

Mapping Structural Change in Manufacturing: Introducing the STiM Database

Yanis Bekhti[‡]

16th October 2025

[Latest version of the paper](#)  | [STiM Database repository](#) 

Abstract. This paper introduces the Structural Transformation in Manufacturing Database (STiM), which provides employment, value-added, and output data at the two-digit manufacturing sub-sectoral level. In its current version, the database covers 12 homogenised industrial activities across more than 145 countries from 1963 to 2019. It is the result of combining several data sources and undertaking substantial harmonisation efforts aimed at ensuring internal, intertemporal, and international consistency. In addition, the database links several price deflators from multiple institutions to adjust nominal measures, thereby producing estimates of real value-added and constant output. To date, this makes the STiM database the most comprehensive dataset for long-run analyses of structural change within manufacturing in both developed and developing economies. The final database and the replication package will be made freely available at the corresponding [GitHub repository](#) .

Keywords: Manufacturing, structural change, database, employment, value-added, deflators.

JEL Codes: O14, C82, N60.

1 Introduction

One of the earliest and perhaps most important tenets of the literature on economic convergence is that development entails structural change (Kuznets, 1973; Allen, 2011). In this setting, industrialisation has long been considered the cornerstone of this transformation, as the reallocation of production factors towards manufacturing has been associated with several direct and indirect gains (Prebisch, 1950; Lewis, 1954; Hirschman, 1958; Kaldor, 1966; Szirmai, 2012). Yet, recent empirical evidence discussing the structural transformation of some developing countries has reshaped the debate on the role of manufacturing in economic development. In the aftermath of the Second World War, the tertiary sector in these economies has often expanded ahead of a sustained industrial base, while manufacturing has begun to decline at much lower income levels than in the past — a pattern now termed premature deindustrialisation (Palma, 2005;

[‡] Catholic University of the Sacred Heart of Milan, Department of Economic Policy.

[‡] Université Paris 1 - Panthéon-Sorbonne, CNRS, IRD, Centre d'Economie de la Sorbonne (CES).

• E-mails: yanis.bekhti@unicatt.it | yanis.bekhti@univ-paris1.fr

Dasgupta and Singh, 2007; Rodrik, 2016; Tregenna, 2016). Put differently, the traditional hump-shaped relationship between income per capita and manufacturing appears to have shifted downward and leftward over time, leading to projected decreases in industrial employment and value-added shares at earlier stages of development. This new time-dependent pattern of structural change hence suggests that industries have become a more challenging route to growth than in the past years, providing developing countries with a narrower window to reap the gains associated with industrialisation.

Given that these findings entail striking normative implications for development policies, a branch of new studies has recently emerged seeking to substantiate, validate, or mitigate these conclusions. Different perspectives have been chosen to contribute directly or indirectly to this ongoing debate, including discussions on whether this pattern still holds when examining the trajectories of new developing countries using exclusive datasets (Kruse et al., 2022; Ngumkeu and Zeufack, 2024; Lautier, 2024) or through reassessments of the contribution of manufacturing to productivity or economic growth (Szirmai and Verspagen, 2015; Cantore et al., 2017; Diao, McMillan et al., 2019; Forero and Tena-Junguito, 2024). Alternatively, further contributions have examined the factors driving successful industrialisation (Haraguchi, Martorano et al., 2019) and gauged the effects of significant post-1990s institutional and technological changes that may have either reduced manufacturing's competitiveness or undermined its historical role as an engine of growth (Baldwin, 2011; Autor et al., 2013; Bogliaccini, 2013).

Yet, so far, the main recurrent limitation of these analyses has been the trade-off between the level of aggregation and the number of countries covered. On the one hand, papers often rely on aggregate measures of industrialisation and deindustrialisation (e.g., manufacturing's share in GDP or employment), which provide broad cross-country coverage but mask heterogeneity across sub-sectors (Felipe et al., 2019). On the other hand, others use granular or firm-level data that provide valuable insights into within-manufacturing dynamics but confine the analysis to a small set of countries, thereby limiting the external validity of the findings (Diao, Ellis et al., 2025). This trade-off has made it difficult to reach generalisable conclusions, as one cannot infer generic properties of manufacturing from a few countries, nor assume that premature deindustrialisation is a universal feature of the sector. Instead, it may be a phenomenon confined to specific traditional activities that were more affected by the technological and institutional changes that have been ongoing since the 1990s. To address these limitations, a few studies have chosen to conduct their analysis using datasets aggregated at the manufacturing sub-sectoral level, as this approach uncovers within-manufacturing dynamics while covering a greater number of countries (Rodrik, 2013; Vaz and Baer, 2014; Andreoni and Tregenna, 2020; Dosi et al., 2021; Bekhti, 2025). Although this is likely the most appropriate level of aggregation for reassessing conclusions about manufacturing's role in recent years, only a few datasets were available until recently. In addition, these datasets were either methodologically weak or incomplete and lacked consistent panel structures, rendering them unsuitable for reliable analysis (Rodrik, 2013).

This paper addresses these methodological limitations by introducing the Structural Transformation in Manufacturing (STiM) Database. The STiM database synthesises, aggregates and harmonises data from major institutional organisations. Its aim is to provide a comprehensive and comparable unbalanced panel dataset on employment, value added, and output at the two-digit manufacturing sub-sectoral level, covering a wide set of countries over the longest time period possible. Constant output and real value-added measures are also provided, as the database links several price deflators from multiple institutions to adjust nominal values and thus avoid conflating prices with quantity. In its current version, the database covers 12 homogenised industrial activities across more than 145 countries from 1963 to 2019. To date, this makes the STiM database the most comprehensive dataset for long-run analyses of structural change within manufacturing in both developed and developing economies.

Beyond this initial release, the code infrastructure is designed to facilitate the updating of major data sources from traditional institutions as new vintages become available. It also supports the seamless incorporation of additional country-specific databases. These will be progressively added to enhance the quality of the estimates, as national statistical offices provide them in response to our requests. In line with this flexible design, the database is released in two versions to accommodate different user needs. The *ready-to-use version* contains the five core variables discussed in the next section (employment, nominal value-added, real value-added, nominal output, and real output), along with a limited set of metadata that identify the country, the year, and the relevant manufacturing sub-sector. This version is accompanied by an additional *detailed version*, which contains several tag variables intended to improve the transparency and tractability of the harmonisation and estimation procedures. These tags allow users to trace the source of each observation, the estimation methods applied (if any), the reference year used in the chain-linking procedure, the final configuration number identifying the reporting methods, the source of the deflator, and the deflator itself. It also includes a set of additional variables that were not retained in the core dataset, such as a regional classification and an adapted OECD industry classification capturing the technological intensity of each of the 12 sub-sectors (OECD, 2003; Vu et al., 2021). Finally, a replication package is provided to ensure full reproducibility and to allow users to adjust any methodological assumptions applied in the estimation procedure, should these be deemed too conservative or too aggressive. The final database and replication files will be made freely available through the corresponding GitHub repository¹.

The construction of the STiM database involved several steps. First, our methodological approach began by retrieving the latest version of the INDSTAT database from UNIDO (Revision 3), which serves as our baseline since it provides the longest temporal and broadest country coverage (UNIDO, 2024). We then downloaded external datasets from major institutions containing data at the manufacturing sub-sector level, namely the 2025 releases of the OECD STAN and EU-KLEMS databases (Horvát and Webb, 2020; Bontadini et al., 2023). After defining common boundaries for the manufacturing

¹The latest version of the STiM database and the replication package can be publicly accessed through the following link: <https://github.com/yanisbkt-econ/STiM-Database> .

sector and extracting only employment, value-added, and output values from these providers, we rescaled units and harmonised currencies. To do so, all monetary values were then converted into current US dollars using period-average bilateral exchange rates obtained from the International Monetary Fund (International Monetary Fund, 2024). This procedure was adopted to ensure consistency with the methodological approach applied in INDSTAT, which had been initially retrieved directly in current US dollars.

Second, even when expressed in the same unit (i.e., current dollars), these sources cannot simply be combined as they often do not rely on the same sample, industrial classifications, or reporting methods. In addition, INDSTAT raw values — which serve as the baseline for every final series — contain several temporal inconsistencies that may be overlooked when downloading directly consolidated variables from UNIDO. As such, every cross-country or intra-country analysis relying on these raw data is likely to be erroneous. To name but a few, national statistical offices may change their industrial classification over time, leading to spurious shifts in the sectoral composition of manufacturing. For instance, they might report values for the textile industry while also including apparel in some years, or report food products while also including beverages and tobacco in other years. In its native form, INDSTAT reports more than 134 different possible combinations of activities at the two-digit level, which vary across countries, variables, and years. Another issue relates to reporting methods that have evolved, following changes in national accounting standards (United Nations, 1968b; International Monetary Fund, 2025). Each national statistical office may change the way employment, value-added, or output is reported over time, leading to unintended breaks in the series. While employment can be reported as the number of employees or number of persons engaged, value-added and output can be reported at basic prices, producer prices, factor values, or unknown classification with potentially several changes occurring within the same country. These two issues (among others) make it impossible to rely directly on INDSTAT or to combine it with external sources without undertaking substantial aggregation and harmonisation efforts. To address these challenges, the second step consisted of developing a systematic and transparent harmonisation procedure, relying on chain-linking techniques and a set of conservative assumptions to ensure internal, intertemporal, and international consistency.

