

Chapter 7: Bartik Instruments

Frontier Topics in Empirical Economics

Lecture notes based on slides by Zibin Huang

College of Business, Shanghai University of Finance and Economics

November 2025

Contents

1	Introduction	3
2	Motivating Examples	4
2.1	Card (2009): Immigration and Inequality	4
2.2	Autor, Dorn, and Hanson (2013): The China Shock	5
3	Goldsmith-Pinkham, Sorkin, and Swift (2020): Share as Instrument	6
3.1	Setup and Two Identities	6
3.2	Two Special Cases	7
3.2.1	Case 1: Two Industries, One Period	8
3.2.2	Case 2: Two Industries, Two Periods	8
3.3	Bartik-GMM Equivalence	9
3.4	Consistency and What Exogeneity Really Requires	10
3.5	The Rotemberg Decomposition	11
3.6	Heterogeneous Treatment Effects: No LATE for Bartik	12
3.7	Empirical Suggestions	13
4	Borusyak, Hull, and Jaravel (2022): Shift as Instrument	14
4.1	Setup	14
4.2	Shock-Level Equivalence	15
4.3	Interpreting the Identification Assumption	16

4.4	Consistency	17
4.5	Empirical Suggestions	18
5	Comparison of the Two Frameworks	19
6	Application: Autor et al. (2013)	20
7	Conclusion	21
A	Homework Problems	22
	Practice Example	24

1 Introduction

In the previous two chapters we developed the general theory of instrumental variables and its extensions beyond the standard LATE framework. In this chapter we zoom in on one particular instrument that is widely used in applied economics: the *Bartik instrument*, also known as the *shift-share instrument*, or SSIV for short. The idea is old but its formal econometric treatment is recent.

The construction is deceptively simple. To instrument a location level endogenous variable such as local employment growth or a local migration inflow, we take a weighted sum of national industry level (or origin country level) shocks, where the weights are baseline shares of each industry or origin group in that location. The intuition is that a location that specialized in a particular industry at some earlier date should be disproportionately affected by a subsequent national shock to that industry, and this differential exposure gives us variation that is arguably exogenous to unobserved local factors.

Bartik instruments show up in a wide range of settings. In labor economics they are used to instrument for local labor demand or for the local inflow of immigrants (Card, 2009). In trade they are used to measure the exposure of US commuting zones to Chinese import growth (Autor et al., 2013). In urban and regional economics they are the workhorse for any question that requires a source of exogenous variation in local shocks.

Two natural questions arise. *How should we use this instrument correctly?* And *what exactly is the identification assumption?* The answer turns out to be more subtle than one might expect from the simplicity of the construction. There are actually two coherent frameworks that give a shift-share instrument a rigorous causal interpretation, and they are not equivalent in general. Goldsmith-Pinkham et al. (2020), henceforth GPSS, treat the shares as the identifying instrument and the shifts as weights. Borusyak et al. (2022), henceforth BHJ, do the opposite: the shifts are the identifying instrument and the shares are weights. Which framework applies to a given application depends on where the exogeneity actually comes from, and reasonable people can and do disagree about that in any particular setting. The good news is that a valid Bartik design can be defended by verifying the assumptions of either framework.

This chapter walks through both frameworks in turn. Section 2 sets up two moti-

vating examples that will accompany us through the chapter. Section 3 develops the GPSS view, including the Bartik-GMM equivalence, the exposure design interpretation, the Rotemberg decomposition into single-industry instruments, and the empirical suggestions that follow. Section 4 develops the BHJ view, including the shock level equivalence, the quasi-random shock assumption, and the inference concerns that come with a shift-share error structure. Section 5 compares the two frameworks and gives guidance on which to use in different settings. Section 6 illustrates with the [Autor et al. \(2013\)](#) China shock application, which both papers use as their running example. Section 7 concludes.

2 Motivating Examples

2.1 Card (2009): Immigration and Inequality

Let us start with two motivating examples of Bartik instrument. [Card \(2009\)](#) studies the effect of the immigrant share of the local labor force on the native-immigrant wage gap. The estimating equation is:

$$y_l = \beta_0 + \beta \ln x_l + \beta_2 C_l + \epsilon_l, \quad (1)$$

where l indexes location (metropolitan area), y_l is the log wage gap between immigrants and natives, x_l is the ratio of immigrant labor to native labor, and C_l collects location level controls.

The immigrant ratio x_l is endogenous. Any positive local productivity shock is likely to raise both the wage gap (through skill sorting, occupational upgrading, and so on) and the immigrant inflow (immigrants respond to labor demand). OLS is therefore biased. Card constructs a shift-share instrument using data for 1980, 1990, and 2000:

$$B_l = \sum_k Z_{lk,1980} \cdot g_k, \quad (2)$$

$$Z_{lk,1980} = (N_{lk,1980}/N_{k,1980}) \times (1/P_{l,2000}), \quad (3)$$

where k indexes home country, $N_{lk,1980}$ is the number of immigrants from country k living in location l in 1980, $N_{k,1980}$ is the total US population of immigrants from k in 1980, $P_{l,2000}$ is the population of l in 2000, and g_k is the total inflow of immigrants from

country k into the United States between 1990 and 2000. The share $Z_{lk,1980}$ measures the base year exposure of location l to origin country k ; the shift g_k measures the size of the national wave from k .

The identification argument runs on two legs. *Relevance* holds because immigrants tend to cluster by origin. Locations that hosted a large Chinese community in 1980 (San Francisco, say) will receive a disproportionate share of the 1990 to 2000 Chinese inflow because new arrivals move to established communities. *Exclusion* holds if the location specific exposure to national origin country flows affects local wages only through the resulting change in the immigrant share and not through other local channels. This second leg is the substantive assumption and it deserves scrutiny.

