{1,\infty}(\mathbf{S}_\infty), M_{2,\infty}(\mathbf{S}_\infty), \Sigma_v, \beta_0).$$

Conditional on $N_{2,\infty}(\cdot)$, $N_{1,\infty}(\hat{\mathbf{S}}_T) \sim \mathcal{N}(0, I_q)$ under H_0 . This implies that the conditional null distribution of CLR_∞ given $N_{2,\infty}(\cdot)$ does not depend on θ or c . Thus, the test that rejects H_0 when $LR_T(\hat{\mathbf{S}}_T) - \kappa_{CLR,\alpha}(\hat{\mathbf{S}}_T) > 0$ is similar at significance level α . $\square$

S.D.2.3 Proof of Proposition S.B.1

The distribution of y is multivariate normal with

$$\mathbb{E}(y) = \bar{Z}(C_{0,T})\theta a'_\beta + X\eta, \quad (\text{S.D.16})$$

independence across rows, and covariance matrix Σ_v for each row. Thus, the density of y is

$$\begin{aligned} (2\pi)^{-T/2} |\Sigma_v|^{-T/2} \exp \left(-\frac{1}{2} \sum_{t=1}^T \left(y_t - a_\beta \theta' \bar{Z}_t(C_{0,T}) - \eta' X_t \right)' \Sigma_v^{-1} \left(y_t - a_\beta \theta' \bar{Z}_t(C_{0,T}) - \eta' X_t \right) \right) \\ (\text{S.D.17}) \\ = (2\pi)^{-T/2} |\Sigma_v|^{-T/2} \exp \left(-\frac{1}{2} \left[\sum_{t=1}^T y_t' \Sigma_v^{-1} y_t - 2\theta' \left(\sum_{t=1}^T \bar{Z}_t(C_{0,T}) y_t' \right) \Sigma_v^{-1} a_\beta \right. \right. \\ \left. \left. - 2\text{Tr} \left(\left(\sum_{t=1}^T X_t y_t' \right) \Sigma_v^{-1} \eta' \right) + \sum_{t=1}^T \left(a_\beta \theta' \bar{Z}_t(C_{0,T}) - \eta' X_t \right)' \Sigma_v^{-1} \left(a_\beta \theta' \bar{Z}_t(C_{0,T}) - \eta' X_t \right) \right] \right). \end{aligned}$$

By the Fisher–Neyman factorization theorem $N(y)$ is a sufficient statistic for $\psi = (\beta, \theta', \gamma', \phi)'$ if and only if the density $\mathcal{L}(y; \beta, \theta, \gamma, \phi, \mathbf{S}_{0,T})$ can be factorized as $\mathcal{L}(y; \beta, \theta, \gamma, \phi, \mathbf{S}_{0,T}) = f_\psi(N(y)) h(y)$ for nonnegative functions $f_\psi(\cdot)$ and $h(\cdot)$. Note that $\bar{Z}(C_{0,T}) = M_X C_{0,T} Z \neq$

$M_X Z$ if $\pi_0 < 1$ and so $\mathcal{L}(y; \beta, \theta, \gamma, \phi, \mathbf{S}_{0,T})$ cannot be factorized as above for $N(y) = [y' M_X Z : y' X]$. Thus, $Z' M_X y$ and $X' y$ are not sufficient statistics for ψ if $\pi_0 < 1$. Therefore when $\pi_0 < 1$, $Z' M_X y$ cannot be sufficient for $(\beta, \theta)'$ if (i) $Z' M_X y$ and $X' y$ are independent, (ii) the distribution of $X' y$ does not depend on $(\beta, \theta)'$ and (iii) the distribution of $Z' M_X y$ does not depend on $(\gamma', \phi)'$.