Finally, the last step of the construction process involved linking several price deflators from multiple institutions to adjust nominal measures, thereby producing estimates of real value-added and constant output. This step was particularly challenging, as deflators are often not available at the two-digit manufacturing sub-sector level for a wide set of countries, particularly developing economies, and over a long period. To mitigate this issue, we first draw on Haraguchi and Amann (2023), who derive sub-sectoral price indices from the Index of Industrial Production (IIP) embedded in INDSTAT. Since the IIP is designed to measure real production growth for each manufacturing activity and thus captures volume changes relative to a baseline year (2015), it indirectly provides price indices that can be retrieved and used to deflate nominal value-added. To complement this deflator, we further retrieve price indices

from the OECD STAN and EU-KLEMS databases, the OECD Main Economic Indicators (OECD, 2016), and the World Development Indicators from the World Bank (World Bank, 2022). Although each source has different levels of aggregation and coverage, they are combined according to a systematic preference order that moves from the most disaggregated to the most aggregated source. We first rely on the INDSTAT-derived price indices specific to each sub-sector, when available, to deflate nominal values. When these are not available, we rely on sub-sectoral deflators stemming from OECD STAN or EU-KLEMS. If these are also missing, we use aggregate producer price indices from the OECD for the manufacturing or industrial sector. In this case, we assume common price dynamics across all manufacturing sub-sectors or, more broadly, across the industrial sector (i.e., manufacturing, mining, construction, water, and electricity). Finally, if none of these options are available, we resort to country-level GDP deflators from the World Bank, thereby assuming common price dynamics across all major sectors (manufacturing, services, and agriculture). This hierarchy is applied consistently to every country by chain-linking the available series and rebasing each price index to 2015.

This paper contributes to the existing literature by extending previous academic efforts to compile and harmonise macroeconomic statistics for a wide range of countries. To our knowledge, very few studies have attempted to build a comprehensive panel at the manufacturing sub-sector level. This is primarily because INDSTAT — the most comprehensive source of information covering manufacturing sub-sectors — has only been freely available since February 2022. So far, the closest attempt to build such a panel was undertaken by Pahl and Timmer (2020). Although their primary focus was not to provide a harmonised database, they nonetheless constructed a panel combining the same data sources as this paper, relying on the same kind of chain-linking procedure. Yet, we believe this paper substantially improves on their pioneering work in several ways. First, we refine the cleaning procedure to detect spurious zeros and outliers in the raw INDSTAT observations. Second, our harmonisation procedure extends coverage back to the 1960s and up to 2019, thereby lengthening the period covered by more than twenty years. To achieve this, we enhance the chain-linking method to maximise the continuity of each country's series by allowing for several possible reference years. Third, we implement ex-post quality checks to test the consistency of the final estimates. It aims to ensure that no spurious outliers remain in the final series due to growth-rate contaminations as the distance from the reference year increases. Fourth, our database is fully reproducible and will be further extended as new sources become available in response to our requests. Lastly, we provide real value-added and constant output measures by linking several price deflators from multiple institutions to adjust nominal measures. This makes the STiM database the most comprehensive dataset for long-run analyses of structural change within manufacturing in both developed and developing economies.

The paper proceeds as follows. Section 2 presents the main manufacturing data sources and discusses the challenges involved in harmonising and combining them. Section 3 presents the industrial classification adopted in the STiM database, which results

in an aggregation into 12 manufacturing sub-sectors. While this approach reduces cross-sector variation, it addresses the issue of unusually reported combined activities and facilitates linkages with external datasets. Section 4 details the chain-linking procedure used to ensure internal, intertemporal, and international consistency. Several new techniques are introduced to maximise the continuity of each country's series while minimising spurious outliers. Section 5 describes the price indices and the deflation procedure, while Section 6 presents some few descriptive statistics for the STiM database. Finally, Section 7 concludes and discusses potential extensions of the database.

2 Presentation of the main data sources

In this section we discuss the main data sources used to construct the Structural Transformation in Manufacturing Database (STiM). We first present the INDSTAT Revision 3 database, which serves as the primary source of information for employment, value added and output at the two-digit manufacturing sub-sector level. We then describe the two external datasets that are linked to INDSTAT to complement estimates and coverage.

2.1 The INDSTAT Revision 3 Database

The main source of the Structural Transformation in Manufacturing Database (STiM) is the INDSTAT Revision 3 dataset, which compiles national industrial surveys and representative censuses collected by the United Nations Industrial Development Organization (UNIDO, 2024). These surveys traditionally exclude firms with fewer than five and sometimes ten employees, depending on censuses. This feature thus confines this database to only formal and registered industrial activities. This extensive data collection, conducted since the 1960s and made freely available in February 2022, provides a wide range of variables at the two-, three-, and four-digit levels of aggregation according to the International Standard Industrial Classification of All Economic Activities (ISIC), Revision 3.1 (United Nations, 2002). For the purpose of this database, we focus exclusively on employment, nominal value added, and nominal output at the two-digit level of aggregation, which allows us to construct the longest possible time series. It should be noted that all monetary variables are retrieved directly in current US dollars. The original values, expressed in local currency units (LCU), were converted to current US dollars using period-average bilateral exchange rates from the IMF (International Monetary Fund, 2024).

The initial raw data cover 186 countries and 31 industrial sub-sectors, spanning the period from 1963 to 2022. Yet, this panel is unbalanced, with many gaps across years, sub-sectors, and variables, so that each country may enter or exit the sample several times during the period. In order to ensure consistency and comparability across all datasets, we first limit the sample to 2019 to avoid the decline in observations due to

ongoing reporting. Secondly, we restrict the sub-sectors to manufacturing activities only, corresponding to those whose ISIC codes range from 15 to 37 (United Nations, 2002). The details of these 23 manufacturing sub-sectors are provided in Table 1. Overall, this very preliminary cleaning step results in a raw sample of 23 manufacturing sub-sectors across 179 countries, with 109,246 observations for employment, 103,837 for output, and 98,230 for value added.

Table 1: ISIC Rev. 3.1 Manufacturing Categories (Section D)

ISIC, Rev. 3.1	Initial ISIC Categories available
15	Food and beverages
16	Tobacco products
17	Textiles
18	Wearing apparel, fur
19	Leather, leather products and footwear
20	Wood products (excl. furniture)
21	Paper and paper products
22	Printing and publishing
23	Coke, refined petroleum products, nuclear fuel
24	Chemicals and chemical products
25	Rubber and plastics products
26	Non-metallic mineral products
27	Basic metals
28	Fabricated metal products
29	Machinery and equipment n.e.c.
30	Office, accounting and computing machinery
31	Electrical machinery and apparatus
32	Radio, television and communication equipment
33	Medical, precision and optical instruments
34	Motor vehicles, trailers, semi-trailers
35	Other transport equipment
36	Furniture; manufacturing n.e.c.
37	Recycling

These three main consolidated variables cannot yet be fully exploited in their current form, as they are merely a combination of the different reporting methods used to measure employment, value added, and output. They indirectly reflect changes in the System of National Accounts (SNA) recommendations, which have evolved over time (United Nations, 1968b; International Monetary Fund, 2025). Importantly, these changes are specific to each country and year, with absolutely no common pattern over time or across regions, as national statistical offices may or may not apply these rules. As such, these three initial consolidated variables are not suitable for empirical analysis, since each change in reporting methods introduces unintended variations, notably level shifts, in the series. As a result, we discard these consolidated series and recover the initial ten unconsolidated variables that correspond to each reporting method (cf. Table 2). Two variables correspond to employment, measured either as the number of employees or as the number of persons engaged. The remaining eight variables correspond to value added and output, each reported under three different valuation methods (basic prices,

producer prices, and factor cost), plus an additional category of unknown classification. Three points warrant particular attention. First, for a given country-year observation, statistical offices report each variable under only one definition or valuation method. For example, employment is reported either as the number of employees or as the number of persons engaged, but never both simultaneously. Similarly, value added and output are reported under only one valuation method at a time. Second, value added and output variables, when available, are not necessarily reported using the same valuation method for a given country-year observation, which introduces further inconsistencies. Third, one should not underestimate the bias each change in reporting methods may introduce.

For example, in Brazil, the statistical office changed the definition of employment in 1996, switching from the number of employees to the number of persons engaged. This led to a sudden increase of hundreds of thousands of workers in activities with a high prevalence of self-employment, such as in the Food and Beverages industry (ISIC 15, Rev. 3.1). Indeed, the reported value nearly doubled, moving from 596,406 employees in 1995 to 1,008,577 persons engaged in 1996. Although less pronounced than in the case of employment measurement, a similar level shift may exist across the three valuation methods for reporting value added and output. In a nutshell, under basic price valuation, subsidies on products are included while taxes are excluded to reflect what the producer actually receives (International Monetary Fund, 2025, p.214). Conversely, under producers' prices, taxes are included and subsidies excluded to assess what the producer charges (International Monetary Fund, 2025, p.214). Lastly, under factor cost, both taxes and subsidies on products are excluded, and value added is obtained by adjusting value added at basic prices for other taxes and subsidies on production (International Monetary Fund, 2025, p.218). Given that taxes and, notably, subsidies are well-known instruments of industrial policy (U. C. V. Haley and G. T. Haley, 2013; Aiginger and Rodrik, 2020), these different valuations methods can indeed lead to discrepancies, depending on each country, year, and sub-sector. The details on the frequency of these ten unconsolidated variables are provided in Table 2

Another concern with this dataset relates to how the raw 23 sub-sectors are defined in the ISIC Rev. 3.1 classification and how this structure is maintained across earlier ISIC revisions. At first glance, the dataset logically imposes a panel structure that includes 23 manufacturing sub-sectors, consistent with the ISIC Rev. 3.1 classification at the time of release. Yet, a closer look reveals that some sub-sectors only began to be reported after a certain year, which often coincides with the introduction of new ISIC revisions. In earlier decades, previous ISIC revisions were in force, and certain industrial activities that are now reported separately were then aggregated into broader categories, reflecting differences in classification norms. For instance, manufacturing activities linked to leather products and the footwear industry (ISIC 19, Rev. 3.1) were only distinguished from the wearing apparel industry (ISIC 18, Rev. 3.1) in the early 1990s. This became possible only with the introduction of the ISIC Rev. 3 classification in 1989 (United Nations, 1989). Consequently, there are no historical data for these

Table 2: Summary statistics of the unconsolidated variables

Variables	N	Mean	Min	Max
Employment				
Employees	93,142	94,508.51	0	10,200,000
Persons engaged	16,104	44,519.07	0	1,871,000
Value added				
Basic prices	38,877	7.93e+08	-6.12e+08	1.85e+11
Producer prices	30,272	3.05e+09	-1.28e+09	5.90e+11
Factor cost	11,252	3.05e+09	-8.29e+08	1.39e+11
Unknown classification	17,829	7.39e+09	-4.89e+09	5.53e+11
Output				
Basic prices	9,754	5.63e+09	0	2.53e+11
Producer prices	45,282	3.76e+09	0	1.30e+12
Factor cost	17,023	5.51e+09	0	5.04e+11
Unknown classification	31,778	1.88e+10	0	2.05e+12

activities prior to this date since national statistical institutes were all following the ISIC Rev. 2 classification then in force (United Nations, 1968a).

