

Know the Economy, Not the Data Lake: Country Structure and the Right Number of Indicators in a GDP Nowcast

Karan Singh Bagavathinathan*

Independent Economist

Tandley Omprakash Sridevi†

Bharathidasan Institute of Management, India

Abstract

How many indicators should a GDP nowcast use, and how should that number be chosen? A common answer is to add every available series and let a factor model or a machine-learning method separate signal from noise. We argue that this misplaces the effort for most economies. The number of indicators that can help is bounded by how many distinct sources of cyclical signal an economy generates, which is a feature of its structure rather than of the available data. We write the choice as a bias–variance trade-off. The accuracy-maximising count k^* rises with the sample length and with the number of orthogonal signal dimensions, and it falls with volatility, so less-diversified economies nowcast best with fewer indicators. We test this on twelve economies. Each faces the same harmonised menu of headline indicators, informed by the blocks that published flagship nowcasts use, over a common 2000 to 2026 out-of-sample window. The optimum is smaller in emerging economies (mean 3.2 against 6.4; $p = 0.026$), and it is explained by economic complexity (0.61) rather than income (0.35, not significant). Norway is the identifying case. It is rich but resource-concentrated, and its optimum sits with the emerging economies, where complexity predicts it and income does not.

Keywords: GDP nowcasting; forecast evaluation; indicator selection; economic complexity; emerging economies; bias–variance trade-off.

JEL: C52, C53, E37, O11, O47.

*Corresponding author. Independent Economist. E-mail: vbk.singh@gmail.com.

†E-mail: sridevi.tandley@bim.edu.

1 Introduction

How many indicators should a nowcast use, and how should one decide? The default modern answer is to let the data decide: assemble every series within reach (hundreds of them) and let a factor model, a shrinkage estimator, or a machine-learning algorithm sift the signal from the noise (Stock and Watson, 2002; Giannone et al., 2008; Bańbura et al., 2013; Bok et al., 2018). On long samples of advanced economies this works, and the marginal indicator rarely hurts. Yet a persistent counter-current runs through the forecasting literature. The forecast-combination puzzle (Stock and Watson, 2004; Smith and Wallis, 2009), the “less-is-more” effect (Gigerenzer and Brighton, 2009), and the illusion of sparsity in macroeconomic prediction (Giannone et al., 2021) all report the same pattern, that simple, parsimonious combinations routinely beat richly parameterised “optimal” ones. This paper argues that the two observations are one phenomenon seen from opposite ends of a single gradient. Placing a given economy on that gradient is a question of macroeconomic understanding rather than data mining. It depends on how many genuinely distinct sources of cyclical signal the economy generates, and on how long and how stably it has been measured. On this view the number of indicators a nowcast should use is a structural property of the economy, to be reasoned about from the composition of its output and the length of its record, rather than a quantity to be maximised over whatever data happen to be available.

We make the mechanism explicit. Write the expected out-of-sample forecast error as the sum of a squared-bias term that falls with the number of indicators k and an estimation-variance term that rises with k . The error is then U-shaped, and its minimiser k^* is (i) increasing in the effective sample length N , (ii) increasing in the number of orthogonal signal dimensions r , and (iii) decreasing in the per-parameter noise scale κ^2 (Section 3). These three comparative statics map onto the stage of an economy’s statistical and structural development. An emerging economy has a short reliable history (small N), a narrow and strongly co-moving formal sector (small r), and elevated volatility (large κ^2), all of which push k^* down. A mature, diversified economy occupies the opposite corner. On this reading the large-dataset tradition is the high- N , high- r limit of a more general relationship rather than a universal prescription.

This distinction matters as nowcasting takes up the tools of machine learning. Gradient boosting, random forests, penalised regressions, and neural networks are increasingly trained on large indicator sets. Where the sample is long, this works well, and it reaches the same place as a hand-chosen compact set. A shrinkage estimator fed many series strips the redundant ones, so automatic and hand-crafted parsimony coincide. The two approaches are two routes to one optimum, and the choice between them is largely a matter of conve-

nience. They part company only where the sample is short. No estimator can extract signal a source does not carry, a limit as old as information theory (Shannon, 1948) and familiar from the no-free-lunch theorem (Wolpert, 1996) and the irreducible bias–variance floor of any estimator (Geman et al., 1992); that more data need not improve a macroeconomic forecast is itself established (Boivin and Ng, 2006). Our contribution is to make this ceiling a cross-country object. Once an economy’s r genuine signal dimensions are spanned, further indicators add collinear noise, and at the short samples typical of emerging economies the estimation-variance term $\kappa^2 k/N$ overwhelms any further bias reduction, so shrinkage can no longer separate weak signal from noise. There, and only there, over-inclusion carries a real accuracy cost, and structural knowledge of the economy does work the data cannot do for the model. This is a boundary result rather than a verdict on any method. The two approaches agree where data are abundant, and they diverge only where data are scarce.

Testing this requires a design that separates the economics of the claim from an artefact of data availability. The design has three features. *First*, rather than run every country through a kitchen-sink of whatever series we could scrape, which would confound the optimum with our collection effort, we give every country the same compact, harmonised menu of headline indicators, one series per standard nowcasting block, with the blocks each country carries informed by its own published flagship nowcast (Section 4). Because the menu is common across countries, cross-country differences in the optimum cannot come from country-specific indicator selection. *Second*, we fix a common 2000–2026 estimation window across all twelve economies, so that differences in the estimated optimum cannot be attributed to differences in sample length; holding N fixed isolates economic structure, r and volatility, as the remaining source of cross-country variation. *Third*, we apply one identical estimator to every country, a LASSO-ordered expanding-window out-of-sample procedure, and read the optimum off the same object, the minimiser of the quarterly out-of-sample error curve. We study the optimal indicator *count* under this uniform quarterly design. We do not exploit within-quarter publication timing or mixed frequency, and we do not produce an operational real-time nowcast. The count question is separable from the timing question, and we address the first.

Three findings follow. The estimated optimum is systematically smaller in emerging than in advanced economies (mean argmin, the argument of the minimum of the out-of-sample error curve, 3.2 versus 6.4; $p = 0.026$). What governs it is economic complexity, not income. The estimated k^* correlates with the Economic Complexity Index (Hidalgo and Hausmann, 2009) at +0.61 ($p = 0.035$), but with the logarithm of GDP per capita at only +0.35 ($p = 0.27$, insignificant). Norway is the case that identifies the distinction. It is the richest economy in the panel, yet because its exports are concentrated in petroleum it is informationally simple, and its estimated optimum sits at the emerging end of the gradient,

where complexity predicts and income does not. The optimal counts also track, in rough order, the size of the nowcasting systems that central banks and statistical agencies field, though we treat that comparison as informal, because published counts are not measured on a common basis across institutions (Section 5).