The construction of B_l decomposes the local immigrant inflow into a location-origin country combination, weighted by baseline local shares and multiplied by national origin country flows. This is a *local share* times *national shift* object, hence the name shift-share, or Bartik, instrument.

2.2 Autor, Dorn, and Hanson (2013): The China Shock

The other running example we use in this chapter is [Autor et al. \(2013\)](#), whose China shock paper is one of the most cited applications of a shift-share instrument in the past two decades. The question is: what is the effect of the surge in Chinese imports on local labor markets in the United States? They construct a location level import exposure variable

$$\Delta IPW_{it} = \sum_j \frac{L_{ijt}}{L_{jt} \cdot L_{it}} \Delta M_{jt},$$

where i indexes commuting zone, j indexes manufacturing industry, and t indexes year. Here L_{ijt} is employment in region i and industry j , L_{jt} is total US employment in industry j , L_{it} is total employment in region i , and ΔM_{jt} is the growth in US imports from China in industry j .

The structure is again a share times a shift. The local industry employment share L_{ijt}/L_{jt} (relative to region size L_{it}) is the local exposure to industry j ; the national import growth ΔM_{jt} is the industry shift. Locations that specialized in industries that were subsequently hit hard by Chinese import competition should have experienced larger declines in local labor market outcomes, and this differential exposure is what the paper exploits.

We will come back to the China shock in Section 6, after we have the two frameworks in hand.

3 Goldsmith-Pinkham, Sorkin, and Swift (2020): Share as Instrument

3.1 Setup and Two Identities

We now step away from the specific examples and set up a general framework. Following [Goldsmith-Pinkham et al. \(2020\)](#), consider the following regression:

$$y_{lt} = D'_{lt}\rho + x_{lt}\beta_0 + \epsilon_{lt}, \quad (4)$$

where l indexes location, t indexes time, D_{lt} is a vector of controls (which may include location and time fixed effects), x_{lt} is an endogenous variable, and β_0 is the parameter of interest. In many applications x and y are growth rates; when they are levels, the researcher should control for location and time fixed effects so that identification works off within-location changes rather than cross sectional levels.

Usually, x_{lt} is correlated with ϵ_{lt} , so we need an instrument. The Bartik construction comes from two algebraic identities. First, we can decompose the location-time level variable x_{lt} into a sum over some finer categorization (industries in the Autor-Dorn-Hanson example, origin countries in the Card example) of a share times a location-category level growth:

Identity 1: Share-Growth Decomposition

$$x_{lt} = Z_{lt}G_{lt} = \sum_{k=1}^K z_{lkt} g_{lkt},$$

where z_{lkt} is the share of category k in location l at time t , and g_{lkt} is the location-category growth rate at t .

This is not an assumption. It is just how one writes a sum. The share z_{lkt} sums to one across k , and g_{lkt} is defined so that the identity holds.

Second, we can decompose the location-category growth into a national industry component plus a residual local component:

Identity 2: National-Local Decomposition

$$g_{lkt} = g_{kt} + \tilde{g}_{lkt},$$

where g_{kt} is the national growth rate of category k at t and $\tilde{g}_{lkt}$ is the location-category deviation from the national trend.

Combining the two identities, one gets

$$x_{lt} = \sum_k z_{lkt} (g_{kt} + \tilde{g}_{lkt}) = \sum_k z_{lkt} g_{kt} + \sum_k z_{lkt} \tilde{g}_{lkt}.$$

Now here comes the substantive move: fix the shares at a baseline period 0 (drop the time subscript on shares), and drop the local shock component $\tilde{g}_{lkt}$ from the sum. What is left is the Bartik instrument:

Definition 3.1 (Bartik Instrument, [Goldsmith-Pinkham et al., 2020](#)).

$$B_{lt} = Z_{l0} G_t = \sum_k \underbrace{z_{lk0}}_{\text{Share}} \underbrace{g_{kt}}_{\text{Shift}}.$$

The first factor is the baseline share of category k in location l ; the second is the national growth of category k at t .

The virtue of this construction is that both factors are, at least on the surface, more plausibly exogenous than x_{lt} itself. The baseline share is fixed in the past and cannot respond to future shocks. The national shift is an aggregate object that averages out any single location's contribution. In some cases, we can also use leave-one-out growth rate to exclude location l when calculating the national growth. So it is close to exogenous with respect to any one location's unobservables. Whether this intuition survives careful scrutiny is exactly what the identification analysis has to establish.

3.2 Two Special Cases

Before we state the general result, it helps to work through two tractable special cases that already convey the main intuition.

3.2.1 Case 1: Two Industries, One Period

Suppose we have only two categories ($K = 2$) and one time period. Since shares sum to one, $z_{l2} = 1 - z_{l1}$, so

$$B_l = z_{l1}g_1 + z_{l2}g_2 = g_2 + (g_1 - g_2)z_{l1}.$$

Now consider the first stage regression of x_l on B_l :

$$x_l = \gamma_0 + \gamma B_l + \eta_l = \underbrace{(\gamma_0 + \gamma g_2)}_{\text{constant}} + \underbrace{\gamma(g_1 - g_2)}_{\text{coefficient}} z_{l1} + \eta_l.$$

Because g_1 and g_2 are constants across locations, using the Bartik instrument B_l in 2SLS is algebraically identical to using the single share z_{l1} as the instrument. All the identifying variation is coming from the cross-location variation in z_{l1} ; the shifts g_1 and g_2 just rescale the coefficient. This is important: in this simplest case, the identification comes from the share.