Unfortunately, these issues of industry combination are not confined to such clear-cut and rational cases. There are many more complex cases where countries report certain sub-sectors jointly for no apparent reason, even though the industrial classification in force at the time does not recommend it. In its raw form, INDSTAT reports more than 134 different possible combinations of activities at the two-digit level, which vary across countries, variables, and years. To give a concrete example, Peru reports correctly employment as the number of persons engaged in the Food and Beverages sector (ISIC 15, Rev. 3.1) without any combination from 1979 (the year of its entry into the database) to 2003. Then, for whatever reason, from 2004 to 2013 it reports the Food and Beverages sector (ISIC 15, Rev. 3.1) combined with Tobacco products (ISIC 16, Rev. 3.1), before reverting to the original classification from 2014 to 2019. It is also in 2014 that Peru began reporting employment as the number of employees instead of persons engaged, thereby adding further noise and united variation to the series. As such, relying directly on the original consolidated variables provided by INDSTAT is likely to lead to erroneous conclusions as it mixes different valuations methods but also because the composition of each sub-sector might changes over time.

Overall, these issues of comparability and consistency across countries, variables, and over time must be addressed before any meaningful analysis can be conducted — the solution proposed to mitigate these issues are discussed in section 4.

2.2 External datasets used to enhance estimates and coverage

To enhance the coverage of the INDSTAT database, we aimed in linking all available external datasets containing data on employment, value added, or output at the two-

digit manufacturing sub-sector level. To our knowledge, only two main public datasets meet these criteria, namely the OECD STAN and EU-KLEMS databases². While these datasets do not cover as many countries as in INDSTAT, they do offer valuable harmonised data for a limited number of countries, particularly high-income economies in the OECD or in the European Union (EU). These two datasets were thus downloaded in their latest available versions, released in 2025. We will update them as new vintages become available and intend to incorporate additional country-specific datasets as they are provided by national statistical offices in response to our requests.

2.2.1 The OECD STAN Database

The OECD STAN database is built primarily on national accounts and industrial surveys harmonised by the OECD (Horvát and Webb, 2020). It provides internationally comparable series on output, value added, and employment that are consistent with the System of National Accounts 2008 (International Monetary Fund, 2009). In practice, it ensures that the definitions and accounting principles used to compile the data are consistent across countries and sub-sectors over time. As such, employment is always retrieved as the number of employees, while value added and output are reported at basic prices allowing direct cross-country comparisons. Yet, all monetary values are originally reported in local currency units (LCU). To ensure comparability with the INDSTAT database, these values are converted into current US dollars using period-average bilateral exchange rates from the IMF (International Monetary Fund, 2024). Regarding the time coverage, it generally begins in the early 1970s and extends to the most recent year available, which we restrict to 2019 to mirror the preliminary cleaning step applied to INDSTAT. Lastly, as briefly discussed, the database only focuses on OECD member countries thus covering a total of 38 countries in its latest version.

In terms of industrial classification, the OECD STAN database uses the NACE Rev. 2 classification, which exactly mirrors the ISIC Rev. 4 classification (European Commission, 2008; United Nations, 2008). More details about the 24 included sub-sectors can be found in Table 3. Yet, for some countries and sub-sectors the OECD STAN does not report data at the individual sub-sector level but rather for aggregated combinations, although such cases remain limited. However, regardless of the country, the database always reports data for the combination of activities related to the manufacture of furniture (ISIC 31, Rev. 4) and other manufacturing activities (ISIC 32, Rev. 4), thereby making it impossible to retrieve separate values for these two sectors. This is the only case in which neither of the two sub-sectors is reported individually. For all other sub-sectors, the OECD STAN database provides individual values, even when combined categories are also reported in cases where national data do not allow a finer distinction.

²At the time of writing this documentation, we identified the BADECON database, which covers eight Latin American countries at the four-digit manufacturing level from 2010 to 2021. This source will be integrated into the STiM database in a forthcoming update.

Table 3: ISIC Rev. 4 Manufacturing Categories (Section C)

ISIC, Rev. 4	Manufacturing Categories
10	Manufacture of food products
11	Manufacture of beverages
12	Manufacture of tobacco products
13	Manufacture of textiles
14	Manufacture of wearing apparel
15	Manufacture of leather and related products
16	Manufacture of wood, etc.
17	Manufacture of paper and paper products
18	Printing and reproduction of recorded media
19	Manufacture of coke and refined petroleum products
20	Manufacture of chemicals and chemical products
21	Manufacture of basic pharmaceutical products, etc.
22	Manufacture of rubber and plastics products
23	Manufacture of other non-metallic mineral products
24	Manufacture of basic metals
25	Manufacture of fabricated metal products, etc.
26	Manufacture of computer, electronic and optical products
27	Manufacture of electrical equipment
28	Manufacture of machinery and equipment n.e.c.
29	Manufacture of motor vehicles, trailers and semi-trailers
30	Manufacture of other transport equipment
31	Manufacture of furniture
32	Other manufacturing
33	Repair and installation of machinery and equipment

2.2.2 The EU KLEMS Database

The EU KLEMS database provides detailed employment, value added, and output data for 27 EU Member States, the United States, Japan, and the United Kingdom (Bontadini et al., 2023). It covers the period 1995-2021, with some gaps and missing values, as is the case with the other external datasets. To ensure comparability with the previous sources, we restrict the sample to 2019 and to the manufacturing sector only. Although the database has been primarily used for studying productivity, it also includes gross value added, employment, and output data at the two-digit manufacturing sub-sector level, making it a suitable candidate to complement the INDSTAT database. This is especially relevant for European countries that are not part of the OECD, such as Bulgaria, Croatia, Cyprus, Malta, and Romania. Just as with the STAN database, the EU KLEMS dataset reports all observations in line with the System of National Accounts (SNA) published in 2008 (International Monetary Fund, 2009). Accordingly, employment is consistently measured as the number of employees, while value added and output are reported jointly at basic prices. Moreover, all monetary values are converted from local currency units (LCU) to current US dollars using period-average bilateral exchange rates from the IMF (International Monetary Fund, 2024).

Table 4: ISIC Rev. 4 manufacturing codes in the EU-KLEMS database

ISIC Rev. 4 Code	Freq.	Percent	Cum.
C10–C12	750	6.67	6.67
C13–C15	750	6.67	13.33
C16–C18	750	6.67	20.00
C19	750	6.67	26.67
C20	750	6.67	33.33
C20–C21	750	6.67	40.00
C21	750	6.67	46.67
C22–C23	750	6.67	53.33
C24–C25	750	6.67	60.00
C26	750	6.67	66.67
C26–C27	750	6.67	73.33
C27	750	6.67	80.00
C28	750	6.67	86.67
C29–C30	750	6.67	93.33
C31–C33	750	6.67	100.00
Total	11,250	100.00	

The database adopts the same industrial classification as the OECD STAN database, which mirrors ISIC Rev. 4 categories (cf. Table 3). However, the EU KLEMS database exhibits even more frequent combinations of activities across sub-sectors. While the OECD STAN database fails to report only two out of 24 sub-sectors individually, the EU KLEMS database does not report up to 18 out of 24 sub-sectors individually for all the sample. The details of the frequency of these ISIC Rev. 4 codes are provided in Table 4. Most of them are systematically retrieved in a combined form across all countries and years. For example, the manufacture of food products (ISIC 10, Rev. 4) is always reported jointly with the beverage industry (ISIC 11, Rev. 4) and the manufacture of tobacco products (ISIC 12, Rev. 4). The only sub-sectors that are consistently reported individually are the manufacture of coke and refined petroleum products (ISIC 19, Rev. 4), the manufacture of chemicals (ISIC 20, Rev. 4), the manufacture of basic pharmaceutical products (ISIC 21, Rev. 4), the manufacture of computer, electronic, and optical products (ISIC 26, Rev. 4), the manufacture of electrical equipment (ISIC 27, Rev. 4) and the manufacture of machinery and equipment (ISIC 28, Rev. 4).

2.3 The two main challenges when linking external datasets to INDSTAT

Overall, two main issues arise when linking these aforementioned datasets to INDSTAT. First, it requires finding an equivalence table between the ISIC Rev. 3.1 and ISIC Rev. 4 classifications. This required building a concordance table to ensure consistency across the two classifications, knowing that a perfect match is impossible because of all the changes introduced between the two revisions (United Nations, 2002; United Nations, 2008). A somewhat further complication arises from the numerous combinations of activities reported in INDSTAT and in some external datasets (e.g., EU-KLEMS). Taken

together, these constraints limit the flexibility when trying to build an equivalence between these datasets.

Second, even with a concordance table, the question remains of how to link these external datasets to INDSTAT. We cannot simply replace values from INDSTAT with those from OECD STAN or EU KLEMS, even if the latter appear more reliable. This is because of differences in levels across sources, even for the same country, year, and sub-sector. Although this might seem counterintuitive, these discrepancies mainly stem from differences in sampling procedures, as the underlying sources and censuses of the three datasets are not necessarily the same. Moreover, series are not always reported using the same valuation method, which further complicates the linkage. While both the OECD STAN and EU KLEMS databases report each variable in a harmonised way, INDSTAT does not, even for large economies covered by all three datasets. For example, in INDSTAT, France, Spain, and the United Kingdom have in recent years reported value added only at factor cost. In such cases, directly substituting values from OECD STAN or EU KLEMS when a missing observation occurs in INDSTAT would introduce biases in levels, both within the time series (before and after the break) and because of differences in valuation methods.