To complete the proof, we verify that (i)–(iii) above hold. For (i), note that $Z' M_X y$ and $X' y$ are (jointly) multivariate normal random matrices and for any $b_1, b_2 \in \mathbb{R}^2$, we have

$$\text{Cov}(Z' M_X y b_1, X' y b_2) = Z' M_X \text{Cov}(y b_1, y b_2) X = Z' M_X b_1' \Sigma_v b_2 I_T X = b_1' \Sigma_v b_2 Z' M_X X = 0,$$

where the second equality uses the independence of the rows of y . Lemma 1(c) of [Andrews, Moreira, and Stock \(2006\)](#) implies (ii) and for (iii), note that the normality of $Z' M_X y$ has

$$\begin{aligned} \mathbb{E}(Z' M_X y) &= Z' M_X \mathbb{E}(y) = Z' M_X (\bar{Z}(C_{0,T}) \theta a'_\beta + X \eta) = Z' \bar{Z}(C_{0,T}) \theta a'_\beta, \\ \text{Var}(Z' M_X y b) &= Z' M_X \text{Var}(y b) M_X Z = Z' M_X b' \Sigma_v b I_T M_X Z, \end{aligned}$$

for any $b \in \mathbb{R}^2$, which do not depend on $(\gamma', \phi)'$. $\square$

S.D.2.4 Proof of Proposition S.B.2

The density of y is given by $\mathcal{L}(y; \beta, \theta, \gamma, \phi, \mathbf{S}_T)$, where $\mathcal{L}(\cdot)$ is defined in (S.D.17) and $\mathbf{S}_T$ is an unknown parameter. Using the same logic as in the proof of Proposition S.B.1, by the Fisher–Neyman factorization theorem $\{\bar{Z}(C_T)' y\}_{\mathbf{S}_T \in \mathcal{S}}$ is a sufficient statistic for $(\beta, \theta)'$ if (i) $\{\bar{Z}(C_T)' y\}_{\mathbf{S}_T \in \mathcal{S}}$ and $X' y$ are independent, (ii) the distribution of $X' y$ does not depend on $(\beta, \theta)'$ and (iii) the distribution of $\{\bar{Z}(C_T)' y\}_{\mathbf{S}_T \in \mathcal{S}}$ does not depend on $(\gamma', \phi)'$. For (i), note that $\bar{Z}(C_T)' y$ and $X' y$ are (jointly) multivariate normal random matrices and for any $b_1, b_2 \in \mathbb{R}^2$, we have

$$\text{Cov}(\bar{Z}(C_T)' y b_1, X' y b_2) = Z' C_T' M_X \text{Cov}(y b_1, y b_2) X = Z' C_T' M_X b_1' \Sigma_v b_2 I_T X = b_1' \Sigma_v b_2 Z' C_T' M_X X = 0,$$

where the second equality uses the independence of the rows of y . Lemma 1(c) of [Andrews, Moreira, and Stock \(2006\)](#) implies (ii) and for (iii), note that

$$\begin{aligned} \mathbb{E}(Z' C_T' M_X y) &= Z' C_T' M_X \mathbb{E}(y) = Z' C_T' M_X (\bar{Z}(C_{0,T}) \theta a' + X \eta) = Z' C_T' \bar{Z}(C_{0,T}) \theta a'_\beta, \\ \text{Var}(Z' C_T' M_X y b) &= Z' C_T' M_X \text{Var}(y b) M_X C_T Z = Z' C_T' M_X b' \Sigma_v b I_T M_X C_T Z, \end{aligned}$$

for any $b \in \mathbb{R}^2$, which does not depend on $(\gamma', \phi')'$. $\square$

S.D.2.5 Proof of Proposition S.B.3

Under the conditions of the proposition, the log-likelihood of y is given (up to a constant) by

$$\begin{aligned} \ell(\beta, \theta, \gamma, \phi, \mathbf{S}_T) & \tag{S.D.18} \\ &= -\frac{1}{2} \text{Tr} \left(\Sigma_v^{-1} \left((y - \bar{Z}(C_T) \theta a'_\beta - X\eta)' (y - \bar{Z}(C_T) \theta a'_\beta - X\eta) \right) \right) \\ &= \text{Tr} \left(\Sigma_v^{-1} a_\beta \theta' \bar{Z}(C_T)' y \right) - \frac{1}{2} \text{Tr} \left(\Sigma_v^{-1} a_\beta \theta' \bar{Z}(C_T)' \bar{Z}(C_T) \theta a'_\beta \right) - \frac{1}{2} \text{Tr} \left(\Sigma_v^{-1} (y - X\eta)' (y - X\eta) \right). \end{aligned}$$