3.2.2 Case 2: Two Industries, Two Periods

The next case adds time. With $K = 2$ and $T = 2$,

$$B_{lt} = g_{1t}z_{l10} + g_{2t}z_{l20} = g_{2t} + (g_{1t} - g_{2t})z_{l10}.$$

Assume we control for time fixed effects. The first stage becomes

$$x_{lt} = \tau_t + \gamma B_{lt} + \eta_{lt} = \underbrace{(\tau_t + g_{2t}\gamma)}_{\tilde{\tau}_t} + z_{l10}(g_{1t} - g_{2t})\gamma + \eta_{lt}.$$

Now the coefficient $g_{1t} - g_{2t}$ is not constant: it varies with t . So the single-share reduction does not go through as cleanly. Using the indicator function $\mathbf{1}(\cdot)$, we can write

$$g_{1t} - g_{2t} = \mathbf{1}(t = 1)(g_{11} - g_{21}) + \mathbf{1}(t = 2)(g_{12} - g_{22}),$$

and the first stage becomes

$$x_{lt} = \tilde{\tau}_t + z_{l10}\mathbf{1}(t = 1)\underbrace{(g_{11} - g_{21})\gamma}_{\tilde{\gamma}_1} + z_{l10}\mathbf{1}(t = 2)\underbrace{(g_{12} - g_{22})\gamma}_{\tilde{\gamma}_2} + \eta_{lt}.$$

The regression now has time fixed effects and two interactions between the baseline share z_{l10} and each time dummy. The identification exploits both the cross-section variation in z_{l10} and the time variation in the industry growth differential $g_{1t} - g_{2t}$.

Here, the second special case has a **policy exposure research design reading**. The growth rate difference $g_{1t} - g_{2t}$ plays the role of a policy shock size at time t : it says how differently industry 1 grew relative to industry 2 in period t . The baseline share z_{l10} plays the role of exposure to that shock: locations with a high share of industry 1 are more exposed. The question the regression asks is whether locations with higher exposure experience different changes in x following the shock.

In the limit case of a pre-shock period ($t = 1$) with no differential growth ($g_{11} - g_{21} = 0$) and a post-shock period ($t = 2$) with a positive differential growth, this is a DID specification: $\tilde{\gamma}_1 = 0$ is a parallel-trend condition, and $\tilde{\gamma}_2 > 0$ measures the effect of the shock.

So Bartik with a baseline share is really a diff-in-diff on differential exposure. Locations differ in their exposure to a common industry shock, and we compare how outcomes evolve for high-exposure and low-exposure locations. That is the identifying variation.

3.3 Bartik-GMM Equivalence

Now let us state the general result. Assume K industries and one period (the $T > 1$ case is a straightforward extension). Stack observations into matrix form: Y and X are $L \times 1$, Z is $L \times K$ (with entry z_{lk0} in row l column k), and G is $K \times 1$ (with g_k in row k). Let $M_D = I - D(D'D)^{-1}D'$ be the annihilator matrix that partials out D , and write $X^\perp = M_DX$, $Y^\perp = M_DY$, $B^\perp = M_DB$ where $B = ZG$ is the $L \times 1$ Bartik instrument.

Tips: The annihilator matrix M_D is used to partial out the impact of control variable D . For any variable matrix X , M_DX gives the OLS regression residual of X on D . Then by Frisch–Waugh–Lovell Theorem, we can estimate the IV/OLS coefficient of interest by running the regressions using residulized X and Y . For more details, please refer to [Hansen \(2022\)](#).

Proposition 3.1 (Bartik-GMM Equivalence, Proposition 1 in [Goldsmith-Pinkham et](#)

al., 2020). Define the Bartik and GMM estimators

$$\hat{\beta}_{Bartik} = \frac{B'Y^\perp}{B'X^\perp}, \quad \hat{\beta}_{GMM} = \frac{X^{\perp'}ZWZ'Y^\perp}{X^{\perp'}ZWZ'X^\perp}.$$

If the GMM weight matrix is $W = GG'$, then $\hat{\beta}_{Bartik} = \hat{\beta}_{GMM}$.

The interpretation is clean. The Bartik instrument $B = ZG$ is a linear combination of the K share instruments $Z_1, \dots, Z_K$ with weights given by the shifts $g_1, \dots, g_K$. Running 2SLS with the single combined instrument B is exactly the same as running an overidentified GMM with the K separate share instruments $Z_1, \dots, Z_K$, using a specific weight matrix $W = GG'$. **In the GPSS framework, therefore, the identifying instruments are the shares, and the shifts serve as the weights that combine them.**

3.4 Consistency and What Exogeneity Really Requires

For the Bartik estimator to be consistent, we need the standard IV conditions to hold at the level of the share instruments. Goldsmith-Pinkham et al. (2020) state the following.

Assumption 3.1 (Relevance). The instrument is correlated with the endogenous regressor after partialling out controls.

Assumption 3.2 (Strict Exogeneity of Shares). For every category k with $g_k \neq 0$,

$$E[\epsilon_{lt} z_{lk0} \mid D_{lt}] = 0.$$

Proposition 3.2 (Consistency, Proposition 2 in Goldsmith-Pinkham et al., 2020). Under Assumptions 3.1 and 3.2, and asymptotics in the location dimension ($L \rightarrow \infty$ with T, K fixed),

$$\text{plim } \hat{\beta}_{Bartik} = \beta_0.$$

The key assumption is exogeneity of the shares. When is it plausible? Suppose y_{lt} is a level rather than a growth rate. Then ϵ_{lt} absorbs all unobserved determinants of the level, including any persistent unobserved characteristics of location l . The baseline share z_{lk0} is almost certainly correlated with these persistent local characteristics. San Francisco's 1980 share of Chinese immigrants is not independent of San Francisco's unobserved 2000 amenities, culture, or productivity. In levels, Assumption 3.2 fails.