The contribution is threefold. We state a falsifiable proposition linking nowcast parsimony to informational maturity and derive its comparative statics. We test it on a cross-country panel under a uniform, precedent-anchored, quarterly out-of-sample protocol, and show that the optimal indicator count varies as predicted and is governed by complexity rather than wealth. Section 3 states the mechanism; Section 4 describes the data and method; Section 5 reports the gradient and the complexity result; Section 6 discusses identification, limitations, and implications.

2 How nowcasting evolved, and the question it never settled

Nowcasting was born of a mismatch in time. The single number that most disciplines monetary and fiscal policy, quarterly GDP, arrives forty-five to seventy days after the quarter it describes, and in revised form for years after that, while the decisions that depend on it cannot wait. The first responses were pragmatic: bridge equations that regressed quarterly GDP on a few monthly indicators once they printed, aggregating high-frequency data up to the quarter by brute force. These worked, but they were *ad hoc*. Each modeller chose a handful of series by judgement, and there was neither a principle for which indicators to include nor a way to update the forecast coherently as data trickled in through the quarter.

The dynamic-factor revolution supplied the missing structure. Giannone et al. (2008) recast nowcasting as the extraction of a latent state from a real-time flow of releases: a small number of common factors summarise a large panel of indicators, the Kalman filter updates the GDP estimate exactly as each new series arrives, and the marginal informational content of any single release can be read off the filter’s news decomposition. This was a genuine paradigm. It turned nowcasting from a craft into a method, and it scaled: because factors compress many indicators into a few dimensions, there was no longer any apparent penalty to adding data. The framework was institutionalised almost immediately (the New York Fed Staff Nowcast, the Atlanta Fed’s GDPNow, the ECB and dozens of other central banks now run descendants of it (Bańbura et al., 2013; Bok et al., 2018)), and with institutionalisation came a working doctrine: when in doubt, add more series. Bok et al. (2018) survey nowcasts drawing on many dozens to hundreds of indicators; the implicit answer to “how many

indicators?” had become “as many as the data flow provides.”

Yet a stubborn counter-current had been running the whole time, and it never reconciled with the maximalist doctrine. Decades of forecasting experience kept producing the opposite lesson: simple, parsimonious combinations beat richly parameterised “optimal” ones. The forecast-combination puzzle (Stock and Watson, 2004; Smith and Wallis, 2009) documented that an equal-weighted average of a few forecasts routinely outperforms the theoretically optimal weighted combination, because the weights themselves must be estimated and the estimation noise swamps the gain. The “less-is-more” literature (Gigerenzer and Brighton, 2009) showed the same in cognition: past a point, more information degrades inference. And most pointedly for the big-data era, Giannone et al. (2021) found, across exactly the kind of macroeconomic prediction problems nowcasting embodies, an “illusion of sparsity”: the data do not even cleanly tell you whether few or many predictors are right, because the posterior over model size is wide and dataset-dependent. The field thus held two incompatible intuitions at once. One said richer is safer; the other said simpler wins. Each was correct within the setting that produced it, and neither specified when that setting applied.

The symptom of this unsettled question is visible in practice, and it is close to chaotic. Central banks nowcasting ostensibly the same object, quarterly GDP, field flagship models whose indicator counts differ by a factor of four or more, from single-digit bridge systems to forty-series factor models, with no agreed principle explaining the spread. Table 1 makes the range concrete: the reference nowcasts we draw on run from eight predictors (Mexico) to thirty-six (the United States). Are the large systems over-engineered, or the small ones under-powered? The literature offers method in abundance but no rule for the one choice every practitioner must make (how many indicators to use), and so the choice is made by convention, by data availability, or by whatever the last paper did. This is the gap the present paper targets. Many estimators already exist; what is missing is the *comparative static* that says how large the indicator set should be, and why it differs across economies.

Our claim is that the two camps are both right, at opposite ends of a single gradient. The maximalist doctrine is the correct limit for economies with long samples, many weakly-correlated sectors, and stable dynamics, the high- N , high- r , low-volatility corner where the advanced-economy literature was built. The parsimony lesson is the correct limit for the opposite corner (short samples, few orthogonal signals, high volatility), where younger and less-diversified economies live. What looks like a contradiction is a comparative static: the optimal indicator count is a function of an economy’s informational maturity rather than a constant to be argued over, rising smoothly from one regime to the other. The next section states that function; the rest of the paper estimates where twelve economies fall along it.

3 A theory of indicator parsimony and informational maturity

Let a nowcast use k indicators drawn from a candidate set, and write the expected out-of-sample mean-squared forecast error as the sum of a squared-bias and an estimation-variance term,

$$\text{MSFE}(k) \approx B^2(k) + \kappa^2 \frac{k}{N}, \quad (1)$$

where N is the effective training length and κ^2 scales the per-parameter estimation noise. The bias term $B^2(k)$ is non-increasing in k : additional indicators can only reduce approximation error. The variance term is increasing in k : each retained indicator must be paid for with an estimated coefficient. If the candidate indicators share an approximate factor structure with r dominant common factors, $B^2(k)$ falls steeply for $k \leq r$ and flattens thereafter, as indicators beyond the r -th are largely collinear with those already included and add little orthogonal signal. Equation (1) is therefore U-shaped in k , with an interior optimum k^* where the marginal reduction in bias equals the marginal estimation cost κ^2/N .

Proposition. *The optimal indicator count k^* is (i) non-decreasing in the training length N ; (ii) non-decreasing in the number of orthogonal signal dimensions r ; and (iii) non-increasing in the per-parameter noise scale κ^2 .*

The three comparative statics map onto the stage of statistical and economic development. An emerging economy has a short reliable data history (small N), a narrow and partially measured formal sector whose indicators co-move strongly with one another and with output (small r), and elevated macroeconomic volatility from structural change and incomplete stabilisation (large κ^2). All three push k^* downward. A mature economy occupies the opposite corner (long samples, many well-measured and loosely correlated sectors, and a stable data-generating process), where k^* is large and the marginal indicator continues to help. The large-dataset nowcasting tradition (Stock and Watson, 2002; Bańbura et al., 2013), developed almost entirely on high- N , high- r economies, is the limiting case of this relationship rather than a universal law, and the recurrent “less-is-more” regularity is the same U-shape seen from the small- N side.