The first issue is addressed by introducing a new aggregation of sub-sectors, which is presented in the next section. The second is resolved by using only the growth rates from STAN and EU KLEMS to extrapolate the original INDSTAT series. As such, no difference in levels is introduced, since we retain the values from INDSTAT. In addition, the problem of varying valuation methods is mitigated by assuming that, regardless of the reporting way, growth rates correctly capture the employment, value-added, and output dynamics of each sector. This issue will be discussed in greater detail in the section on cleaning and chain-linking procedures.

3 Aggregation into 12 manufacturing sub-sectors

In order to link the different data sources and address the aforementioned issues of combined activities, we had no choice but to aggregate the data into broader manufacturing sub-sectors. We proceed in two steps. First, we examined the 134 combinations of activities initially reported in INDSTAT and grouped sub-sectors according to the most frequent combinations observed. As discussed previously, these were mostly attributable to classification changes from ISIC Rev. 2 to ISIC Rev. 3 in 1989, as well as to unusual reporting practices. This procedure resulted in an aggregation into 12 manufacturing sub-sectors, identical to those used by Pahl and Timmer (2020) in their seminal paper. With respect to INDSTAT, observations corresponding to combinations that did not fit within this aggregation were dropped from the analysis. However, they represented less than 2% of total observations across each of the ten unconsolidated variables, which constitutes a negligible loss. The main trade-off of such aggregation is the reduction of cross sub-sector variation that could have been exploited by researchers, as the number of sub-sectors is reduced from 23 to 12. At the same time, this

ensures a much higher degree of consistency and avoids comparing sub-sectors across years for which composition was changing.

Second, moving from ISIC Rev. 3.1 to ISIC Rev. 4 while remaining at the two-digit level inevitably leads to some misallocations. Mismatches between the two classifications necessarily arise due to changes at the three- and four-digit levels, even after several attempts to aggregate the initial 23 sectors into 12, 10, or 8 sub-sectors. Given that changes between ISIC Rev. 3.1 and ISIC Rev. 4 are too substantial and that no optimal solution exists, we decided to follow the correspondence applied by the OECD in their STAN database Horvát and Webb (2020, p.44). Although not perfect, this solution appeared to us to be the most reasonable given that their database have been used several times in the literature. Importantly, while their correspondence departs from the 23 initial sub-sectors, it remains compatible with our aggregation into 12 sub-sectors. This allows us to get a full correspondence between the classification that will be used in the STiM database (i.e., the new 12 sub-sectors) with ISIC Rev. 3, ISIC Rev. 3.1, ISIC Rev. 4 and NACE 2. Details are provided in Table 5.

Table 5: Aggregation into 12 manufacturing sub-sectors

STiM, V1	ISIC, Rev. 3.1	ISIC, Rev. 3	ISIC, Rev. 4	NACE 2
1	15, 16	15, 16	10+11, 12	10+11, 12
2	17, 18, 19	17, 18, 19	13, 14, 15	13, 14, 15
3	20	20	16	16
4	21, 22	21, 22	17, 18	17, 18
5	23	23	19	19
6	24	24	20+21	20+21
7	25	25	22	22
8	26	26	23	23
9	27, 28	27, 28	24, 25	24, 25
10	29, 30, 31, 32, 33	29, 30, 31, 32, 33	28, 26, 27, 26, 26	28, 26, 27, 26, 26
11	34, 35	34, 35	29, 30	29, 30
12	36, 37	36	31+32+33	31+32+33

Note: This correspondence table is indirectly based on Pahl and Timmer (2020) and Horvát and Webb (2020) as discussed in the core of the paper. Labels for each sub-sector are provided in Table 6.

When applying this classification to all data sources, it results to the loss of additional information beyond those already dropped because of an incompatibility between the current combination activity and the new classification (e.g., as in INDSTAT or EU-KLEMS). Indeed, the construction of the 12 aggregated sub-sectors logically requires complete information on all underlying components. For instance, when creating the new sub-sector “Food, beverages and tobacco products”, if an observation is missing for either the manufacturing of food and beverages or the tobacco industry, then the entire aggregated observation for this aggregate sub-sector is not retained. This constraint ensures that all industrial activities are always consistently comparable across years. To complement, the database we also link each of the 12 sub-sectors to their technological intensity following the OECD taxonomy (OECD, 2003) and Vu et al. (2021). This allows

anyone to make use of the database for analyses related to manufacturing sophistication and technological upgrading. The correspondence between each sub-sector and its technological intensity is provided in Table 6, where we also attach reconstructed labels for the 12 sub-sectors. For the sake of flexibility and transparency, we also highlight that the replication files have been designed to allow users to easily adapt the re-aggregation process if needed. Moreover, for tractability purposes, two tag variables are created to link each new sub-sector to its original components in ISIC Rev. 3.1 and ISIC Rev. 4, respectively. These variables are included in the *detailed version* of the STiM database.

Table 6: New aggregated labels and linking to OECD taxonomy

STiM, V1	New label	OECD Taxonomy
1	Food, beverages and tobacco products	Low-tech industries
2	Textiles, wearing apparel, leather products, fur	Low-tech industries
3	Wood products (excl. furniture)	Low-tech industries
4	Paper products, printing and publishing	Low-tech industries
5	Coke, refined petroleum products, nuclear fuel	Med-tech industries
6	Chemicals and chemical products	High-tech industries
7	Rubber and plastics products	Med-tech industries
8	Non-metallic mineral products	Med-tech industries
9	Basic and fabricated metal products	Med-tech industries
10	Machinery, equipment and electronic products	High-tech industries
11	Transport equipment	High-tech industries
12	Other manufacturing and recycling	High-tech industries

Note: The correspondence between each sub-sector and its technological intensity is derived from OECD (2003) and adapted following Vu et al. (2021). The STiM code number 12 — Other manufacturing and recycling — can be classified as either medium-tech or high-tech industries, since the sub-sectors it aggregates are typically assigned to both categories. To balance the classification (with four sub-sectors in each technological category), we assign it to high-tech industries. Nevertheless, scholars should test whether their results are sensitive to this choice.

4 Cleaning and chain-linking procedure

While the aggregation scheme opens up the possibility of combining data from several sources and solves unusual combination of activities, the INDSTAT database on which we mostly rely still exhibit several inconsistencies. As such, this section discusses the initial cleaning procedure to detect spurious zeros and outliers, describe the chain-linking procedure aimed at maximising the consistency of estimates, details how we integrate external datasets and present some ex-post quality checks carried out on the final series.

4.1 Initial cleaning procedure

The initial cleaning procedure discussed below is carried out on all raw series before any aggregation. It only concerns the INDSTAT database whose series are the most prone to inconsistencies. Except aggregating them in twelve sub-sectors, external

sources (OECD STAN & EU-KLEMS) remain untouched as they are already cleaned and harmonised by the institutions providing them.

Raw unconsolidated series from INDSTAT primarily face three major issues. First, some series contain spurious zeros that do not reflect the actual absence of activity but rather missing values that have been coded as zeros by the data providers (Pahl and Timmer, 2020). Second, some series contain duplicates, meaning that the same value is reported several times for a given country-subsector combination across different years. We address these cases differently depending on the type of variable. Third, some series contain extremely large and unrealistic outliers that might contaminate the growth rates used in the chain-linking procedure. While no explanations are provided, we assume that these values either result from reporting errors, manipulations, or incorrect conversions when using period-average exchange rates to retrieve values in current dollars. Fourth, we also deal with missing values by linearly interpolating them under restricted conditions.

4.1.1 Detecting spurious zeros and dropping negative values

To begin with, we set all negative values to missing, as they do not make sense in the context of employment, value added, nor output. As shown in Table 2, these cases are confined to the unconsolidated value-added variables. We then follow Pahl and Timmer (2020) to detect zeros that are likely to correspond to missing values. First, we consider a zero spurious if it appears between two positive observations. Second, if a zero is reported in one of the three main variables of interest while at least one of the other two records a positive value, we also treat it as spurious. Third, if a zero is followed by a positive value such that the sub-sector suddenly accounts for more than 5% of total manufacturing, this zero is likely to reflect a missing value rather than a true observation. Similarly, a zero is treated as missing if it follows a year in which the sub-sector accounted for more than 5% of total manufacturing, since such an abrupt disappearance is much unlikely. Finally, we also treat zeros as missing when they occur between two missing values.

4.1.2 Dealing with duplicates

Some statistical offices in certain countries tend to report identical observations for the same sub-sectors over several consecutive years. While this may reflect a lack of updates and should therefore be treated as missing in most cases, we also consider that in some sub-sectors this phenomenon might be rational when reporting employment. Accordingly, we treat duplicates differently depending on the variable. If duplicates appear in employment for no more than two consecutive years within a given sub-sector, we retain them only if the total sum of manufacturing employment changes across these years. In this case, we assume that the duplicates reflect a temporary stagnation in that sub-sector. If duplicates in employment extend beyond two consecutive years, or if the total sum of manufacturing employment does not change, we consider only

the earliest year as valid, with all subsequent values set to missing. For unconsolidated value-added and output series, the procedure is more straightforward as we do not tolerate any duplicates for any country-sub-sector combination. As with employment, we thus retain only the earliest year and set all subsequent duplicates to missing.