Levels vs Changes: A Practical Warning

Initial shares are typically *not* mean independent of unobserved outcome *levels*. Persistent local characteristics that shape 1980 shares also shape 2000 outcomes. Initial shares are more plausibly mean independent of unobserved outcome *changes*. Once we difference out the level, we are left with idiosyncratic shocks over the sample period that are less likely to be tied to the baseline industry mix. When using a Bartik instrument, either use a growth variable as the outcome, or control for location and time fixed effects. Never run a Bartik IV in cross-sectional levels without location fixed effects.

Substantively, Assumption 3.2 is an *exposure design* assumption, close in spirit to parallel trends in DID. Locations that were more exposed to an industry are more affected by shocks to that industry, but the assumption says there is no systematic difference in *other* unobserved shocks between high-exposure and low-exposure locations. Think of Shanghai and Hong Kong versus Shenyang and Wuhan. Shanghai and Hong Kong specialize in finance, so a global financial crisis hits them harder. For the Bartik instrument to be valid, we need to assume that in the absence of the crisis shock, the outcome trends in the four cities would have been parallel. Any unobserved shocks that hit finance-heavy cities differentially would invalidate the design.

3.5 The Rotemberg Decomposition

The GMM equivalence result tells us that a Bartik IV is really a combination of many single-industry instruments. A natural question follows: which industry is doing the identification work? A Bartik point estimate could be driven by one dominant industry that happens to have a large national shock, or it could be a genuine average across many industries with comparable weights. This matters for interpretation and for robustness. Goldsmith-Pinkham et al. (2020) give a clean decomposition.

Proposition 3.3 (Rotemberg Decomposition, Proposition 3 in Goldsmith-Pinkham et al., 2020). *The Bartik estimator can be written as*

$$\hat{\beta}_{Bartik} = \sum_k \hat{\alpha}_k \hat{\beta}_k,$$

where

$$\hat{\beta}_k = (Z_k' X^\perp)^{-1} Z_k' Y^\perp$$

is the just-identified IV estimator that uses only the share Z_k of industry k as an instrument, and the weights

$$\hat{\alpha}_k = \frac{g_k Z'_k X^\perp}{\sum_{k'} g_{k'} Z'_{k'} X^\perp}$$

are the Rotemberg weights. The weights sum to one.

Each $\hat{\beta}_k$ is what we would get if we used a single share instrument Z_k ($L \times 1$ vector with each row representing the share of industry k in one specific location) to instrument X . The Rotemberg weight $\hat{\alpha}_k$ measures how much industry k 's single-instrument IV estimate contributes to the combined Bartik estimate. An industry with a big Rotemberg weight is one whose identification assumption is critical for the overall estimate: if that industry's share is not really exogenous, the Bartik estimate will be badly biased. An industry with a small Rotemberg weight can be misspecified without much damage.

Practical Guidance

Report the industries with the highest Rotemberg weights. Discuss the plausibility of Assumption 3.2 for those industries in particular. If the whole identification is really riding on one or two shares, be honest about it and defend those assumptions specifically.

A useful conceptual clarification: the Rotemberg decomposition is different from the GMM equivalence in Proposition 3.1. The GMM equivalence says that Bartik in 2SLS with a single combined instrument B equals GMM with K share instruments in a *joint estimation*, using $W = GG'$. The Rotemberg decomposition says that the Bartik estimator can be rewritten as a weighted average of K *separately estimated single-instrument IV estimators*. Both statements are true and both use the same object, but they describe different aspects of the estimator.

3.6 Heterogeneous Treatment Effects: No LATE for Bartik

In the LATE framework of Imbens and Angrist (1994), we saw that a single binary IV with monotonicity identifies the average treatment effect for compliers. Does a Bartik instrument have an analogous LATE interpretation? The answer, developed in Goldsmith-Pinkham et al. (2020), is negative in general. Even in a restricted heterogeneous effect model, the Bartik IV does not have a LATE interpretation. This is not surprising given what we have learned in Chapter 6.

First, for the single share instrument case, in a restricted heterogeneous effect model with a set of monotonicity style assumptions, the Bartik IV can be written as a combination of location level treatment effects with positive weights.

Second, however, for the combined Bartik instrument, monotonicity of each single share instrument is not enough for us to have positive weights. The reason is that different share instruments can push locations in different directions, and combining them can produce non-monotone shifts. Therefore, the weights can be negative. Negative weights make the estimator uninterpretable as any convex average of underlying effects.

You will see the same story unfold in later chapters when we discuss DID with staggered treatment timing and heterogeneous effects. Simple regressions turn out to average local treatment effects in ways that can be pathological when the underlying effects are complicated (heterogeneous, dynamic, etc...), and this is not a coincidence: it is an intrinsic feature of trying to summarize heterogeneous causal effects with a single regression coefficient.

3.7 Empirical Suggestions

The GPSS framework leads to several concrete pieces of empirical advice.

Always include location and time fixed effects. As discussed above, initial shares are plausible instruments for outcome *changes*, not levels. In a panel setting, location and time fixed effects deliver the difference specification automatically. In a cross section, use a growth outcome variable rather than a level.

Focus on high Rotemberg weight industries. Report the top industries by weight and defend the exogeneity assumption for each one. If most of the weight is on a single industry, the paper is really identified off that industry, and the discussion should reflect that.

Test 1: Balance on covariates. Suppose there is a location level control d_l that plausibly predicts changes in y through some channel other than x . Test whether the baseline shares are correlated with d_l . Because the shares are supposed to affect y only through x , an economically or statistically significant correlation is a red flag. [Goldsmith-Pinkham et al. \(2020\)](#) suggest also controlling for higher-level aggregate shares (for instance, if industry is defined at the 4-digit level, control for 2-digit shares) to soak up broad category composition.