Of the three comparative statics, N is the one our design deliberately *holds fixed*: by estimating every economy over the same 2000–2026 window we shut down the sample-length channel, so that any remaining cross-country variation in k^* must come from r (signal dimensionality) and κ^2 (volatility). Both are properties of economic structure (how many distinct sectors an economy has, and how turbulent its output is), and both are summarised, empirically, by how diversified and complex the economy is. This is why the theory predicts

that k^* should track a measure of economic *complexity* rather than a measure of *wealth*: two economies at the same income can differ greatly in r (a diversified manufacturer versus a single-commodity exporter), and it is r , not income, that Equation (1) says governs the optimum. Section 5 tests exactly this.

4 Data and method

4.1 Precedent-best indicators

The threat to any cross-country comparison of optimal indicator counts is that the “optimum” merely reflects which series the researcher happened to assemble. We neutralise this in two steps. First, we do not scrape: every country faces the same compact, harmonised menu of headline indicators, one series for each standard nowcasting block (real activity, prices, the labour market, trade, a financial and monetary tier, business and consumer surveys, and, for open economies, an external or global block). Because the menu is common, cross-country differences in the estimated optimum cannot come from country-specific indicator selection.

Second, we let published flagship nowcasts inform which blocks each country carries, so the menu is precedent-informed rather than arbitrary. We read each country’s central-bank or academic nowcast of record, namely the New York Fed’s factor model for the United States (Bok et al., 2018), the ifo Institute’s nowcast and the Bundesbank factor model for Germany (ifo Institute, 2024; Schumacher and Breitung, 2008), the Bank of England’s release-augmented model for the United Kingdom (Anesti et al., 2022), and analogously for the other nine (Table 1; full sourcing in the Supplement). We use these nowcasts to fix the country-specific features of the menu: the national survey each economy relies on (the ifo Institute business climate survey and the ZEW (Leibniz Centre for European Economic Research) sentiment indicator for Germany, the Bank of Japan Tankan survey for Japan, and the National Institute of Economic Research (NIER) tendency survey for Sweden), whether an external block belongs (weighted heavily in Canada’s flagship, absent for the United States, which is itself the world driver), and the target definition (Mainland GDP for Norway). Indicators whose short samples would truncate the common window are dropped (Section 4). The candidate pool per country is therefore modest (10–15 series), common in its backbone and country-appropriate at the margin. We do not reproduce each flagship’s full variable list, which runs from eight series to several thousand and is not comparable across institutions; we take the block structure that guides these systems, not their size. This has a cost we state plainly: a few flagships make deliberate exclusions our common menu does not follow (Norway’s and Mexico’s models carry no price or financial series; Japan’s and

Indonesia’s exclude the CPI), and we retain those blocks for cross-country comparability rather than match any single flagship exactly.

4.2 Panel, window, and transformations

The panel covers twelve economies: seven advanced (United States, Germany, United Kingdom, Japan, Canada, Sweden, Norway) and five emerging (Brazil, Mexico, South Africa, Indonesia, Turkey). The target is quarterly real GDP growth (for Norway, Mainland GDP, which excludes petroleum extraction and shipping). Indicators are aggregated to quarterly frequency and enter as year-on-year growth rates for series with seasonal structure (industrial production, retail sales, trade, CPI) and as levels or first differences for rates and spreads. The common estimation window is 2000Q1–2026Q1, giving N between 87 and 103 usable quarters per country after transformation, long enough that the sample-length channel is not binding, and equalised across countries by construction. Two housekeeping rules keep the window clean: scattered interior gaps in step-function series (policy rates, ragged releases) are forward-filled using past-information only, and any candidate indicator whose late start would truncate a country below 80 usable quarters is dropped rather than allowed to shorten the sample (this removes, e.g., Indonesian retail sales, which begin only in 2012). All data are drawn from publicly redistributable sources (FRED, OECD, Eurostat, IMF, BIS, and national statistical offices and central banks); the Supplement documents every series, source, and transformation.

4.3 Estimator and the optimum

For each country we estimate the parsimony curve with a single, uniform procedure. Indicators are first ordered by the sequence in which they enter a LASSO path (Tibshirani, 1996) fitted on the residuals of an AR(2) baseline, so that the most informative series enter first. We then walk k from 0 (the AR(2) baseline) upward, and at each k compute the pseudo-real-time mean-squared forecast error over an expanding window: at each quarter the model is re-estimated on data through the prior quarter and used to predict the current one, with no full-sample standardisation and hence no look-ahead. This yields, for every country, an out-of-sample error curve $\text{RMSE}(k)$ whose shape and minimiser we read directly.

We report two summaries of each curve. The *argmin* is the indicator count that minimises the out-of-sample error, our primary estimate of k^* . The *one-standard-error* count is the smallest k within one standard error of the minimum, a deliberately conservative selector. We lead with the argmin because the one-standard-error rule compresses almost every country to $k = 1$ – 2 and so discards the cross-country variation that is the object of study; the argmin

retains it. Two caveats attach to the level of the argmin and we state them plainly. It is bounded above by the size of the (curated, 10–15 series) candidate pool, so it cannot and should not be read as commensurate with a published flagship count that draws on hundreds of series; and its absolute value is less robust than the *ordering* of arguments across countries and the *shape* of the curve (a sharp upturn, where over-inclusion is penalised, versus a flat basin). Accordingly our inference rests on cross-country comparisons (the emerging-versus-advanced gap, the rank correlation with published counts, and the complexity gradient), not on the absolute number for any single country.

5 Results

5.1 The parsimony–maturity gradient

Table 1 reports, for each of the twelve economies, the sample length N , the one-standard-error count, and the argmin. The pattern is systematic. Among advanced economies the argmin averages 6.4 and ranges from 2 (Norway) to 10 (United Kingdom); among emerging economies it averages 3.2 and never exceeds 4. The difference in argmin between the two groups is significant despite the small panel (two-sample t -test, $p = 0.026$). The estimated argmin also tracks, in rough order, the size of the nowcasting systems these countries’ central banks and statistical agencies field, the advanced economies running larger flagship models and the emerging economies more compact ones. We report that alignment only informally. Published predictor counts are not measured on a common basis across institutions: some flagship models are curated factor systems of a few dozen series, some are big-data systems that select from thousands, and some are model suites, so a formal correlation across them would compare unlike objects. The reproducible statement is the argmin gradient itself.

The *shape* of the curves reinforces the level result. In emerging economies the error curve turns up sharply once a few indicators are in. This over-inclusion penalty is the signature of a small- r , large- κ^2 economy, in which extra indicators buy variance without bias reduction. In advanced economies the curve settles into a flat basin, and adding indicators past the optimum costs little, because the richer factor structure keeps supplying weakly correlated signal. Norway is the instructive advanced-economy exception. Its curve behaves like an emerging economy’s, reaching its minimum at $k = 2$ and rising thereafter, and we return to it below because it is the case that separates our explanation from the obvious alternative.