4.1.3 Removing outliers

To identify outliers in each unconsolidated series, we apply several filters. These mainly concern value-added and output series, as employment is less likely to exhibit extreme variations. First, we retrieve total manufacturing value added in current dollars from the World Development Indicators (WDI) database (World Bank, 2022) and merge it with the INDSTAT series. Although the two sources are not fully comparable, since they rely on different sampling methods, the total value added produced by the manufacturing industry in a given country-year cannot reasonably be exceeded by the value added of only one of its sub-sectors. We therefore assume that even if the valuation between the two sources is not perfectly aligned, such differences cannot justify a sub-sector exceeding the total value added produced by the whole manufacturing sector. Accordingly, we set to missing any sub-sector observation whose value added exceeds the corresponding total manufacturing value added from the WDI. If this occurs, we also set the output (when available) to missing for the same country, sub-sector, and year. Additionally, we check whether the sum of value added across the 12 manufacturing sub-sectors exceeds a country's total manufacturing value added in a given year. Given differences in valuation methods, we allow for a certain margin of tolerance. As such, if the sum of sub-sectors exceeds 175% of total manufacturing value added, we set all sub-sector observations to missing for that country-year, as it is not possible to determine whether the discrepancy originates from one or several sub-sectors. In these cases, we also set output to missing for all sub-sectors in that country-year.

Second, we repeat the same procedure using total current GDP from the same database (World Bank, 2022). While this represents a less stringent filter than the one based on total manufacturing value added, it covers a broader range of countries and extends back to the 1960s, thereby allowing for a more comprehensive temporal and spatial check. Additionally, we compute the share of each sub-sector in total manufacturing value added and analyse how this share changes from one year to the next. We flag cases where the share of a sub-sector in GDP increases or decreases by more than 5 percentage points relative to the previous year. Although not perfect, we then investigate these cases individually to assess whether they are plausible, whether they correspond to major inflationary peaks or political crises, and whether such differences in levels persist over time. When no explanation can be found for abrupt changes, and they do not concern the oil subsector, we set the value-added and output observations to missing for that country-subsector-year. Out of 24 cases identified, we ultimately set 8 as missing. The following cases are set to missing: Chile, Basic metals (1974); Côte d'Ivoire, Chemicals (1989); Fiji, Food and beverages (2001); Singapore, Machinery and

equipment (1989 and 1990); Thailand, Textiles (1990); Venezuela, Chemicals (1998); and Sierra Leone, the entire series.

4.2 Linear interpolation

Lastly, we perform a careful linear interpolation at the country, sub-sector, and variable levels for the INDSTAT Database³. This interpolation is applied within each unconsolidated series to preserve the specific levels associated with each valuation method, but it occurs *only* after the aggregation into twelve sub-sectors. This sequencing ensures that the interpolation relies on values that correspond to the correct industrial activity as defined by the panel, rather than on one of the 134 combinations discussed earlier. However, we restrict interpolation to gaps of up to five consecutive years in a given series to avoid adding noise to the final estimates.

At this stage of the procedure, it applies to 3.6% of all employment observations, 0.3% of value-added observations, and 0.2% of output observations. For tractability purposes, a tag variable is created to identify which values have been interpolated in the additional *detailed version* of the STiM database.

4.3 Chain-linking procedure

In short, to ensure internal and intertemporal consistency within each country and across the three main variables of interest (employment, value added, and output), we implement a chain-linking procedure based on the growth rates of each series. This procedure, also used by Pahl and Timmer (2020), is standard in the literature on national accounts (Kruse et al., 2022). The idea is to establish an internal reference year common to all final variables for each country. Starting from this country-specific baseline, we reconstruct the consolidated series by extrapolating both backward and forward using the growth rates of the underlying unconsolidated series. In other words, this method assumes that the growth rates observed in each unconsolidated series correctly reflect the actual dynamics of the corresponding activity, such that combining these growth rates from a common reference year produces consistent series. The resulting consolidated series are internally consistent, since they are all derived from the same baseline and growth rates, and “intertemporally” reliable, since no breaks in valuation methods occur once fixed by the reference year. Ultimately, to ensure international comparability, all countries should share the same valuation method in the reference year (across the three main variable), which is the objective of our procedure, although this may not always be achievable given differences in national statistical office practices. This section aims at providing a detailed description of the chain-linking procedure.

³We don’t apply this linear interpolation to the external datasets as they already have their own cleaning, harmonisation and inputting procedure.

4.3.1 Choosing the reference year

The first and crucial step of the chain-linking procedure involves determining a reference year that is specific to each country. This year serves as the baseline from which all other years are extrapolated using growth rates from the combined unconsolidated series. However, to ensure internal and intertemporal consistency, the reference year must be the same across the three main variables of interest within each country. Since we also aim to maximise both the length of the final series and the international comparability of the estimates, we design an algorithm to select the optimal reference year for each country.

- Step 1. Identifying the longest streak. We first tag the longest sequence of consecutive years during which each of the three main variables (employment, value added, and output) is available and covers at least one sub-sector. We then compute a score by summing the three variables, which allows us to identify the longest possible streak of coverage for each country.

- Step 2. Prioritising valuation methods. We then establish, ex ante, a preferred hierarchy of valuation methods for each variable. As recommended by the SNA (International Monetary Fund, 2009; International Monetary Fund, 2025), we give preference to employment series measured as the number of employees over those measured as the number of persons engaged. For value added and output, we prefer series expressed at basic prices over those at producer prices, and producer prices over factor prices. If none of these preferred methods are available, we fall back on the classification reported as “unknown.”

- Step 3. Counting available observations. For each year, we compute the total number of non-missing observations across our 10 unconsolidated series that covers three main variables. Since we have 12 sub-sectors, the maximum possible count for one year is 36. Note that we ensure to lock final configurations when doing this count as the interpolation have created, sometimes, some overlap between two different reportings methods from the same variable.

- Step 4. Selecting the reference year. When having these pieces of information, we can proceed to the final step which consists in selecting the reference year. For each country, the reference year is chosen in the following order of priority:

1. The year must belong to the longest streak of consecutive years identified in the very first step. This avoids selecting isolated years with very limited surrounding data, which is particularly important since the extrapolation of values is highly sensitive to missing observations in the growth-rate series.
2. Among these years, we first select those when employment is reported as the number of employees. If none exist, we retain years when employment is reported as the number of persons engaged.

3. Within this set, we further restrict to years when value added is reported according to our preferred order of valuation methods (basic prices, then producer prices, then factor prices, then unknown classification).
4. Still retaining these years, we apply the same preference order to output. If possible, output should be valued using the same method chosen for value added to ensure consistency. If this is not possible, we relax this condition and select the second preferred method for output, and so on (basic prices, then producer prices, then factor prices, then unknown classification).
5. Among the remaining years, we choose the one with the highest count of non-missing observations across the three variables.
6. Lastly, if multiple years still satisfy all these conditions, we select the most recent one as the reference year.

Note that, for tractability purposes, one tag variable is created to identify the reference year selected for each country, along with three additional tag variables that record the valuation method used for each of the three main variables in that year. These variables are particularly useful for researchers who wish to avoid conducting cross-country analyses on value added, output, or employment when the underlying levels are not based on the same valuation method. These tag variables are included in the *detailed version* of the STiM database.

4.3.2 Combining the growth rates from all available sources

Once a unique reference year is selected for each country, we proceed to the second step of the chain-linking procedure, which involves combining the growth rates of the three main variables of interest. At this stage, external sources are integrated into the process thanks to the aggregation discussed in the previous section. We thus construct three final growth-rate series — one for employment, one for value added, and one for output. Each combined growth-rate series encompasses all available growth rates from every dataset and valuation method. Importantly, these combined growth-rate series are built from pre-computed growth rates within each unconsolidated source. In other words, for INDSTAT, we first compute growth rates within each valuation method and then combine them, rather than merging raw values across valuation methods to derive growth rates. This approach consistent with international recommendations ensures that each rate reflects a consistent definition of the underlying variable.

For instance, the final employment growth-rate series that gonna be used to reconstruct employment levels is constructed according to the following order of priority. If available, the final series relies on employment growth rates expressed as the number of employees from INDSTAT or the employment growth rates expressed as the number of persons engaged, as reported by UNIDO. If not available, it collects the growth rates from external sources, namely OECD STAN or EU KLEMS. The same logic applies to

value-added and output growth-rate series, where the order of priority is determined by the preferred valuation methods discussed previously (International Monetary Fund, 2009; International Monetary Fund, 2025). Note that when two series rely on the same valuation method, we always prioritise INDSTAT over OECD STAN and EU KLEMS since INDSTAT act as the primary source for the STiM database. Full details regarding the construction of each combined growth-rate are outlined in Table 7.

Table 7: Order of priority for constructing the combined growth-rate series

Priority	Source	Employment	Source	Value added/Output
1	INDSTAT	Persons employed	INDSTAT	At basic prices
2	INDSTAT	Persons engaged	INDSTAT	At producers' prices
3	OECD STAN	Persons employed	INDSTAT	At factor costs
4	EU KLEMS	Persons employed	INDSTAT	Unspecified valuation
5	-	-	OECD STAN	At basic prices
6	-	-	EU KLEMS	At basic prices

Note: We respect recommendations outlined by the System of National Accounts (SNA) since 2008 (International Monetary Fund, 2009; International Monetary Fund, 2025). When valuation methods are the same, we always prioritise INDSTAT over OECD STAN and EU KLEMS. See main text for more details.

So far, the final combined growth-rate series for each of the three main variables (employment, value added, and output) is primarily constructed from INDSTAT, which accounts for more than 90% of all observations in the final series. OECD STAN and EU KLEMS contribute only marginally to filling gaps in the INDSTAT series. This might change in future versions of the STiM database as more external sources are planning to be integrated. Note that, for tractability purposes, three tag variables are created to identify the source of the growth rate for each of the three main variables in each country-year. It allows researchers to easily track the source of the final estimated values. These tag variables are included in the *detailed version* of the STiM database.