Test 2: Pre-trends. If the data include a pre-shock period, the specification with baseline shares interacted with pre-period dummies is a DID with growth rate as the shock size. Under the exposure design interpretation, the pre-period coefficient $\tilde{\gamma}_1$ should be zero. Test this for both the overall Bartik IV and for the single-industry instruments with the highest weights.

Test 3: Overidentification. The Bartik-GMM equivalence says that Bartik is an overidentified GMM with shares as instruments. Standard overidentification tests (Hansen’s J) can be run at the level of the individual share instruments. As always, a rejection has different potential interpretations: it could be at least one instrument is invalid, or the treatment effect is heterogeneous across instruments, or some other model misspecification. The test cannot distinguish them, so it should be interpreted with care. [Goldsmith-Pinkham et al. \(2020\)](#) does not recommend this test.

4 Borusyak, Hull, and Jaravel (2022): Shift as Instrument

The GPSS framework treats shares as the identifying instrument, with the shifts serving as weights. [Borusyak et al. \(2022\)](#) propose the opposite reading: treat the shifts as the identifying instrument, with the shares serving as weights. The identification argument then rests on random assignment of the industry level shocks, not on exogeneity of the location shares.

This is a substantively different framework. It is not that one is right and the other wrong; they capture different empirical situations. When exogeneity plausibly comes from the shares, GPSS is the right framework. When exogeneity plausibly comes from the shifts, BHJ is the right framework.

4.1 Setup

BHJ consider a location level regression

$$y_l = \beta x_l + w_l' \gamma + \epsilon_l,$$

where w_l is a vector of controls. The shift-share instrument has the general form

$$z_l = \sum_k s_{lk} g_k, \quad k = 1, 2, \dots, K,$$

with s_{lk} the share of category k in location l and g_k the national shift for category k . The moment condition that validates the instrument is

$$E \left[\sum_l z_l \epsilon_l \right] = 0.$$

The question is: what does this look like when we plug in the shift-share definition of z_l ?

4.2 Shock-Level Equivalence

Substitute $z_l = \sum_k s_{lk} g_k$ into the moment condition:

$$E \left[\sum_l z_l \epsilon_l \right] = E \left[\sum_l \sum_k s_{lk} g_k \epsilon_l \right].$$

Exchange the order of summation and factor out g_k , which does not depend on l :

$$\begin{aligned} E \left[\sum_l z_l \epsilon_l \right] &= E \left[\sum_k \sum_l s_{lk} g_k \epsilon_l \right] = E \left[\sum_k g_k \sum_l s_{lk} \epsilon_l \right] \\ &= E \left[\sum_k g_k \left(\frac{\sum_l s_{lk} \epsilon_l}{\sum_l s_{lk}} \right) \left(\sum_l s_{lk} \right) \right] \\ &= E \left[\sum_k s_k g_k \bar{\epsilon}_k \right], \end{aligned}$$

where we have defined

$$s_k = \sum_l s_{lk}, \quad \bar{\epsilon}_k = \frac{\sum_l s_{lk} \epsilon_l}{\sum_l s_{lk}}.$$

The quantity s_k is the total exposure to category k summed across locations. The quantity $\bar{\epsilon}_k$ is a share-weighted average of the location errors, aggregated up to the shock (category) level; being a ratio, it is invariant to the normalization.

This algebraic manipulation delivers a remarkable observation. The location level mo-

ment condition is equivalent to the shock (industry) level moment condition

$$E\left[\sum_k s_k g_k \bar{\epsilon}_k\right] = 0.$$

So even though we thought we were running an IV at the location level, the identifying moment is really at the shock level, where the instrument g_k is required to be uncorrelated with a share-weighted aggregate of the location errors. BHJ formalize this as follows.

Proposition 4.1 (Shock-Level Equivalence, Proposition 1 in [Borusyak et al., 2022](#)). *The shift-share IV estimator $\hat{\beta}$ equals the second stage coefficient of a s_k -weighted shock-level IV regression using the shifts g_k as the instrument in estimating*

$$\bar{y}_k = \alpha + \beta \bar{x}_k + \bar{\epsilon}_k,$$

where $\bar{v}_k = \sum_l s_{lk} v_l / \sum_l s_{lk}$ is the exposure-weighted average of a variable v_l across locations.

This equivalence has two immediate uses. First, it identifies the true unit of identification. Even though we have L location observations, the identifying variation is really over K industries. If K is small, the effective sample size for the identifying moment is small, regardless of how big L is. Second, it clarifies what exogeneity means. The identifying assumption is that the shocks g_k are as good as randomly assigned, in the sense that they are uncorrelated with the industry level unobservable $\bar{\epsilon}_k$.

4.3 Interpreting the Identification Assumption

Let us make the assumption concrete. Suppose we are estimating the effect of a US tariff on employment in China. The tariff shocks g_k vary by industry (steel, textiles, semiconductors, etc.). The location aggregate $\bar{\epsilon}_k$ is a share-weighted average of local unobserved supply shocks, aggregated across locations by how important industry k is in each place.

For the BHJ identification to hold, we need $g_k \perp \bar{\epsilon}_k$: industries that experience a rise in the US tariff should not systematically face different labor supply shocks in the regions where they concentrate. If Chinese steel is intensively produced in Hebei, then a US tariff on steel is a big shock to Hebei. We must assume that Hebei was

not simultaneously receiving some other unobserved shock, such as a change in the Gaokao enrollment quota, that would independently affect Hebei's labor market and be correlated with the tariff.

This is a substantive assumption about the sources of variation in g_k . It is often defensible when the shocks are driven by conditions in some other country (US tariffs, from the perspective of Chinese labor markets) or by aggregate technology, but the researcher has to make the case.

4.4 Consistency

BHJ establish consistency under two assumptions.