Table 1: The parsimony–maturity gradient: estimated optimal indicator count (argmin of the out-of-sample error curve on each country’s precedent-best set, 2000Q1–2026Q1).

Country	Group	N	k^* (1-SE)	argmin
United Kingdom	Advanced	99	3	10
Canada	Advanced	99	2	8
Germany	Advanced	99	2	8
United States	Advanced	103	3	6
Japan	Advanced	103	1	6
Sweden	Advanced	96	5	5
Norway	Advanced	96	2	2
Indonesia	Emerging	89	2	4
Mexico	Emerging	96	1	4
South Africa	Emerging	99	2	4
Brazil	Emerging	96	1	2
Turkey	Emerging	87	1	2
<i>Advanced mean</i>				6.4
<i>Emerging mean</i>				3.2

Notes: argmin is the indicator count minimising the expanding-window out-of-sample RMSE on each country’s precedent-best candidate set; $k^*(1\text{-SE})$ is the smallest count within one standard error of that minimum. Advanced-vs-emerging difference in argmin: two-sample t -test $p = 0.026$. The published predictor counts of each country’s reference nowcast are listed in the Supplement and discussed informally, as they are not measured on a common basis across institutions.

5.2 Complexity, not wealth

Why does the optimum sit where it does? The theory says k^* should track the number of orthogonal signal dimensions r (a structural property, well proxied by how diversified and complex an economy’s production is) rather than how rich the economy is. We test this directly by correlating the estimated argmin with two anchors: the logarithm of GDP per capita (a wealth measure) and the Economic Complexity Index of Hidalgo and Hausmann (2009) (a measure of the diversity and sophistication of the export basket). Table 2 reports the result. The argmin correlates with complexity at Spearman $+0.61$ ($p = 0.035$) but with income at only $+0.35$ ($p = 0.27$, not significant); in a linear fit, $\text{argmin} = 3.53 + 1.74 \text{ ECI}$ ($R^2 = 0.33$), whereas income explains barely a tenth of the variation. Informational maturity, as the theory predicts, is about structural complexity rather than about being rich.

5.3 Norway: the case that identifies the mechanism

Income and complexity are positively correlated across most economies, so a single case in which they diverge does a great deal of identifying work. Norway is that case. On GDP

Table 2: What explains the optimum: economic complexity versus income.

Anchor (correlated with estimated argmin)	Spearman ρ	p
Economic Complexity Index (ECI)	+0.61	0.035
log GDP per capita	+0.35	0.272

Notes: $n = 12$ economies. ECI from the Harvard Growth Lab Atlas of Economic Complexity; GDP per capita (PPP, constant international dollars) from the World Bank. Linear fits: $\text{argmin} = 3.53 + 1.74 \text{ ECI}$ ($R^2 = 0.33$, $p = 0.050$); $\text{argmin} = -8.72 + 1.30 \log(\text{GDPpc})$ ($R^2 = 0.12$, $p = 0.272$). ECI and GDP-per-capita values are entered from published sources and reported in the Supplement.

per capita it is the *richest* economy in the panel (about \$114,000, PPP); on the Economic Complexity Index it is one of the *least* complex of the advanced group ($\text{ECI} \approx 0.36$, below Mexico), because its exports are dominated by petroleum and its Mainland economy by a handful of large sectors. Wealth and complexity therefore make opposite predictions for Norway’s optimum: a wealth story predicts a large k^* near those of the United States or Germany, while a complexity story predicts a small one near the emerging cluster. The data are unambiguous. Norway’s argmin is 2, tied for the lowest in the entire panel, and its error curve turns up like an emerging economy’s. The income-based prediction misses Norway by the widest margin of any country; the complexity-based prediction places it correctly. It is this divergence, more than any average correlation, that warrants reading the gradient as a gradient of informational *maturity* (structural signal dimensionality) rather than of development or income.

5.4 Reading the gradient: what puts each economy where it is

The gradient is not a statistical accident; each country’s position follows from what its economy is made of. Norway, discussed above, is the clearest case, but the same logic reads across the panel and is the substance of what we mean by “knowing the economy.”

At the high- k^* end sit the deep, diversified producers. Germany’s output spans motor vehicles, machinery, chemicals, and electrical equipment, each with a partly independent cycle and each separately surveyed (ifo, ZEW); Japan’s export basket is the most complex in the panel (precision machinery, electronics, vehicles), consistent with its top ECI; Sweden pairs a diversified small-open industrial base (vehicles, telecom equipment, forestry, pharmaceuticals) with a rich survey apparatus (NIER); and the United Kingdom adds a large, distinct services and financial sector on top of a diversified goods economy. These economies generate many weakly-correlated signals (high r), so their flagship nowcasts keep adding indicators usefully (published counts 25–36) and our argmin runs high. The United

States belongs here too: its curated candidate pool caps the measured argmin at six, but its thirty-six-series published nowcast reveals the true ceiling.

At the low- k^* end sit the concentrated economies, and they arrive there by two distinct routes. The first is *resource concentration*: Norway (petroleum), Brazil (soybeans, iron ore, crude), South Africa (mining), and Indonesia (coal, palm oil) each key their activity to one or two global commodity cycles, so their domestic indicators (output, exports, the currency, fiscal revenue) move together and carry few independent signals. Low r , low k^* , regardless of income: Norway is rich and Indonesia is not, yet both sit near the bottom. The second route is *volatility*: Turkey has a moderately complex manufacturing base, but recurrent currency and inflation crises inflate κ^2 , so the estimation-variance term dominates early and the optimum collapses to two, through the volatility channel of Equation (1) rather than the signal channel.

Two economies are instructive intermediate cases. Mexico’s manufacturing is genuinely sophisticated, but its cycle is largely a satellite of United States industrial production (vehicles and electronics assembled for the US market), so few *independent* domestic signals survive once US activity is in the model, a compact optimum despite real complexity. Canada is the panel’s most honest tension: resource-weighted and only moderately complex on the index, yet its argmin is high, because a diversified manufacturing base, dense and well-measured data, and strong US spillovers still supply many signals. Canada and the United Kingdom are the two economies our complexity anchor fits least well (Section 5), and we flag them rather than bury them: the gradient is a central tendency, not a deterministic law, and the residuals are themselves informative about how signal dimensionality and raw complexity can come apart.