4.3.3 Assuming constant labour-productivity to bridge minor gaps in growth rates combined series

Whenever a change in reporting methods occurs, it mechanically generates a break in the combined growth-rate series of a single variable (e.g., the combined growth rate of all value-added measures). This happens because growth rates are computed separately within each unconsolidated series before being combined. This approach is the only way to ensure that each rate effectively captures a consistent definition of the underlying variable. By contrast, combining raw values first and then computing growth rates would imply calculating changes across reporting methods, thereby conflating methodological breaks with genuine economic dynamics. Although this procedure follows most international recommendations, it has one drawback in our setting. Since there are two to four possible reporting methods for a single final measure,

each change in valuation creates a one-year gap in the final growth-rate series of the variable concerned. As a result, when a country frequently changes its reporting methods over time, the final growth-rate series will systematically contain (at least) one-year gaps at each valuation break because each segment originates from a distinct unconsolidated source. To address this issue and prevent the extrapolation process from stopping at each missing observation, we assume constant labour productivity to bridge these minor gaps in the growth-rate series. We make this assumption in line with Pahl and Timmer (2020) and apply it only when two growth-rate series can be bridged. In other words, if two value-added growth-rate series corresponding to different reporting methods are separated by a one-year gap in the final combined value-added growth rate series, we assume that the missing growth rate for that year equals the growth rate of employment in the same year. The same logic applies to output series, which rely on the growth rate of employment to fill one-year gaps (i.e., assuming constant labour productivity per worker for output). Alternatively, if a one-year gap occurs in the employment growth-rate series, we assume that the missing growth rate equals the growth rate of value added, without relying on output. This choice reflects the closer conceptual relationship between employment and value added.

In the current version of the STiM database, we tolerate up to three missing values between two growth-rate series. This means that we assume, at most, three consecutive years of constant labour productivity within the same series to bridge such gaps. This assumption is introduced to prevent the extrapolation process from being interrupted too frequently, particularly for countries that frequently change their reporting methods or take time to adapt their reporting systems. If more than three consecutive years are missing, the constant labour productivity assumption is not applied, meaning that the remaining part of the series is lost. In addition, it also implies that if a country changes in the same year the valuation methods for value-added/output and employment, no bridging is possible since no alternative variable is available.

As an example, this was the case for Paraguay in 2010. The national statistical institute switched from measuring employment as the number of employees to the number of persons engaged, while simultaneously changing the valuation of value added and output from producers' prices to an unspecified method. As a result, no bridging is possible in this case, since the growth rate of employment cannot be used to fill the gap in value-added or output growth rates, and vice versa. Importantly, the labour productivity computed from 2010 onward — based on the new definitions of both employment and valuation methods — cannot be used to link pre-2010 and post-2010 segments, as doing so would artificially merge two incompatible measurement levels and bias the resulting chained series. We therefore cannot recover values after 2009 in the absence of external datasets whose growth rates could be used to fill the gap for at least that year. If such growth rates were available, it would then be possible to use those from INDSTAT in subsequent years.

Although the loss of some series separated by more than three years cannot be entirely avoided, we recall that we mitigate this issue by selecting the reference year within the longest continuous non-missing streak of the series. This approach prevents us from relying solely on valuation-method preferences, which could otherwise result in very short final series when the preferred valuation is reported for only a few years (e.g., Eswatini's case). Overall, in our final estimates, assumptions of constant labour productivity corresponding to only 0.7% of observations in the final employment growth-rate series, 0.9% in the final value-added growth-rate series, and 1% in the final output growth-rate series. Nonetheless, for transparency, three tag variables are created to identify, for each final series, the values estimated under the assumption of constant labour productivity growth. These tag variables are included in the *detailed version* of the STiM database.

4.3.4 Backward and forward extrapolation

The final step of the chain-linking procedure involves reconstructing the consolidated series by extrapolating both backward and forward from the reference year using the combined growth-rate series. The reference-year value (X_{t_0}) for each of the three consolidated series is always taken from the INDSTAT database. The extrapolation is performed as follows, where g_t denotes the growth rate stemming from the combined employment, value-added, and output growth-rate series:

$$X_t = \begin{cases} X_{t-1} \times (1 + g_t), & \text{for } t > t_0 \quad (\text{forward extrapolation}) \\ X_{t+1} \div (1 + g_{t+1}), & \text{for } t < t_0 \quad (\text{backward extrapolation}) \end{cases} \quad (1)$$

4.3.5 Additional reference years to recover extra observations

Once this first extrapolation is complete, we perform additional extrapolations by allowing a given country to have several reference years if the initial streak is interrupted by a series of missing values. However, each new reference year must be consistent with the valuation method of the first one in order to preserve internal and intertemporal consistency. To be more precise, we adopt as many additional reference years as there are complete breaks common to all three consolidated variables in a given country. Recall that this procedure is performed after an initial attempt to avoid these gaps through the CLP hypothesis and the initial five-year interpolation. Each new reference year is locked to the same configuration (employment concept and price/valuation) as the original reference year and is chosen once per complete post-break streak. Within that streak, we apply the same selection filters as for the first reference year. We then chain-link starting from this new anchor using growth rates within that streak.

If, in a given streak, the exact configuration is not simultaneously available for all three variables, we nonetheless define a single common reference year for the streak, anchoring it to the largest admissible subset according to the following priority order. Suppose two out of three variables share the same configuration as the original reference

year. In that case, we prefer to recover (i) employment and value added if available, followed by (ii) employment and output, and lastly (iii) output and value added. If no pair is available, we allow single-variable anchors in the following preference order: (i) employment, followed by (ii) value added, and then (iii) output. In all cases, the reference year is unique and applies simultaneously to the three variables within the same streak, ensuring intertemporal consistency. In all cases, the reference year is unique and applies simultaneously to the three variables within the same streak, ensuring internal consistency. This is crucial to avoid distortions in the series, as failing to respect this condition would misalign levels and distort key ratios — notably labour productivity when measured as value added per worker for example. It is therefore impossible to assign a reference year to only one variable, while the series for another variable continue from the previous chain-linking. Put differently, additional reference years are introduced only when a total break occurs — namely, a common discontinuity in employment, value added, and output that cannot be bridged by interpolation or the CLP hypothesis.

The case of Tunisia illustrates why we aim to implement this procedure. In Tunisia, there is a common break across employment, value added, and output from 1982 to 1988 that cannot be bridged by interpolation or CLP. As such, we initially end up with two separate streaks: 1963-1981 and 1989-2019. Given the algorithm used to select the reference year, the first anchor is set in the longest streak (1989-2019) and corresponds to 2002 which maximise the coverage and valuation methods. In that year, value added and output are reported at basic prices, while employment is measured as the number of persons engaged. To avoid losing the whole 1963-1981 series, we introduce a second reference year for the earlier streak, locking it to the same configuration as in 2002 — producers' prices and number of employees. Since several matching years are available, we then choose among these years the one with the highest count of non-missing observations and the most recent one in this streak, which is 1981. This procedure recovers the whole 1963-1981 block while preserving internal consistency (the same valuation/definitions within each anchor) and intertemporal consistency with the most recent streak. It thus extends the usable series without mixing configurations.

4.4 Ex-post quality checks

The last step of the chain-linking procedure involves performing quality checks to ensure that the final series remain broadly consistent with the initial data whenever the latter are available (i.e., values just before chain-linking). One drawback of this procedure is that, in some cases, values may differ from their initial level not because of differences in valuation methods, but because growth rates can occasionally distort the series. This effect tends to be amplified the further the observation is from the reference year (eg., Fuel and gas sector in Thailand).

As such, to ensure that the final series remain broadly consistent with the initial data, we perform two ex-post quality checks. First, when computing the ratio between the

final and initial series, we set the raw data to missing whenever the ratio falls outside the $[0.5, 2.0]$ interval, meaning we tolerate at most a halving or doubling relative to the original values. Second, we discard observations whose internal shares differ by more than ± 3 percentage points compared to the shares computed from the raw data. These thresholds are somewhat quite arbitrary, but they are designed to strike a balance between preserving the integrity of the original data and ensuring the reliability of the final series. When doing such test, we loose less than 3% of observations in each final consolidated variable.

5 Deflating procedure

This first subsection aims in describing the various sources of price indexes used to deflate nominal value-added and output data in the STiM database. The second subsection details the deflation methodology applied to link all these sources.

5.1 Price indexes sources

The series for value added and nominal production are deflated using price indices organised hierarchically according to their degree of sectoral specificity and conceptual proximity to the preferred valuation methods adopted in the STiM database. Priority is always given to deflators derived from INDSTAT, the OECD (STAN), and KLEMS. In their absence, we rely on more aggregated indices — namely, producer price (PPI) or wholesale price indices (WPI) — either at the manufacturing aggregate or country level. When these are unavailable, GDP implicit deflators from the UNSD and the World Bank are used as proxies, particularly for the earliest years.

5.1.1 INDSTAT and the Index of Industrial Production (IIP)

Natively, INDSTAT reports value added and output in nominal terms. To disentangle price from quantity movements, we follow Haraguchi and Amann (2023) recommendations. We therefore exploit the Index of Industrial Production (IIP) embedded in the INDSTAT database, which measures real production growth by manufacturing division relative to the 2015 benchmark. To be more precise, let us denote nominal output O for country i , sub-sector s , and year t as the following product between prices and quantity.

$$O_{ist} = P_{ist} \times Q_{ist} \quad (2)$$

If we consider that $\tilde{O}_{ist}$ denotes constant-price output valued at base-year prices P_{isb} for a given baseline b , we can express the IIP as follows:

$$\text{IIP}_{ist} = \frac{\tilde{O}_{ist}}{O_{isb}} = \frac{P_{isb} \times Q_{ist}}{P_{isb} \times Q_{isb}} = \frac{Q_{ist}}{Q_{isb}}. \quad (3)$$

Hence, the relative output price is identified from nominal output and the IIP:

$$\frac{P_{ist}}{P_{isb}} = \frac{O_{ist}/O_{isb}}{IIP_{ist}} \quad (4)$$

As such, retrieving the embedded price index (PI) at a base year b can be done in the following way:

$$PI_{ist} = 100 \times \frac{O_{ist}/IIP_{ist}}{O_{isb}/IIP_{isb}} \quad (5)$$

Note that PI_{ist} is computed before any ex-post cleaning or re-aggregation is done so that the decomposition respects the original accounting identity. In other words, PI_{ist} is retrieved using the initial consolidated raw output series, which mixes different valuation methods. While this preserves coherence with the initial framework, the resulting index is later applied to the cleaned nominal output and value-added series.

