

Regional and dynamical effects of the "digitalization" phenomenon under an mAIcroscope based on Google Trends data – FULL TEXT

László Pitlik (0000-0001-5819-0319), László Pitlik Jr. (0000-0002-8058-9577), Dániel Váradi (0000-0001-9610-8566) - KJU, Hungary

ABSTRACT

Keywords: similarity analysis, benchmarking, comparison, automation

History/context: The department of computer science (focusing on AI) already presented 2 papers (2025) using the same methodology¹. A third one could be evaluated as a less innovative approach; however, even this context free process demonstrates the real potential of the non-LLM-based AI (c.f. footnote[7]).

Aims and methodology: Here and now, the first Panel (digital transformation) is the focus. The authors propose the following basic hypothesis: each country (10+) and/or each year (2004-2025) could be evaluated as the same. This means that each country/year could lead to the same **digitalization-index**. The (own) anti-discrimination-based similarity analysis², makes it possible to detect optimized models (staircase-functions) to determine which countries/years have higher/lower **digitalization-indices** (let alone: which countries/years are norm-like or not-evaluable based on the available dataset). The benchmarks are subjectively weighted scoring models. AI needs qualitative and quasi unlimited data. Google Trends can be seen as one of the sources capable of meeting these expectations. Google Trends data have regional and time-series-oriented dimensions. In parallel, the Conference³ and especially the organizers⁴ represent a region-oriented challenge. The focus of the Conference⁵ highlights the importance of comparative approaches. These parameters are familiar from historical pre-conditions. The Conference has 7 panels⁶. Each could be addressed using the objectivity-oriented methodology previously applied in Hungary.

Targeted groups/Utility - Results: Moody's, Fitch Ratings and other think-tanks create, for example country-ranking-solutions based on subjective steps,

¹ (c.f. see Taylor-Swift: <https://miau.my-x.hu/miau2009/index.php3?x=e0&string=taylor> / see IT-security <https://miau.my-x.hu/miau2009/index.php3?x=e0&string=III20>)

² (c.f. https://miau.my-x.hu/miau/196/My-X%20Team_A5%20fuzet_EN_jav.pdf)

³ (c.f. 13TH IIMS INTERNATIONAL CONFERENCE, 2025 - <http://iimsconference.com/Conference/13th-iims-international-conference-2025>)

⁴ (DE, IN, LA, LK, VN, MY, KH, BD, NZ + HU as guest) ←Google Trends codes for regions

⁵ (c.f. Influence of Technology, Governance and Culture in Building Societies: **Global Experiences**)

⁶ (Panel A: Digital Transformation in Socio-Economic Development / Panel B: Advancement of Technology in Building Human Capital / Panel C: Management Skills in Business Development, Particular Reference to MSME / Panel D: Good Governance – A Prerequisite for Social Balancing / Panel E: Cultural Diversity in Social Development / Panel F: Technological Innovations in Health Care and Environmental Management / Panel G: Responsible Management Practices & Sustainable Society)

weights, and parameters. However, there is an appropriate online analytical tool⁷, capable of delivering generally acceptable similarities between objects. **These objective similarities make it possible to avoid double standards and other forms of subjective distortions in economics and/or in social challenges.**

INTRODUCTION

Human experts are capable of generating expertise quasi for virtually all challenges – and such expertise is always a product of human intuition processes. Large Language Models (LLMs), such as ChatGPT, Copilot, etc., appears to do seemingly the same: they generate texts (i.e., expertise) based on the probabilistic machine intuition processes. However, machine intuition depends on human textual patterns.

The problem in both cases is simple: the resulting expertise is often of low quality compared to ideal cases where each part of the expertise is derived in a data-driven and consistency-oriented manner.

In this article (following the conference instructions), the authors focus on GEO-characteristics based on previous scientific activities (c.f. Chapter: Review of Literature).

REVIEW OF LITERATURE

The authors (in different constellations) focused on 3 GEO-levels across 5 projects: settlement-level (cf. Budapest-project), regional level (cf. Mezőföld-project), country-level (cf. EU-homogeneity-project, Taylor-Swift-project, IT-security-project). In detail:

The Budapest-project identified which districts could or should be excluded from Budapest and which neighbouring settlements could or should be included to Budapest, in order to demonstrate AI-competencies based on statistical data: cf. VÁRADI, D., PITLIK, L., (2025).

The Mezőföld-project applied similar exclusions and/or inclusions, but in the case of a non-administrative region, “Mezőföld”, to create a regional robot-expert: cf. VÁRADI, D., KULCSÁR, L., PITLIK, L., (2024).

⁷ (e.g. component-based object comparison for objectivity - https://miau.my-x.hu/myx-free/index_en.php3) - LLM-solutions (nowadays called as AI-solutions) are not capable of creating objectivity-oriented models, because the corpus (written sources) behind the complex learning processes are irrational, corrupt – as the human thinking in general. LLMs are only a part of AI. Numeric input-output processes are more robust especially they have own QA-layers...