6 Discussion

What the panel establishes, and what it does not. With twelve economies our inference is a cross-section of twelve, and we are deliberate about which claims that supports. The emerging-versus-advanced gap in the optimum ($p = 0.026$) and the complexity-over-wealth result (+0.61 versus +0.35) are cross-country statements that a twelve-economy panel can carry; the absolute level of any single country’s argmin, bounded as it is by a curated candidate pool, is not something we lean on. The fixed 2000–2026 window removes sample length as a confound by construction, so the surviving variation is structural, which is why complexity rather than income does the explanatory work.

Identification. The mechanism is identified off structure rather than income, and Norway is the central case: a rich but resource-concentrated economy whose optimum sits with the emerging group, exactly where complexity predicts and income does not. The commodity emerging economies (Brazil, South Africa, Indonesia) reinforce the same point from the other side: low complexity, low optimum, regardless of income rank.

Robustness to a fully real-time protocol. The baseline computes the LASSO variable ordering and the feature standardisation once on the full sample; these fix only the order and scale of the candidates, while the forecast *scoring* is already strictly out-of-sample. To confirm the gradient is not an artefact of that choice, we re-ran the entire panel under a fully real-time variant that re-ranks and re-standardises *within each expanding-training window*, so no future information touches any step. The gradient is qualitatively robust: the advanced optimum remains well above the emerging one (real-time means 8.1 versus 5.0), and the per-country argmins agree closely with the baseline (Spearman +0.77, $p = 0.004$). The strict two-group significance does weaken (the advanced-versus-emerging t -test rises to $p = 0.13$) because fold-by-fold re-ranking injects extra estimation noise into an already small twelve-country cross-section; we report this openly. The qualitative gradient (direction, rank, and complexity association) is thus not an artefact of full-sample ordering, while the precise significance of the mean gap is sensitive to protocol at $n = 12$, a sensitivity the larger-panel and within-country extensions below are designed to resolve. The per-country real-time optima are tabulated in the Supplementary Material.

Limitations and next steps. Four are worth naming. The panel is small, and extending it (Korea, Australia, Poland, and other economies with documented flagship nowcasts) would sharpen the complexity gradient and stabilise the significance just discussed. The argmin is a compressed, candidate-pool-bounded estimate of k^* ; a within-country design that traces k^* as each economy’s sample lengthens would test comparative static (i), that k^* rises with N , directly, and would convert the cross-section into a panel with far more observations. A direct per-country factor count would corroborate the Economic Complexity Index as the proxy for r , but is not identified on the narrow precedent-best panels: the Bai–Ng (Bai and Ng, 2002) information criterion saturates at its maximum and the eigenvalue-ratio estimator (Ahn and Horenstein, 2013) is unstable on ten-to-fifteen series, so we rely on the ECI as the structural proxy for r and leave a directly estimated factor count to the wider panel. And the ECI and GDP-per-capita anchors are entered from published sources; their rankings are well established but their exact values should be re-verified against source for the final replication archive. None of these disturbs the central result: the optimal number of indicators in a GDP

nowcast is a comparative static of informational maturity rather than a constant, smaller for younger and less-diversified economies, and governed by economic complexity rather than wealth.

Data and replication. The panel is assembled entirely from publicly redistributable sources and the estimation code applies one uniform procedure to every country; the Supplement documents each series and source, and the replication archive will be deposited on acceptance.

References

- Seung C. Ahn and Alex R. Horenstein. Eigenvalue ratio test for the number of factors. *Econometrica*, 81(3):1203–1227, 2013.
- Nikoleta Anesti, Ana Beatriz Galvão, and Silvia Miranda-Agrippino. Uncertain kingdom: Nowcasting GDP and its revisions. *Journal of Applied Econometrics*, 37(1):42–62, 2022.
- Marta Bańbura, Domenico Giannone, Michele Modugno, and Lucrezia Reichlin. Now-casting and the real-time data flow. In *Handbook of Economic Forecasting, Volume 2A*, pages 195–237. Elsevier, 2013.
- Jushan Bai and Serena Ng. Determining the number of factors in approximate factor models. *Econometrica*, 70(1):191–221, 2002.
- Jean Boivin and Serena Ng. Are more data always better for factor analysis? *Journal of Econometrics*, 132(1):169–194, 2006.
- Brandyn Bok, Daniele Caratelli, Domenico Giannone, Argia M. Sbordon, and Andrea Tambalotti. Macroeconomic nowcasting and forecasting with big data. *Annual Review of Economics*, 10:615–643, 2018.
- Stuart Geman, Elie Bienenstock, and René Doursat. Neural networks and the bias/variance dilemma. *Neural Computation*, 4(1):1–58, 1992.
- Domenico Giannone, Lucrezia Reichlin, and David Small. Nowcasting: The real-time informational content of macroeconomic data. *Journal of Monetary Economics*, 55(4):665–676, 2008.
- Domenico Giannone, Michele Lenza, and Giorgio E. Primiceri. Economic predictions with big data: The illusion of sparsity. *Econometrica*, 89(5):2409–2437, 2021.

- Gerd Gigerenzer and Henry Brighton. Homo heuristicus: Why biased minds make better inferences. *Topics in Cognitive Science*, 1(1):107–143, 2009.
- César A. Hidalgo and Ricardo Hausmann. The building blocks of economic complexity. *Proceedings of the National Academy of Sciences*, 106(26):10570–10575, 2009.
- ifo Institute. ifoCAST: Nowcasting German GDP growth. <https://www.ifo.de/en/ifoCAST>, 2024. Operational nowcast; the ifo Institute reports monitoring on the order of 300 series grouped into seven blocks; the exact indicator count is not publicly documented.
- Christian Schumacher and Jörg Breitung. Real-time forecasting of German GDP based on a large factor model with monthly and quarterly data. *International Journal of Forecasting*, 24(3):386–398, 2008.
- Claude E. Shannon. A mathematical theory of communication. *Bell System Technical Journal*, 27(3):379–423, 1948.
- Jeremy Smith and Kenneth F. Wallis. A simple explanation of the forecast combination puzzle. *Oxford Bulletin of Economics and Statistics*, 71(3):331–355, 2009.
- James H. Stock and Mark W. Watson. Macroeconomic forecasting using diffusion indexes. *Journal of Business & Economic Statistics*, 20(2):147–162, 2002.
- James H. Stock and Mark W. Watson. Combination forecasts of output growth in a seven-country data set. *Journal of Forecasting*, 23(6):405–430, 2004.
- Robert Tibshirani. Regression shrinkage and selection via the lasso. *Journal of the Royal Statistical Society, Series B*, 58(1):267–288, 1996.
- David H. Wolpert. The lack of a priori distinctions between learning algorithms. *Neural Computation*, 8(7):1341–1390, 1996.