Assumption 4.1 (Quasi-random Shock Assignment, Assumption 1 in [Borusyak et al., 2022](#)). Conditional on the shares s and the industry level unobservable $\bar{\epsilon}$, the shocks have a common mean:

$$E[g_k \mid \bar{\epsilon}, s] = \mu.$$

Assumption 4.2 (Many Uncorrelated Shocks, Assumption 2 in [Borusyak et al., 2022](#)). The industry-level Herfindahl index of *normalized* exposure goes to zero,

$$E\left[\sum_k \tilde{s}_k^2\right] \rightarrow 0, \quad \tilde{s}_k = \frac{s_k}{\sum_{k'} s_{k'}},$$

and the shocks are pairwise uncorrelated given exposure,

$$\text{Cov}[g_k, g_{k'} \mid \bar{\epsilon}, s] = 0 \text{ for } k \neq k'.$$

The first assumption says the shocks look randomly assigned conditional on the exposures. The second says exposure should not be too concentrated in a few industries, and that shocks are independent across industries. This second assumption is the substantive analogue of $L \rightarrow \infty$ in a standard IV setting: we need many effectively independent draws for the LLN to work. In this framework the effective sample size is at the shock (industry) level, not the location level.

Proposition 4.2 (Consistency, Proposition 3 in [Borusyak et al., 2022](#)). *Under Assumptions 4.1 and 4.2 plus standard regularity conditions,*

$$\hat{\beta} \xrightarrow{p} \beta.$$

4.5 Empirical Suggestions

Three important pieces of practical advice come out of the BHJ framework.

Inference: watch out for standard errors.

Adao et al. (2019) show that traditional (heteroskedasticity-robust or clustered-by-location) standard errors for a shift-share IV are typically wrong. The reason is that the location level observations are not i.i.d.: the same national shock g_k enters the regression for every location that has any share in industry k , so the errors ϵ_l have a common component and are mechanically correlated across locations with similar industry mixes. Traditional standard errors ignore this correlation and understate uncertainty, sometimes by a lot. BHJ show that the equivalent shock-level regression (Proposition 4.1) does not have this problem: at the industry level, the standard errors from the aggregated regression are valid. This gives a clean practical recipe: run the shock level regression and read off the standard error. The Stata package `ssaggregate` automates the aggregation.

Effective sample size and Herfindahl index.

When using the BHJ framework, we must validate Assumption 4.2 for consistency. Intuitively, we must have a large number of industries compared with locations and shocks should not concentrate on only a few industries. In practice, we can directly calculate the industry-level Herfindahl index of normalized exposure to make sure it is small enough.

Balance tests.

Location-level balance. Regress a pre-determined location level control r_l (GDP, population, education) on the SSIV z_l . A significant coefficient is a red flag. This can be combined with an Oster-style bound to assess how much unobserved selection would be needed to overturn the estimate. This is similar to the balance test proposed by GPSS.

Location-level balance using an industry-level confounder. Take an industry level covariate r_k , construct the location aggregate $r_l = \sum_k s_{lk} r_k$, and regress that on the SSIV.

Shock-level balance. Aggregate a location level confounder r_l to the industry level via $r_k = \sum_l s_{lk} r_l / \sum_l s_{lk}$, and regress that on the shock g_k . In the BHJ framework this is

arguably the most fundamental balance test, because it directly tests the shock-level exogeneity assumption. If shocks g_k are supposed to be quasi-randomly assigned across industries, they should not correlate with any observable industry level characteristic aggregated from location data.

5 Comparison of the Two Frameworks

We now have two frameworks that give a shift-share IV a rigorous causal interpretation. They are not competitors so much as different maps of the same territory. Which map to use depends on the terrain of the specific application. Specifically, in one study, there is no need for us to show both shares and shifts are exogenous. As long as the identification assumptions for **either** of these two frameworks are validated, we are safe.

Comparison of GPSS and BHJ

- **GPSS** ([Goldsmith-Pinkham et al., 2020](#))
 - Equivalence: joint GMM using shares as instruments, shifts as weights.
 - Design: exposure DID with shares as differential exposure.
 - Assumption: locations with different shares have parallel trends absent the shock.
- **BHJ** ([Borusyak et al., 2022](#))
 - Equivalence: shock-level regression with shifts as instruments, shares as weights.
 - Design: random assignment of shifts across industries.
 - Assumption: industry level shocks are uncorrelated with industry-aggregated location unobservables.

Under what conditions should we favor one framework over the other?

Prefer GPSS when:

- The exogenous variation is coming from the shares (historical settlement patterns, legacy industrial structure, or some other pre-determined local composition).
- The design is an exposure design in which differential exposure to a common shock is the identifying variation.

- There is a small fixed number of industries, and the sampling is over locations (K fixed, $L \rightarrow \infty$).
- The paper focuses on a few specific industries whose shares carry policy meaning.

Prefer BHJ when:

- The exogenous variation is coming from the shifts (industry level shocks that plausibly look random).
- There is a small fixed number of locations, and the sampling is over industries (L fixed, $K \rightarrow \infty$).
- The second stage error has a shift-share structure that generates mechanical correlation between the SSIV and the location error, so that inference needs the shock-level correction.

6 Application: [Autor et al. \(2013\)](#)

The canonical modern application of Bartik instruments is the China shock paper of [Autor et al. \(2013\)](#). Both [Goldsmith-Pinkham et al. \(2020\)](#) and [Borusyak et al. \(2022\)](#) use it as their running empirical example, precisely because it is representative of how the technique is used in trade and labor economics.

The research question is: what is the effect of the surge in Chinese imports on local labor market outcomes in the United States? The independent variable is a location i level import exposure measure at time t , ΔIPW_{it} , constructed as

$$\Delta IPW_{it} = \sum_j \frac{L_{ijt}}{L_{jt} \cdot L_{it}} \Delta M_{jt}.$$

where i is commuting zone in the U.S. and j refers to industry. L_{ij} is the number of labor forces working in industry j in commuting zone i . ΔM_{jt} is the change in U.S. imports from China in industry j . The regression outcomes are commuting zone level manufacturing employment, wages, and various welfare and public assistance measures.