The EU-project derived homogeneity-index-values for countries and years to enable objective discussion of complex and/or abstract phenomena such as homogeneity: cf. VÁRADI, D., PITLIK, L., (2023).

The Taylor-Swift-Project constructed a measurement scale for similarities across 55 countries and 21 years based on the keyword “Taylor Swift” in order to build a benchmark compared to human-expertise-based approaches used during the course: cf. PITLIK, L. (2025).

And finally, the IT-security project (quasi) replicated the methodology of the Taylor-Swift-project to demonstrate which countries and/or time periods can be identified where the societies of the Google-observed countries focused more intensively on the term “IT-security”: cf. PITLIK, L., RIKK, J. (2025).

Behind all these projects, the author’s own AI-development (called similarity analysis – COCO: component-based object comparison for objectivity) was integrated into the analytical processes to ensure optimized and anti-discrimination-based derivations of values, that prior to appropriate mathematical modelling, could only be imagined and executed through human intuition. However, human intuition processes are never capable of delivering arbitrarily high levels of consistency in terms of the logical interpretation quality. The same risk can/must also be considered in case of LLM solutions...

The previous projects and their associated literatures help clarify the details necessary for reproducibility. One general overview should be highlighted here and now regarding the roots of similarity analysis: cf. PITLIK, L. (2014): Similarity analysis is a group of algorithms that enables the derivation of production functions based on staircases, anti-discrimination oriented models, and explorative models. All these AI-oriented optimization modules of the online tool (<https://miau.my-x.hu/myx-free/>, <https://miau.my-x.hu/myx-free/coco/index.html>, <https://miau.my-x.hu/myx-free/index.php3?x=e0>) can be used freely for limited OAM (object-attribute-matrix) versions.

OBJECTIVES OF THE STUDY

As demonstrated by the highlighted elements of the previous projects, it is possible to construct robot-experts for elementary questions/hypotheses. In the current context, the focused international conference with a lot of interesting countries and conference-keywords such as comparison, naturally encourages the adaptation of the previous methodological experiences. The goal is to derive a digitalisation index for countries ultimately enabling a fully objective discussion about of different interpretation layers of digitalisation.

Why is it important to be or become objective? The answer is simple, and must be simple: the science has only one priority: to be/become objective! Knuth

(probably:-) said: “*Science is what we understand well enough to explain to a computer; art is everything else.*” In other words: we must be capable of transforming/translating our human ideas into source code!

METHODOLOGY

A publication should ideally always support the reproducibility, but this is not possible when publication space is limited. Therefore, the analytical steps can be presented here in a shortened form:

Data-asset: The keyword “digitalisation” as a by-Google-pre-interpreted topic (c.f. <https://trends.google.com/trends/explore?date=all&q=%2Fm%2F0227jd>) enables the creation of a time series (2004-2025) for each of the highlighted countries ($62=55+1+6^8$), reflecting the intensity of search activity related to this phenomenon. Each country’s time-series allows for the analysis of internal development processes. The Google Trends dataset can also be analysed using human intuition processes (c.f. Annex#1-2-3-4). However, this analytical potential can be critically limited (cf. robot-eye concept).

Directions: The higher the annual activity level for each year, the higher the digitalisation-index! This trivial expectation is the only human-defined parameter in the entire analytical process. The AI-based (anti-discrimination-based, optimized, objective) digitalisation-index is therefore a mirrored version of human intuition regarding digitalisation of countries. LLM-interpretations are also a mirrored version of the human intuitions, but the two mirroring process are fundamentally different: LLM-based solutions generate textual outputs from for textual prompts (inputs) using probabilistic calculations in the “middleware” layer. In contrast, anti-discrimination-based calculations use numerical prompts (inputs) and produce numerical outputs. The LLM-based calculations are infected/influenced by human intuitions due to their training on human-generated texts. The numerical solution, however, is pure mathematics. The numerical outputs should ideally be reformulated into human-readable texts - something that is still not fully achievable through LLMs, as the necessary textual patterns are not yet available in sufficient volume (cf. source code generation, which is already feasible due to the abundance of training data).

Comparison: The data-driven approach is based on a simplified hypothesis: Can each country be interpreted as a norm-like object in comparison to others? If not,

⁸ +1=ALL countries, 55=countries of the previous projects c.f. https://miau.my-x.hu/miau/326/digitalisation/digitalisation_de.xlsx, 6=new countries (VN, NZ, LK, LA, KH, BD) + new/old-countries: 4 (DE, HU, IN, MY) *** The 27 EU-countries should quasi always be analysed in an EU-affected project + 27 other countries – randomly chosen + USA because of the Taylor-Swift-project = 55 benchmark-countries...

an artificial similarity scale can be constructed norm-like objects on the centre. Such a scale is necessary for meaningful comparisons.

Countries (objects) with higher index values (later: over-norm-group) are considered more digitalized, while those with values below the norm (later: under-norm-group) are considered less digitalized. A fourth group includes not-interpretable objects, where function-symmetry-based quality assurance within the similarity analyses fails to establish rational relationships between mirrored input constellations. This means that some of the countries may not be evaluated based on the available raw data. Notably, this kind of self-correction mechanism is never part of the human intuition processes!