Supplementary Material

Data Provenance, Construction, and Robustness
for *Know the Economy, Not the Data Lake: Country Structure and the
Right Number of Indicators in a GDP Nowcast*

This Supplement documents the construction of the twelve-country panel used in the main paper: the data channels and their licensing status (§1); the country-by-country flagship nowcast from which each precedent-best indicator set is inherited, with the exact indicators used (§2); the per-indicator source, native frequency, transformation, and publication lag (§3); the per-country source institutions and binding sample history (§4); the two complexity/wealth anchors and their sources (§5); and the sample-construction rules and estimator settings (§7). The guiding principle throughout is that every series used in estimation is drawn from a *publicly redistributable* source, so that the replication archive reproduces all tables and figures without proprietary data.

Contents

1	Data channels and licensing	1
2	Precedent-best indicator sets by country	2
3	Sources by indicator	4
4	Sources and binding history by country	5
5	Complexity and wealth anchors	6
6	Real-time robustness (train-only ordering and standardisation)	6
7	Sample construction and estimator settings	7

1 Data channels and licensing

The panel is built backbone-first from free, redistributable providers, with proprietary terminals used only to cross-check and never as the source of record.

- **Channel 1: public / redistributable (the estimation backbone).**

- *FRED* (Federal Reserve Bank of St. Louis): GDP, trade, FX, CPI, rates, unemployment; scriptable via the FRED API. Covers the majority of the panel.
 - *DBnomics*: free aggregator of FRED, OECD, Eurostat, IMF, BIS and national providers under one Python API; used to fill series that FRED froze or never carried (current industrial production, retail sales, and CPI via OECD/Eurostat; Banco Central do Brasil SGS; etc.). Points to public providers, so redistribution-friendly.
 - *OECD Data Explorer* (SDMX): current vintages of the Production & Sales, Prices, and Key Short-Term Indicators datasets, including the OECD-mirror series FRED stopped updating around 2023–2025.
 - *IMF* (IFS, DOTS) and *BIS*: monetary aggregates, credit-to-the-private- non-financial-sector, cross-country policy rates, NEER, and emerging-market trade/FX.
 - *National statistical offices and central banks*, the residual gaps: e.g. BPS and Bank Indonesia (Indonesia), IBGE and BCB (Brazil), INEGI and Banxico (Mexico), StatsSA and SARB (South Africa), TCMB (Turkey), SSB and Norges Bank (Norway; Mainland GDP), Statistics Canada and the Bank of Canada, Destatis/ONS/METI/SCB for the European and Japanese activity series.
- **Channel 2: proprietary terminals (cross-check only).** Bloomberg and Refinitiv Datas-
stream were used to assemble and verify the financial-market tier (equity indices, government
yields, FX) and to spot-check ragged emerging-market releases. Under their licences these
data *cannot* be redistributed; every series retained from this channel is re-sourced to a public
equivalent (FRED/OECD/IMF/BIS/national) for the replication archive.

2 Precedent-best indicator sets by country

Every country faces the same harmonised menu of headline indicators, one series per standard nowcasting block, rather than a country-specific scrape (main paper, §3.1). Published flagship nowcasts inform which blocks each country carries (the national survey it relies on, whether an external or global block belongs, and the target definition), but we do not reproduce a flagship’s full variable list, which is not comparable in size across institutions. Table 1 lists, for each economy, the reference nowcast, its published predictor count (for context only; see the caveat below), and the harmonised candidate indicators we estimate on. “US/global” denotes the shared external block (US industrial production, WTI oil, a global commodity index, VIX).

On the published counts. These span two orders of magnitude, from eight series in Mexico’s bridge model to over two thousand candidate series in Brazil’s LASSO system, and are not measured on a common basis: some are curated factor panels, some big-data selection systems, some model suites. We therefore report them for context only and do not use them as a comparison target. Two counts warrant a note. Germany’s operational ifoCAST does not publish an exact indicator count

(the ifo Institute reports monitoring on the order of 300 series and grouping the nowcast inputs into seven blocks); the academic Bundesbank benchmark of Schumacher and Breitung (2008) uses 52 series. We list the ifoCAST block figure and treat the precise German count as unresolved. Several flagships also make deliberate exclusions our common menu does not follow (Norway’s and Mexico’s models carry no price or financial series, and Japan’s and Indonesia’s exclude the CPI), and we retain those blocks for cross-country comparability at the cost of not matching any single flagship exactly.

Table 1: Precedent-best indicator sets and their sources.

Country	Reference flagship nowcast	Pub. k	Inherited candidate indicators
United States	Bok et al. (2018), NY Fed 4-block DFM	36	Industrial production, retail sales, unemployment, CPI, exports, imports, business confidence, consumer confidence, policy rate, 10y yield, WTI oil. (No global block; the US <i>is</i> the world driver.)
Germany	ifo Institute (2024) (ifo Institute); Schumacher and Breitung (2008), Bundesbank large factor model	21 / 52	IP, retail sales, unemployment, CPI, exports, imports, business (ifo) & consumer confidence, equity index, 10y yield, policy rate, FX; US IP.
United Kingdom	Anesti et al. (2022), BoE release-augmented DFM	33	IP, retail sales, unemployment, CPI, exports, imports, business & consumer confidence, equity index, 10y yield, policy rate, FX; US IP.
Japan	Hayashi and Tachi (2022), mixed-frequency DFM	19	IP, retail sales, unemployment, CPI, exports, imports, business (Tankan) & consumer confidence; US IP, WTI oil.
Canada	Chernis and Sekkel (2017), BoC 2-block DFM	23	IP, retail sales, unemployment, CPI, exports, imports, FX, business confidence, equity index, 10y yield, policy rate; US IP, WTI oil, commodity index, VIX.
Sweden	Karlsson et al. (2025), Örebro Univ., indicator selection for Swedish GDP	35	IP, retail sales, unemployment, CPI, exports, imports, business (NIER) & consumer confidence, equity index, 10y yield, policy rate, FX; US IP.
Norway	Luciani and Ricci (2014), Bayesian DFM (Mainland GDP); Aastveit and Trovik (2012)	14	IP, retail sales, unemployment, CPI, exports, imports, business confidence, equity index, 10y yield, policy rate, FX; WTI oil, US IP. Target = Mainland GDP .
Indonesia	Luciani et al. (2018), DFM	10	IP, CPI, exports, imports, FX, business & consumer confidence, policy rate, 10y yield, equity index; commodity index, WTI oil. (Retail sales dropped; starts 2012.)
Mexico	Gálvez-Soriano (2018), Banxico bridge equations	8	IP, retail sales, CPI, exports, imports, FX, business & consumer confidence, policy rate, 10y yield, equity index; US IP, WTI oil.