The problem for OLS is the reverse causality: a US productivity shock could increase Chinese imports (through demand for cheap inputs) while also affecting local labor markets directly. Autor, Dorn, and Hanson address this with an instrumental variables strategy: instrument ΔIPW_{it} (US) with an analogous exposure variable that uses

Chinese imports to eight other high-income countries as the shift. The idea is that Chinese import growth to Germany, Japan, and so on reflects supply-side factors in China (productivity, exchange rates, WTO accession) rather than demand-side factors in the United States. This is a typical Bartik IV.

In the GPSS reading, the identifying instruments are the initial US industry shares L_{ijt}/L_{jt} : commuting zones with more manufacturing employment in industries that will later experience large Chinese import growth are more exposed to the shock. The parallel-trend assumption is that, absent the China shock, commuting zones with different industry compositions would have followed parallel trajectories.

In the BHJ reading, the identifying instruments are the industry level Chinese import shocks ΔM_{jt} . The quasi-random assignment assumption is that the industry-by-industry pattern of Chinese import growth (to other high-income countries) is not systematically related to U.S. industry-level unobservables that also affect the aggregated US local labor market outcomes.

Both readings are defensible for the China shock, and both papers show explicitly how their frameworks apply. Students are strongly encouraged to read [Autor et al. \(2013\)](#) together with the discussion sections of [Goldsmith-Pinkham et al. \(2020\)](#) and [Borusyak et al. \(2022\)](#) to see how the same empirical setting is reinterpreted through the two lenses.

7 Conclusion

The Bartik or shift-share instrument is a workhorse tool in applied economics, especially in the trade and labor subfields. Its construction is simple: a weighted sum of national industry shocks, with weights given by baseline local industry shares. Its identification, however, requires care.

We have introduced two frameworks that give the shift-share design a rigorous causal interpretation. The first, [Goldsmith-Pinkham et al. \(2020\)](#), treats shares as instruments and shifts as weights. Bartik is equivalent to an over-identified GMM using the K shares. The identifying assumption is exogeneity of the shares, which is plausible for outcome changes but almost never for outcome levels, so specifications must include location and time fixed effects or use a growth outcome. It is essentially one type of exposure design. The Rotemberg decomposition tells us which industries drive the

identification and highlights where the exogeneity assumption most matters.

The second framework, [Borusyak et al. \(2022\)](#), treats shifts as instruments and shares as weights. Bartik is equivalent to a shock-level regression using shifts. The identifying assumption is quasi-random assignment of the industry shocks. Standard errors must account for the shift-share error structure.

The two frameworks are not competitors but complements. A good Bartik paper argues explicitly which framework applies and defends the corresponding assumptions. Applications where exogeneity flows from historical settlement or long-established industrial composition fit the GPSS reading. Applications where the shocks are plausibly random draws from an industry level distribution fit the BHJ reading. The China shock of [Autor et al. \(2013\)](#) is a case in point: both papers use it and both frameworks can be defended.

In the next chapter we turn from IV to difference-in-differences, which shares more with the Bartik exposure design than the surface would suggest. We will see the same themes recur: heterogeneous effects that a simple regression cannot summarize, careful design-based reasoning about where identifying variation actually comes from, and modern econometric tools that clean up the messiness.

A Homework Problems

1. Verifying the Bartik-GMM Equivalence.

Consider the special case of $K = 2$ industries and one time period with no controls (so D is empty and $X^\perp = X$). Suppose there are L locations with observed data $(y_l, x_l, z_{l1}, z_{l2})$ and national shifts (g_1, g_2) . Let $B_l = z_{l1}g_1 + z_{l2}g_2$.

- (a) Write out the Bartik estimator $\hat{\beta}_{\text{Bartik}} = B'Y/B'X$ in summation notation.
- (b) Write out the GMM estimator with weight matrix $W = GG'$ where $G = (g_1, g_2)'$. Simplify.
- (c) Show algebraically that the two coincide.
- (d) Now substitute the constraint $z_{l2} = 1 - z_{l1}$. Show that the Bartik estimator reduces to a just-identified IV estimator using z_{l1} alone.

2. Rotemberg Weights and Interpretation.

Consider a Bartik instrument with $K = 3$ industries. Suppose the single-industry IV estimates and unnormalized weights are:

Industry k	$\hat{\beta}_k$	Unnormalized weight $g_k Z'_k X^\perp$
1	2.0	0.6
2	1.0	0.3
3	-3.0	0.1

- Compute the normalized Rotemberg weights $\hat{\alpha}_k$ and verify they sum to one.
- Compute the combined Bartik estimator $\hat{\beta}_{\text{Bartik}} = \sum_k \hat{\alpha}_k \hat{\beta}_k$.
- Suppose the exogeneity assumption is violated for industry 1 but holds for industries 2 and 3. Which industry's misspecification would do the most damage to the Bartik estimate, and why? Which industries could be misspecified with relatively little damage?
- Now suppose we relabel the industries so that industry 3 is called industry 1 (swap the labels). Do the Rotemberg weights and $\hat{\beta}_{\text{Bartik}}$ change?

3. Deriving the Shock-Level Equivalence.

Consider the BHJ setup with one period, no controls, and a shift-share instrument $z_l = \sum_k s_{lk} g_k$. Assume $\sum_l s_{lk} = s_k$ for each k .