Maximizing this log-likelihood with respect to θ under H_0 yields

$$\tilde{\theta}(C_T) = \left(\bar{Z}(C_T)' \bar{Z}(C_T) \right)^{-1} \bar{Z}(C_T)' y \Sigma_v^{-1} a_{0,\beta} (a'_{0,\beta} \Sigma_v^{-1} a_{0,\beta})^{-1},$$

so that the concentrated likelihood function under H_0 is

$$\begin{aligned} \ell(y; \beta_0, \tilde{\theta}(C_T), \gamma, \phi, \mathbf{S}_T) & \tag{S.D.19} \\ &= (a'_{0,\beta} \Sigma_v^{-1} a_{0,\beta})^{-1} \text{Tr} \left(\Sigma_v^{-1} a_{0,\beta} a'_{0,\beta} \Sigma_v^{-1} y' \bar{Z}(C_T) \left(\bar{Z}(C_T)' \bar{Z}(C_T) \right)^{-1} \bar{Z}(C_T)' y \right) \\ &\quad - \frac{1}{2} (a'_{0,\beta} \Sigma_v^{-1} a_{0,\beta})^{-2} \text{Tr} \left(\Sigma_v^{-1} a_{0,\beta} a'_{0,\beta} \Sigma_v^{-1} y' \bar{Z}(C_T) \left(\bar{Z}(C_T)' \bar{Z}(C_T) \right)^{-1} \bar{Z}(C_T)' y \Sigma_v^{-1} a_{0,\beta} a'_{0,\beta} \right) \\ &\quad - \frac{1}{2} \text{Tr} \left(\Sigma_v^{-1} (y - X\eta)' (y - X\eta) \right) \\ &= \frac{1}{2} (a'_{0,\beta} \Sigma_v^{-1} a_{0,\beta})^{-1} a'_{0,\beta} \Sigma_v^{-1} y' \bar{Z}(C_T) \left(\bar{Z}(C_T)' \bar{Z}(C_T) \right)^{-1} \bar{Z}(C_T)' y \Sigma_v^{-1} a_{0,\beta} \\ &\quad - \frac{1}{2} \text{Tr} \left(\Sigma_v^{-1} (y - X\eta)' (y - X\eta) \right). \end{aligned}$$

Maximizing (S.D.19) with respect to $\mathbf{S}_T$ is equivalent to maximizing

$$(a'_{0,\beta} \Sigma_v^{-1} a_{0,\beta})^{-1} a'_{0,\beta} \Sigma_v^{-1} y' \bar{Z}(C_T) \left(\bar{Z}(C_T)' \bar{Z}(C_T) \right)^{-1} \bar{Z}(C_T)' y \Sigma_v^{-1} a_{0,\beta}, \tag{S.D.20}$$

making the MLE of $\mathbf{S}_T$ equal to the maximizer of (S.D.20) over $\mathbf{S}_T \in \mathcal{S}$.

To complete the proof, we show that maximizing $M_{2,T}(\mathbf{S}_T) = N_{2,T}(\mathbf{S}_T)' N_{2,T}(\mathbf{S}_T)$ over $\mathbf{S}_T \in \mathcal{S}$ is asymptotically equivalent to maximizing (S.D.20) over $\mathbf{S}_T \in \mathcal{S}$ under the conditions