Table 1, continued

Country	Reference flagship nowcast	Pub. k	Inherited candidate indicators
South Africa	Botha et al. (2021), SARB model suite	25	IP, retail sales, CPI, exports, imports, FX, business & consumer confidence, policy rate, 10y yield, equity index; commodity index.
Brazil	Gonçalves (2022), BCB, electronic-payments + LASSO	68	IP, retail sales, CPI, exports, imports, FX, business & consumer confidence, policy rate, equity index; commodity index, WTI oil. (10y yield dropped; starts 2007.)
Turkey	Modugno et al. (2016) (Fed FEDS); Gül and Kazdal (2021) (TCMB), factor-MIDAS	14	IP, CPI, exports, imports, FX, business & consumer confidence, equity index; commodity index, WTI oil. (10y yield and retail sales dropped; start too late.)

3 Sources by indicator

Table 2 gives the primary public route, native frequency, model transformation, and typical publication lag for each indicator type, applied consistently across the twelve economies. Where FRED’s mirror was frozen, the current vintage is taken from OECD or the national office via DBnomics.

Table 2: Per-indicator source, frequency, transformation, and lag.

Indicator	Primary public source (route)	Freq.	Transform	Lag (d)
Quarterly real GDP (target)	National NSO; OECD QNA; FRED	Q	YoY %	45–70
Industrial production	FRED / OECD MEI (Production of Total Industry)	M	YoY %	30–45
Retail sales	FRED / OECD MEI; national	M	YoY %	35–50
CPI (headline)	FRED / OECD Prices; national	M	YoY %	10–20
Exports (goods)	IMF DOTS / OECD MEI / FRED	M	YoY %	30–50
Imports (goods)	IMF DOTS / OECD MEI / FRED	M	YoY %	30–50
Unemployment rate	FRED / OECD MEI / ILO	M	level	20–45
Bank credit (private)	BIS credit series; IMF IFS	M/Q	YoY %	30–45
Policy rate	BIS policy rates; central bank	M	level	0
10-year govt. yield	FRED / OECD MEI / national	M	level	0
Equity index	Exchange; public index history	M	YoY %	0
Exchange rate (USD)	FRED; BIS	M	YoY %	0
Business confidence	OECD BCI; national (ifo, Tankan, NIER)	M	level	20–30
Consumer confidence	OECD CCI; national	M	level	20–30
US industrial production (global)	FRED	M	YoY %	30–45
WTI crude oil (global)	FRED / EIA	M	YoY %	0–1

Indicator	Primary public source (route)	Freq.	Transform	Lag (d)
Commodity price index (global)	IMF / public index	M	YoY %	0–5
VIX (global)	Cboe / FRED	M	level	0

4 Sources and binding history by country

Table 3 records the national source institutions, the usable sample length N entering estimation (after transformation and the 80-quarter rule of §7), and country-specific notes. The common estimation window is 2000Q1–2026Q1; N varies slightly across countries because of transformation losses and dropped late-starting series.

Table 3: Per-country source institutions, usable N , and notes.

Country	National source institutions	N	Notes
United States	BEA, BLS, Federal Reserve, Census	103	Richest data; no global block (US is the driver).
Germany	Destatis, Bundesbank, ifo, ZEW	99	The ifo Institute business climate survey and the ZEW (Leibniz Centre for European Economic Research) sentiment indicator are the flagship surveys.
United Kingdom	ONS, Bank of England	99	Release-augmented flagship; CIPS/PMI available.
Japan	Cabinet Office, METI, BOJ	103	Bank of Japan Tankan survey replaces PMI; IIP + trade core.
Canada	Statistics Canada, Bank of Canada	99	High foreign-signal share; US/global block added.
Sweden	SCB, Riksbank, NIER	96	National Institute of Economic Research (NIER) tendency survey; SEK crosses.
Norway	SSB, Norges Bank	96	Target = Mainland GDP; Brent control.
Brazil	IBGE, Banco Central do Brasil	96	BCB-SGS API; FGV confidence. 10y yield dropped (starts 2007).
Mexico	INEGI, Banxico	96	IGAE activity; strong US-IP linkage.
South Africa	StatsSA, SARB	99	Mining production; commodity price index; 2016 break.
Indonesia	BPS, Bank Indonesia	89	IP relatively weak; retail sales dropped (starts 2012).
Turkey	TurkStat, TCMB	87	Capacity utilisation; FX-regime breaks; 10y yield & retail dropped.

5 Complexity and wealth anchors

The two anchors used in the complexity-versus-wealth test (main paper, §4.2) are reported in Table 4. The Economic Complexity Index (ECI) is from the Harvard Growth Lab’s *Atlas of Economic Complexity* (export-basket sophistication, most recent available vintage ≈ 2022). GDP per capita is the World Bank’s PPP series in constant international dollars (≈ 2023). Values are entered from these published sources; their cross-country *rankings* are well established and are what the rank correlations use, but the exact figures should be re-verified against source for the final replication archive.

Table 4: Economic Complexity Index and GDP per capita, by economy (sorted by ECI).

Economy	Estimated argmin	ECI (Harvard)	GDP p.c. PPP (\$)
Japan	6	+2.27	45,000
Germany	8	+1.95	63,000
Sweden	5	+1.62	62,000
United States	6	+1.52	74,000
United Kingdom	10	+1.42	54,000
Mexico	4	+1.06	22,000
Canada	8	+0.52	57,000
Turkey	2	+0.52	37,000
Norway	2	+0.36	114,000
Brazil	2	−0.05	18,000
South Africa	4	−0.10	14,500
Indonesia	4	−0.35	14,000

Notes: Norway is the identifying case, the highest GDP per capita in the panel but a low ECI (below Mexico), and an estimated argmin of 2 that sits with the emerging cluster, matching the complexity anchor and contradicting the wealth anchor. Spearman correlations with the estimated argmin: ECI +0.61 ($p = 0.035$); log GDP per capita +0.35 ($p = 0.27$).