- Show that the location-level moment condition $E[\sum_l z_l \epsilon_l] = 0$ can be rewritten as the shock-level moment condition $E[\sum_k s_k g_k \bar{\epsilon}_k] = 0$, where $\bar{\epsilon}_k = \sum_l s_{lk} \epsilon_l / \sum_l s_{lk}$.
- Explain in words the difference between the location level exogeneity assumption $E[z_l \epsilon_l] = 0$ and the shock level exogeneity assumption $E[g_k \bar{\epsilon}_k] = 0$. Which of the two is easier to defend in the context of the China shock?
- Suppose one industry accounts for 60% of total exposure, $\tilde{s}_1 = 0.6$, while the remaining $K - 1$ industries share the remaining 40% equally. Discuss what this concentration does to Assumption 4.2 and to the effective sample size for identification.

Practice Example: Bartik IV with Simulated Shift-Share Data

In this practice example we simulate a shift-share environment, verify that OLS is biased, verify that the Bartik instrument recovers the true structural coefficient, and compute the Rotemberg decomposition. The setup follows the GPSS framework, but the equivalent shock-level regression from BHJ is computed at the end as a cross check.

Data generating process. Consider $L = 300$ locations and $K = 6$ industries. Baseline shares z_{lk0} are drawn from a Dirichlet distribution so they sum to one for each location. National industry shifts g_k are measured in percentage points and fixed at

$$(g_1, g_2, g_3, g_4, g_5, g_6) = (15, -3, 2, 8, -5, 3),$$

so industry 1 has a large positive shock, industry 4 a moderate positive shock, and the rest are small or negative. Measuring the shifts (and hence x and y) in percentage points is not cosmetic: because B_l is a convex combination of the g_k , writing the same values as decimal fractions would put the variance of the instrument two orders of magnitude below the noise variance in x , the first stage F statistic would fall below one, and no IV method could work. A local unobserved shock $\epsilon_l \sim \mathcal{N}(0, 1)$ is drawn independently of everything else. The endogenous variable is

$$x_l = \sum_k z_{lk0} g_k + \lambda \epsilon_l + \eta_l,$$

where $\lambda = 0.8$ generates the endogeneity and $\eta_l \sim \mathcal{N}(0, 0.5)$ is a pure first-stage shock. The outcome is

$$y_l = \alpha + \beta_0 x_l + \epsilon_l,$$

with $\beta_0 = 1.0$ the true parameter.

Predicted results. OLS regresses y_l on x_l and picks up $\beta_0 + \lambda \text{Var}(\epsilon_l) / \text{Var}(x_l)$, which for these parameters is roughly $1 + 0.8/4.6 \approx 1.17$. Bartik IV using $B_l = \sum_k z_{lk0} g_k$ should recover $\beta_0 \approx 1$. Rotemberg weights should be dominated by industry 1 (largest positive shift), with moderate weights on industries 4 and 5 (industry 5 pairs a negative shift with a negative share covariance, which produces a positive weight) and small weights, possibly negative, on the rest. The BHJ shock-level regression, with six

observations, should give an identical point estimate to the location-level Bartik IV.

Analysis steps. The do-file (`src/chapter 7/ch7_bartik.do`) performs:

1. Simulate the shift-share DGP.
2. Descriptive statistics: verify shares sum to one, tabulate the fixed shifts, summarize x and y .
3. OLS of y on x (biased).
4. Bartik IV: construct B_l , run 2SLS.
5. Single-industry IV: for each k , run IV using z_{lk0} alone; collect $\hat{\beta}_k$.
6. Rotemberg weights: compute $\hat{\alpha}_k$ and verify $\sum_k \hat{\alpha}_k \hat{\beta}_k = \hat{\beta}_{\text{Bartik}}$.
7. BHJ shock-level equivalence: aggregate y and x to the industry level via share-weighted averages, run the shock-level IV using g_k as the instrument, and verify the coefficient equals the Bartik estimate.
8. Key lessons summary.

Key lessons.

- OLS is biased when the endogenous variable is correlated with the outcome error. The direction and magnitude of the bias depend on the sign and size of λ .
- The Bartik instrument B_l is exogenous by construction, because both the baseline shares and the national shifts are drawn independently of ϵ_l . Bartik IV recovers β_0 .
- The Rotemberg decomposition reveals which industries drive the estimate. Industries with large shifts and shares that are systematically correlated with x get large weights; industries with small shifts or shares uncorrelated with x get small weights.
- The BHJ shock-level equivalence gives the same point estimate as the location-level Bartik IV, confirming Proposition 4.1. In real applications this equivalence delivers valid standard errors via `ssaggregate`, which the location level regression cannot.

The do-file is located at `src/chapter 7/ch7_bartik.do` and the simulated dataset at `data/chapter 7/ch7_bartik.dta`.

References

- Adao, Rodrigo, Michal Kolesár, and Eduardo Morales. 2019. “Shift-share Designs: Theory and Inference.” *Quarterly Journal of Economics*, 134(4): 1949–2010.
- Autor, David H., David Dorn, and Gordon H. Hanson. 2013. “The China Syndrome: Local Labor Market Effects of Import Competition in the United States.” *American Economic Review*, 103(6): 2121–2168.
- Borusyak, Kirill, Peter Hull, and Xavier Jaravel. 2022. “Quasi-experimental Shift-share Research Designs.” *Review of Economic Studies*, 89(1): 181–213.
- Card, David. 2009. “Immigration and Inequality.” *American Economic Review*, 99(2): 1–21.
- Goldsmith-Pinkham, Paul, Isaac Sorkin, and Henry Swift. 2020. “Bartik Instruments: What, When, Why, and How.” *American Economic Review*, 110(8): 2586–2624.
- Hansen, Bruce. 2022. *Econometrics*. Princeton University Press.
- Imbens, Guido W. and Joshua D. Angrist. 1994. “Identification and Estimation of Local Average Treatment Effects.” *Econometrica*, 62(2): 467–475.