of the proposition. To see this, first note that

$$\begin{aligned}
 \Sigma_{N_1 N_2}(\mathbf{S}) &= \lim_{T \rightarrow \infty} T^{-1} \sum_{t=1}^T \mathbb{E} \left[v_t b_0 \bar{Z}_t(C_T)' \bar{Z}_t(C_T) a'_{0,\beta} \Sigma_v^{-1} v_t' \right] \\
 &= \lim_{T \rightarrow \infty} T^{-1} \sum_{t=1}^T \bar{Z}_t(C_T)' \bar{Z}_t(C_T) a'_{0,\beta} \Sigma_v^{-1} \mathbb{E} [v_t' v_t] b_0 \\
 &= \lim_{T \rightarrow \infty} T^{-1} \sum_{t=1}^T \bar{Z}_t(C_T)' \bar{Z}_t(C_T) a'_{0,\beta} b_0 = 0
 \end{aligned} \tag{S.D.21}$$

for $\mathbf{S} = \lim_{T \rightarrow \infty} T^{-1} \mathbf{S}_T$, where the first equality follows from the i.i.d. assumption, the second from the assumption of fixed regressors, the third from $\mathbb{E} [v_t' v_t] = \Sigma_v$ and the fourth from $a'_{0,\beta} b_0 = 0$. Thus, applying Lemma S.D.3, $N_{2,T}(\mathbf{S}_T)$ is asymptotically equivalent to $\Sigma_{N_2}^{-1/2}(\mathbf{S}_T) T^{-1/2} \bar{Z}(C_T)' y \Sigma_v^{-1} a_{0,\beta}$ under Assumption 6.4, implying that $M_{2,T}(\mathbf{S}_T)$ is asymptotically equivalent to $T^{-1} a'_{0,\beta} \Sigma_v^{-1} y' \bar{Z}(C_T) \Sigma_{N_2}^{-1}(\mathbf{S}_T) \bar{Z}(C_T)' y \Sigma_v^{-1} a_{0,\beta}$. The result then follows from

$$\begin{aligned}
 \Sigma_{N_2}(\mathbf{S}) &= \lim_{T \rightarrow \infty} T^{-1} \sum_{t=1}^T \mathbb{E} \left[v_t \Sigma_v^{-1} a_{0,\beta} \bar{Z}_t(C_T)' \bar{Z}_t(C_T) a'_{0,\beta} \Sigma_v^{-1} v_t' \right] \\
 &= \lim_{T \rightarrow \infty} T^{-1} \sum_{t=1}^T \bar{Z}_t(C_T)' \bar{Z}_t(C_T) a'_{0,\beta} \Sigma_v^{-1} \mathbb{E} [v_t' v_t] \Sigma_v^{-1} a_{0,\beta} = \lim_{T \rightarrow \infty} T^{-1} \bar{Z}(C_T)' \bar{Z}(C_T) a'_{0,\beta} \Sigma_v^{-1} a_{0,\beta}
 \end{aligned}$$

for $\mathbf{S} = \lim_{T \rightarrow \infty} T^{-1} \mathbf{S}_T$, in analogy with (S.D.21). $\square$

References

- ANDREWS, D. W. K., M. MOREIRA, AND J. H. STOCK (2006): “Optimal Two-Sided Invariant Similar Tests for Instrumental Variables Regression,” *Econometrica*, 74(3), 715–752.
- CASINI, A., A. MCCLOSKEY, L. ROLLA, AND R. PALA (2025): “Supplement Not For Online Publication to ”Dynamic Local Average Treatment Effects in Time Series”,” *Unpublish manuscript*.
- LEHMANN, E. L., AND J. P. ROMANO (2005): *Testing Statistical Hypotheses*. Springer-Verlag New York, 3 edn.
- LEWIS, D. J. (2022): “Robust Inference in Models Identified via Heteroskedasticity,” *Review of Economics and Statistics*, 104(3), 510–524.
- MAGNUSSON, L. M., AND S. MAVROEIDIS (2014): “Identification Using Stability Restrictions,” *Econometrica*, 82(5), 1799–1851.

- NAKAMURA, E., AND J. STEINSSON (2018): “High Frequency Identification of Monetary Non-Neutrality: The Information Effect,” *Quarterly Journal of Economics*, 133(3), 1283–1330.
- RIGOBON, R., AND B. SACK (2003): “Measuring the Reaction of Monetary Policy to the Stock Market,” *The Quarterly Journal of Economics*, 118(2), 639–669.