6 Real-time robustness (train-only ordering and standardisation)

The baseline estimator (§7) fixes the LASSO variable ordering and the feature standardisation once on the full sample, while scoring strictly out-of-sample. Table 5 reports a fully real-time re-run in which both the ranking and the standardisation are recomputed *inside each expanding-training window*, so no future information enters any step. The estimated optima rise (real-time re-ranking is noisier, so it keeps more indicators) but the cross-country pattern is preserved: advanced economies remain well above emerging (means 8.1 versus 5.0), the real-time and baseline argmins are strongly rank-correlated (Spearman +0.77, $p = 0.004$), and the ranking still tracks published flagship counts (Spearman +0.53, $p = 0.08$). The advanced-versus-emerging t -test rises to $p = 0.13$: with twelve economies and fold-by-fold re-ranking, the strict significance of the mean gap is protocol-sensitive, even though its direction and rank are not.

Table 5: Real-time robustness: baseline vs. fully real-time argmin (train-only ordering and standardisation).

Country	Group	Baseline argmin	Real-time argmin	Published k
United Kingdom	Advanced	10	11	33
Canada	Advanced	8	11	23
Germany	Advanced	8	11	26
United States	Advanced	6	11	36
Japan	Advanced	6	3	25
Sweden	Advanced	5	6	35
Norway	Advanced	2	4	14
Indonesia	Emerging	4	4	10
Mexico	Emerging	4	9	8
South Africa	Emerging	4	7	18
Brazil	Emerging	2	2	10
Turkey	Emerging	2	3	14
<i>Advanced mean</i>		6.4	8.1	27
<i>Emerging mean</i>		3.2	5.0	12

Notes: Real-time variant re-ranks (LASSO) and re-standardises within each expanding-training window. Spearman(real-time, baseline) = +0.77 ($p = 0.004$); Spearman(real-time, published k) = +0.53 ($p = 0.08$); advanced-vs-emerging t -test $p = 0.13$.

7 Sample construction and estimator settings

Transformations. Series with seasonal structure (industrial production, retail sales, trade, CPI, credit, equity indices, FX) enter as year-on-year growth rates, which remove seasonality by construction; rates and spreads (policy rate, 10-year yield, unemployment, confidence indices, VIX) enter in levels or first differences. Two lags of GDP growth form the AR(2) baseline against which all indicators are evaluated.

Window and gaps. The common estimation window is 2000Q1–2026Q1. Scattered interior gaps in step-function or ragged-release series are forward-filled using past information only (no look-ahead); only leading missing rows are dropped. Any candidate indicator whose late start would truncate a country’s usable sample below 80 quarters is removed rather than allowed to shorten the window. This affects Indonesian retail sales (from 2012), the Brazilian 10-year yield (from 2007), and the Turkish 10-year yield and retail sales.

Estimator. For each country the candidate indicators are ordered by their entry sequence along a LASSO path (Tibshirani, 1996) fitted on AR(2) residuals. The parsimony curve $RMSE(k)$ is then computed by expanding-window pseudo-real-time evaluation: at each quarter the model is re-estimated on data through the prior quarter and used to predict the current quarter, with standardisation fit on the training portion only (no full-sample look-ahead). The *argmin* of this curve is the primary estimate of k^* ; the *one-standard-error* count (the smallest k within one standard

error of the minimum) is reported as a conservative alternative. The argmin is bounded above by the size of the curated candidate pool and is therefore not directly comparable in level to a published flagship count drawn from hundreds of series; the cross-country *ordering* of arguments and the *shape* of the curve (sharp upturn versus flat basin) carry the inference.

Reproducibility. All estimation data are from the Channel-1 public sources above; the panel-assembly and estimation scripts apply one uniform procedure to every country. The full replication archive (data-build code, estimation code, and the assembled panel) will be deposited in the journal’s data archive on acceptance.

References

- Knut Are Aastveit and Tørres Trovik. Nowcasting Norwegian GDP: The role of asset prices in a small open economy. Working Paper 2007/9, Norges Bank, 2012.
- Nikoleta Anesti, Ana Beatriz Galvão, and Silvia Miranda-Agrippino. Uncertain kingdom: Nowcasting GDP and its revisions. *Journal of Applied Econometrics*, 37(1):42–62, 2022.
- Brandyn Bok, Daniele Caratelli, Domenico Giannone, Argia M. Sbordone, and Andrea Tambalotti. Macroeconomic nowcasting and forecasting with big data. *Annual Review of Economics*, 10: 615–643, 2018.
- Byron Botha, Monique Reid, Tim Olds, Daan Steenkamp, and Rossouw van Jaarsveld. Nowcasting South African GDP using a suite of statistical models. Working Paper WP/21/01, South African Reserve Bank, 2021.
- Tony Chernis and Rodrigo Sekkel. A dynamic factor model for nowcasting canadian GDP growth. Staff Working Paper 2017-2, Bank of Canada, 2017.
- Oscar de Jesús Gálvez-Soriano. Nowcasting Mexican GDP using factor models and bridge equations. Working Paper 2018-06, Banco de México, 2018.
- Raquel Nadal Cesar Gonçalves. Nowcasting Brazilian GDP with electronic payments data. Working Paper 564, Banco Central do Brasil, 2022.
- Selçuk Gül and Abdullah Kazdal. Nowcasting and short-term forecasting Turkish GDP: A factor-MIDAS approach. Working Paper 21/11, Central Bank of the Republic of Turkey, 2021.
- Fumio Hayashi and Yuta Tachi. Nowcasting Japan’s GDP. Working paper, National Graduate Institute for Policy Studies (GRIPS), 2022.
- ifo Institute. ifoCAST: Nowcasting German GDP growth. <https://www.ifo.de/en/ifoCAST>, 2024. Operational nowcast; the ifo Institute reports monitoring on the order of 300 series grouped into seven blocks; the exact indicator count is not publicly documented.

- Sune Karlsson, Stepan Mazur, and Mariya Raftab. Identifying useful indicators for nowcasting GDP in Sweden. Working Paper 4/2025, Örebro University School of Business, 2025.
- Matteo Luciani and Lorenzo Ricci. Nowcasting Norway. *International Journal of Central Banking*, 10(4), 2014.
- Matteo Luciani, Madhavi Pundit, Arief Ramayandi, and Giovanni Veronese. Nowcasting Indonesia. *Empirical Economics*, 55(2):597–619, 2018.
- Michele Modugno, Barış Soybilgen, and Ege Yazgan. Nowcasting Turkish GDP and news decomposition. Finance and Economics Discussion Series (FEDS) 2016-044, Federal Reserve Board, 2016.
- Christian Schumacher and Jörg Breitung. Real-time forecasting of German GDP based on a large factor model with monthly and quarterly data. *International Journal of Forecasting*, 24(3): 386–398, 2008.
- Robert Tibshirani. Regression shrinkage and selection via the lasso. *Journal of the Royal Statistical Society, Series B*, 58(1):267–288, 1996.