In this single-deflation setting, we set $b = 2015$ for all countries and sub-sectors, as it is the native reference year of the IIP index. Since the IIP is initially reported for the 22 original ISIC manufacturing sub-sectors, the next step implies aggregating these indices to match the STiM classification into 12 broader sub-sectors. This aggregation is carried out by computing a weighted average of the individual price indices, where the weights correspond to the nominal output of each original sub-sector in each year. As such, if we note S one of the new STiM sub-sectors composed of a set of ISIC sub-sectors j , the price index for a STiM sub-sector S in year t is given by:

$$PI_{S,t} = \frac{\sum_{j \in S} O_{j,t} \times PI_{j,t}}{\sum_{j \in S} O_{j,t}} \quad (6)$$

This rule is applied for S1 (15–16), S2 (17–19), S4 (21–22), S9 (27–28), S10 (29–33), and S11 (34–35). For S12 (36–37), a weighted average cannot be performed since ISIC 37 is not covered by the IIP. As a result, we apply the price index of ISIC 36 to the entire group. The remaining sectors (S3, S5, S6, S7, and S8) correspond directly to a single ISIC sub-sector, so their price index is directly inherited from the corresponding ISIC sub-sector. More details about the aggregation are provided in the previous section (i.e., Table 3).

However, a complete weighted average is not always possible. To illustrate, consider the STiM sub-sector S2, which aggregates three underlying ISIC sub-sectors (17, 18, 19). In theory, the final price index for S2 would require the three corresponding price indices and the output for each of these sub-sectors. Yet, if only two price indices (or output values) out of the three required to aggregate correctly into S2 are available, we compute a partial weighted average over those two. If only one is available out of the three, we use that one as a proxy for the whole sector. In other words, when a complete weighted average cannot be computed, we assume that the missing sub-sectors j within S exhibit the same price dynamics as the available ones. If none is available, we use the next-best alternative source of deflator (see below).

Note that for tractability purposes, we provide a tag variable that documents whether the price index for a given sector is based on a full or partial weighted average. This tag is only available when the deflator used is derived from the IIP. The tag is only available in the *detailed version* of the STiM database.

5.1.2 OECD and EU-KLEMS databases

When INDSTAT's IIP implicit deflators are unavailable, we turn to the OECD STAN and EU-KLEMS price indices (Horvát and Webb, 2020; Bontadini et al., 2023). Both sources provide deflators for value added and output series at the two-digit ISIC level. Compared with INDSTAT, these deflators are series-specific, allowing, when available, for the implementation of a double-deflation method. Yet, the coverage of these databases is more limited in terms of countries and years, as discussed previously, so that, in practice, only very few observations are ultimately deflated using distinct price indices. To match the STiM sub-sectors, we apply the same aggregation procedure as described above for INDSTAT and rebase all indices to 2015.

All the remaining sources of deflators described below are used to extend the series backwards in time, but they cease to be sub-sector specific. They are therefore used as proxies for the unobserved real price dynamics of the various manufacturing sub-sectors.

5.1.3 OECD Manufacturing PPI Price Indexes

In the initial release of OECD (2016), specific manufacturing Producer Price Indices (PPIs) were provided for OECD countries. These deflators are reported at the level of the aggregate manufacturing industry, meaning that no distinction is possible across sub-sectors. When available — and when the previous sources are not — these indices serve as proxies for the overall price dynamics of manufacturing and are applied uniformly across all sub-sectors. The index is rebased to 2015.

5.1.4 IMF PPI and WDI Price Indexes

When the previous sources are unavailable, we rely on the Producer Price Index (PPI) and the Wholesale Price Index (WPI) from the IMF's International Financial Statistics (IFS) database (International Monetary Fund, 2024). Both sources provide country-level price indices. The PPI captures price changes at the producer level, typically reflecting factory prices of domestically produced goods, while the WPI measures wholesale transaction prices. Consequently, the PPI is preferred as a more accurate proxy for industrial producer price dynamics. However, in both cases, these indices also reflect price movements from other sectors of the economy. All indices are rebased to 2015 and applied uniformly across all sub-sectors.

5.1.5 Implicit GDP Deflators from UNSD and the World Bank

Lastly, when none of the previous sources are available, we resort to the implicit GDP deflator from the United Nations Statistics Division (UNSD) and the World Bank (World Bank, 2022). Both sources provide country-level indices, which are rebased to 2015 and applied uniformly across all sub-sectors. While these indices are the least suitable for deflating manufacturing value added and output, they are often the only available source for many developing economies, particularly for earlier years. When both sources are available, we prioritise the UNSD deflator over the World Bank's, as the former is generally more standardised and harmonised across countries. In both cases, the price indices reflect price fluctuations in each country's local currency.

For each source, a tag variable is included in the database to identify the deflator used for each country, sub-sector, and year. This tag is only available in the *detailed version* of the STiM database.

5.2 Deflation procedure

To link all these price indices into a single index for both value added and output, we chain-link them from 2015 (the baseline) using the growth rate of each selected index. The hierarchy of growth rates matches that described above. In other words, when available, the growth rate of the IIP-based index is used to extend the series backwards and forwards from 2015. When it is not available, we use the growth rate of the OECD/EU-KLEMS index, followed by the OECD manufacturing PPI, the IMF PPI/WPI, and finally the GDP deflator from UNSD and the World Bank. The extrapolation is performed for each country i , sub-sector s and year t . It is conducted separately for value added (VA_{ist}) and output (Y_{ist}) since OECD STAN and EU-KLEMS provide distinct indices for each series. The extrapolation matches equation 1.

Once the chained index is constructed, it is applied to the cleaned nominal value-added and output series following the cleaning and chain-linking procedures describes in section 4. As such, to obtain real value-added and constant output:

$$VA_{ist}^{\text{real}} = \frac{VA_{ist}^{\text{nom}}}{PI_{ist}^{\text{VA}}} \times 100 \quad ; \quad Y_{ist}^{\text{real}} = \frac{Y_{ist}^{\text{nom}}}{PI_{ist}^{\text{Y}}} \times 100 \quad (7)$$

When looking at the distribution of value-added deflators, the majority of observations are sourced from INDSTAT, which together accounts for around 40.2% of all entries — 30.57% from complete weighted averages and 9.60% from partial weighted averages. A substantial share, 29.73%, relies on the UNSD GDP deflator, which fills numerous data gaps, particularly for developing economies and earlier periods. Smaller but significant contributions come from the IMF PPI (9.03%) and IMF WPI (6.81%), which serve as alternative proxies when industrial-specific deflators are unavailable. The OECD PPI at the manufacturing level represents 4.71% of observations, while the OECD STAN accounts for 4.15%, and the EU-KLEMS database contributes a marginal

0.35%, reflecting its more limited temporal and geographical coverage. Finally, 2.84% of the observations use the World Bank GDP deflator. The proportions are nearly the same for output deflators.

6 Final database

At the end of the cleaning and chain-linking procedure, we obtain a final database covering 12 manufacturing sub-sectors across more than 145 countries from 1963 to 2019. The final STiM database contains 56,801 observations for value added, 58,574 for output, and 58,414 for employment (i.e., Table 8). We are able to provide real value added and constant-price output for 45,641 and 47,366 observations, respectively. The deflation procedure thus allows us to obtain real estimates for around 80% of the total observations, considering both value added and output. Yet, it should be noted that all the cleaning, aggregation, and harmonisation procedures described in the previous sections lead to a loss of approximately half of the initial observations when compared with the original raw data at the 22-sector level (i.e., Table 2). After the chain-linking

Table 8: Statistics regarding valuation methods in the STiM Database

Configuration	Value added		Output		Employment	
	N	%	N	%	N	%
Basic prices	13,375	23.55	11,409	19.48	—	—
Producers prices	23,639	41.62	28,061	47.91	—	—
Factor values	16,007	28.18	10,810	18.46	—	—
Unknown prices	3,780	6.65	8,294	14.16	—	—
Number of employees	—	—	—	—	56,911	97.43
Number of engaged	—	—	—	—	1,503	2.57
Total	56,801	100.00	58,574	100.00	58,414	100.00

procedure, the final valuation structure of the STiM database shows that most observations are valued at producers' prices (41.62% and 47.91% of total observations for value added and output, respectively). These are followed by basic prices (23.55% and 19.48%) and factor cost valuations (28.18% and 18.46%). The remaining observations (6.65% and 14.16%) correspond to cases where the valuation method is unknown. In other words, despite the efforts to prefer basic prices for both value added and output, the most common valuation method in the final database is producers' prices. This is partly explained by our choice to restrict the initial reference year to the longest initial streak of observations in the original algorithm designed to select the optimal reference years. While these cross-country differences in valuation methods may weaken the international consistency of the database, any resulting bias is now country-specific and time-invariant thanks to the chain-linking that preserves initial level differences. Whatever its sign, such bias would therefore be theoretically captured by including country fixed effects in any econometric analyses. Accordingly, we strongly recom-

mend that all empirical analyses using the STiM database include country dummies to account for level differences arising from valuation methods. Regarding employment, almost all observations refer to the number of employees (97.43%), while only 2.57% refer to the number of persons engaged.

7 Conclusion

This paper sets out to present the construction of the STiM database, which provides employment, value-added, and output data at the two-digit manufacturing sub-sectoral level for more than 145 countries from 1963 to 2019. The database is the result of combining several data sources and undertaking substantial harmonisation efforts aimed at ensuring internal, intertemporal, and international consistency. In addition, it links several price deflators from multiple institutions to adjust nominal measures and produces consistent estimates of real value-added and real output. To date and to our knowledge, this makes the Structural Transformation in Manufacturing database (STiM) the most comprehensive dataset for long-run analyses of structural change within manufacturing in both developed and developing economies.

Besides contributing to the existing literature by extending previous academic efforts to compile and harmonise macroeconomic statistics for a wide range of countries, the database also aims to promote further research avenues on how manufacturing has evolved over the last six decades. Moreover, it might be also used to explore how manufacturing can still be leveraged in our modern economies to foster inclusion, growth, and development.

Finally, we recall that the final database comes under two versions. The *ready-to-use version* contains the five core variables discussed in the paper (employment, nominal value added, real value added, nominal output, and real output), along with metadata identifying the country, year, and relevant manufacturing sub-sector. This release is also complemented by a *detailed version*, which only includes tag variables designed to enhance the transparency and tractability of the harmonisation and estimation procedures. The database and the replication package will be made freely available through the corresponding GitHub repository, enabling researchers to replicate our results, build upon our work, and contribute to the improvement of the database. This also serves to reaffirm that the current release is only the first iteration of the STiM database (V1.0), and we welcome feedback and suggestions for future updates and enhancements. Additionally, further data sources are planned to be integrated as they become available, thereby enhancing both the quality and coverage of the database.

References

- Aiginger, K. and D. Rodrik (2020). 'Rebirth of Industrial Policy and an Agenda for the Twenty-First Century'. In: *Journal of Industry, Competition and Trade* 20.2, pp. 189–207.
- Allen, R. C. (2011). *Global economic history: a very short introduction*. Vol. 282. Oxford University Press, USA.
- Andreoni, A. and F. Tregenna (2020). 'Deindustrialisation reconsidered: Structural shifts and sectoral heterogeneity'. In: .
- Autor, D. H., D. Dorn and G. H. Hanson (2013). 'The China Syndrome: Local Labor Market Effects of Import Competition in the United States'. In: *American Economic Review* 103.6, pp. 2121–2168.
- Baldwin, R. (2011). 'Trade And Industrialisation After Globalisation's 2nd Unbundling: How Building And Joining A Supply Chain Are Different And Why It Matters'. In: *NBER Working Papers*.
- Bekhti, Y. (2025). *Manufacturing in Structural Change: Patterns and Internal Reconfigurations*.
- Bogliaccini, J. A. (2013). 'Trade Liberalization, Deindustrialization, and Inequality: Evidence from Middle-Income Latin American Countries'. In: *Latin American Research Review*, p. 28.
- Bontadini, F., C. Corrado, J. Haskel, M. Iommi and C. Jona-Lasinio (2023). 'EUKLEMS & INTANProd: industry productivity accounts with intangibles'. In: *Sources of growth and productivity trends: methods and main measurement challenges*, Luiss Lab of European Economics, Rome.
- Cantore, N., M. Clara, A. Lavopa and C. Soare (2017). 'Manufacturing as an engine of growth: Which is the best fuel?' In: *Structural Change and Economic Dynamics* 42, pp. 56–66.
- Dasgupta, S. and A. Singh (2007). 'Manufacturing, Services and Premature Deindustrialization in Developing Countries: A Kaldorian Analysis'. In: *Advancing Development*. Ed. by G. Mavrotas and A. Shorrocks. London: Palgrave Macmillan UK, pp. 435–454.
- Diao, X., M. Ellis, M. McMillan and D. Rodrik (2025). 'Africa's Manufacturing Puzzle: Evidence from Tanzanian and Ethiopian Firms'. In: *The World Bank Economic Review* 39.2, pp. 308–340.
- Diao, X., M. McMillan and D. Rodrik (2019). 'The Recent Growth Boom in Developing Economies: A Structural-Change Perspective'. In: *The Palgrave Handbook of Development Economics*. Ed. by M. Nissanke and J. A. Ocampo. Cham: Springer International Publishing, pp. 281–334.
- Dosi, G., F. Riccio and M. E. Virgillito (2021). 'Varieties of deindustrialization and patterns of diversification: why microchips are not potato chips'. In: *Structural Change and Economic Dynamics* 57, pp. 182–202.
- European Commission, ed. (2008). *NACE Rev. 2: statistical classification of economic activities in the European Community*. Luxembourg: Publications Office. 1 p.
- Felipe, J., A. Mehta and C. Rhee (2019). 'Manufacturing matters... but it's the jobs that count'. In: *Cambridge Journal of Economics* 43.1, pp. 139–168.

- Forero, D. and A. Tena-Junguito (2024). 'Industrialization as an engine of growth in Latin America throughout a century 1913–2013'. In: *Structural Change and Economic Dynamics* 68, pp. 98–115.
- Haley, U. C. V. and G. T. Haley (2013). *Subsidies to Chinese Industry: State Capitalism, Business Strategy, and Trade Policy*. Oxford University Press.
- Haraguchi, N. and J. Amann (2023). 'Expanded Real Value Added Data for Manufacturing: A New Approach to Measuring Sub-Sectoral Manufacturing Development'. In.
- Haraguchi, N., B. Martorano, M. Sanfilippo and A. Shingal (2019). 'Manufacturing growth accelerations in developing countries'. In: *Review of Development Economics* 23.4, pp. 1696–1724.
- Hirschman, A. O. (1958). *The strategy of economic development*. Yale studies in economics. New Haven: Yale University Press. 217 pp.
- Horvát, P. and C. Webb (2020). 'The OECD STAN Database for industrial analysis: Sources and methods'. In: *OECD Science, Technology and Industry Working Papers*.
- International Monetary Fund (2009). *System of National Accounts 2008*. Washington, D.C: International Monetary Fund. 1 p.
- (2024). *International Finance Statistics (IFS) (April 2024 Edition)*. Version 1.0.
- (2025). *System of National Accounts 2025*. Washington, D.C: International Monetary Fund.
- Kaldor, N. (1966). *Causes of the slow rate of economic growth of the United Kingdom: an inaugural lecture*. University of Cambridge. Inaugural lectures. London: Cambridge University Press. 40 pp.
- Kruse, H., E. Mensah, K. Sen and G. de Vries (2022). 'A Manufacturing (Re)Naissance? Industrialization in the Developing World'. In: *IMF Economic Review*.
- Kuznets, S. (1973). 'Modern Economic Growth: Findings and Reflections'. In: *The American Economic Review* 63.3, pp. 247–258.
- Lautier, M. (2024). 'Manufacturing still matters for developing countries'. In: *Structural Change and Economic Dynamics* 70, pp. 168–177.
- Lewis, W. A. (1954). 'Economic Development with Unlimited Supplies of Labour'. In: *The Manchester School* 22.2, pp. 139–191.
- Nguimkeu, P. and A. Zeufack (2024). 'Manufacturing in structural change in Africa'. In: *World Development* 177, p. 106542.
- OECD (2003). *OECD Science, Technology and Industry Scoreboard 2003*. Paris: Organisation for Economic Co-operation and Development.
- (2016). *Main Economic Indicators - complete database*.
- Pahl, S. and M. P. Timmer (2020). 'Do Global Value Chains Enhance Economic Upgrading? A Long View'. In: *The Journal of Development Studies* 56.9, pp. 1683–1705.
- Palma, J. G. (2005). 'Four sources of de-industrialisation and a new concept of the Dutch Disease'. In: *Beyond Reforms: Structural Dynamics and Macroeconomic Vulnerability*. Ed. by J. A. Ocampo. Vol. 3. Stanford, CA: Stanford University Press, pp. 71–116.
- Prebisch, R. (1950). *The Economic development of Latin America, and its principal problems*. Lake Success: N.Y. : United Nations, Department of economic affairs. 59 pp.

- Rodrik, D. (2013). 'Unconditional Convergence in Manufacturing'. In: *The Quarterly Journal of Economics* 128.1, pp. 165–204.
- (2016). 'Premature deindustrialization'. In: *Journal of Economic Growth* 21.1, pp. 1–33.
- Szirmai, A. (2012). 'Industrialisation as an engine of growth in developing countries, 1950–2005'. In: *Structural Change and Economic Dynamics* 23.4, pp. 406–420.
- Szirmai, A. and B. Verspagen (2015). 'Manufacturing and economic growth in developing countries, 1950–2005'. In: *Structural Change and Economic Dynamics* 34, pp. 46–59.
- Tregenna, F. (2016). 'Deindustrialization and premature deindustrialization'. In: *Handbook of Alternative Theories of Economic Development*, pp. 710–728.
- UNIDO (2024). *INDSTAT, Revision 3*. Version 2024.
- United Nations (1968a). *International Standard Industrial Classification of All Economic Activities (ISIC), Rev.2*. Statistical Papers (Ser. M). New York: UN.
- (1968b). *System of National Accounts 1968*.
- (1989). *International Standard Industrial Classification of All Economic Activities (ISIC), Rev.3*. Statistical Papers (Ser. M). New York: UN.
- (2002). *International Standard Industrial Classification of All Economic Activities (ISIC), Rev.3.1*. rev. 3.1. Statistical Papers (Ser. M) 4, rev.3.1. New York: UN.
- (2008). *International Standard Industrial Classification of All Economic Activities (ISIC), Rev.4*. Statistical Papers (Ser. M). s.l: United Nations. 1 p.
- Vaz, P. H. and W. Baer (2014). 'Real exchange rate and manufacturing growth in Latin America'. In: *Latin American Economic Review* 23.1, p. 2.
- Vu, K., N. Haraguchi and J. Amann (2021). 'Deindustrialization in developed countries amid accelerated globalization: Patterns, influencers, and policy insights'. In: *Structural Change and Economic Dynamics* 59, pp. 454–469.
- World Bank (2022). *World Development Indicators (WDI) (2022Q1 Edition)*. Version 1.0.