The challenge comparing country profiles using Google-Trends data can also be interpreted as a form of robot-eye-development (see Annex#1: visualized development versions). While naïve human eyes are efficient, their interpretations are often suboptimal. Naïve characteristics will be represented later through a simple aggregation form: the classic average-building – as benchmark. From a classical or statistical perspective, this comparison can also be viewed as a clustering challenge.

RESULTS & DISCUSSION

Results (see Annex#6-7-8-9): for more details see https://miau.my-x.hu/miau/326/digitalisation/digitalisation_de.xlsx

To begin, we must assess the levels of analytical risk: the naïve averages of the ranking values (see Annex#5) and the optimized similarities (calculated using COCO online-tool: Y0 module) yield a (Pearson) correlation coefficient of 0,9985. This indicates that the robot-eye and the human intuition produce relatively similar results (see later B=Basic-version).

However, differences do exist between the naïve and optimized digitalisation ranking values. The ranking range spans from 1 to 62. The maximum difference is +5, the minimum is -2. More specifically: 47 countries received the same evaluation in both the naïve and optimized approaches. The maximum difference of “+5” occurred in one single case: NO=Norway where the naïve evaluation assigned a rank of 16 while the optimized evaluation assigned a rank of 11. The minimum differences were observed in three countries: BE=Belgium, CA=Canada, and NL=Netherlands – each of which had a lower digitalisation level in the optimized version than in the naïve one.

The optimized modes (inverse and direct – based on the mirrored input ranking values) are error-free. This means that each country can be evaluated in the optimized way.

The number of the less digitalized (under-norm-countries): 35 There are no norm-like countries. Therefore, the number of over-norm-countries: 27 (35+27=62).

The under-norm-countries showed smaller differences between naïve and optimized approaches: +1 vs -1. The over-norm-countries can logically be characterized with the +5 vs -2 (extreme) interval.

Importantly, every country is correctly classified into either the over-norm or under-norm group—there are no misclassified cases.

As we can observe, countries form more and less similar groups based on the objective scale of anti-discrimination-oriented similarity analyses. Politically correct labels such as developed countries or emerging countries are not trivially applicable here and/or in general - as was also the case in previous projects. Fortunately, this allows for a more unbiased view. For example:

Brazil ranks among the top three (2nd place), ahead of the United States (3rd place) and behind France (1st place).

Germany (naïve vs optimized ranking: 12 vs 13) is not in the top group, nor is Hungary (37:37), which is also outside the over-norm group. No Central and Eastern European Country (CEEC) appears in the top group. Only 12 EU-countries are in the over-norm-group (FR>AT>BE>DK>NL>FI>DE>PT>PL>SK>IE>ES). Hungary does not have the lowest ranking value (cf. LU>BG>EE>HR>LV>GR>**HU**>IT>CZ>RO>LT>SE>SI>CY>MT). This suggests that the EU cannot be considered a homogeneous group. If the EU lacks homogeneity, it implies that global homogeneity is even more likely than commonly assumed. Significant processes are unfolding across different continents. While this observation may be less striking today - given recent global events such as the COVID-19 pandemic, the UK-RU war, the rise of BRICS countries, and political shifts in the US—it remains essential to view diversity through the lens of objective analysis, or “robot eyes,” especially since human perception often overlooks it.

Additional highlights:

Vietnam is already in the over-norm-group (27:27). India: 24:24 (c.f. China 21:21). Worth highlighting: Nigeria 18:18 and/or Mexico: 23:23

Interestingly, Japan’s consciousness regarding digitalisation, as reflected in Google Trends data, places it in the under-norm group.

All examined countries are listed in Annex #6.

The high correlation between the naïve and the optimized approach ensures that no surprising effects (e.g., countries switching between under-norm and over-norm groups) are observed. However, the relatively small differences between naïve and optimized ranking values (+5 vs. -2) occur only in the over-norm-group. As a result, assigning e.g. a “country of the year” label is not always straightforward.

Results (Annex#10):

There are four analytical versions: B-version (Basic version – see above) considers only years, T-version (Trend-version) includes trend-effects, D-version (stDeviation-version) includes standard-deviation-effects, A-version (All-effect version) combines all effects.

The B-version represented the static-cumulative status of the digitalisation. The T-version reflects the dynamic status concerning digitalisation. The D-version present the stability of digitalisation’s importance, and the A-version offers a complex, aggregated evaluation.

It becomes evident that a single effect (e.g., trend vs. standard deviation) can have significantly different impacts compared to previous versions. Moreover, the addition of new attributes does not necessarily lead to greater impact than a single attribute, due to the balancing relationships between them.

The (Pearson) correlation values between the naïve version and the B-T-D-A-versions demonstrate, why the anti-discrimination-oriented analyses are necessary: $0.99 > 0.96 > 0.49 > 0.47$. The correlations between the optimized index values are always over 0.98. However, the correlations values between the naïve B and the further (T-D-A) solutions are: $0.97 > 0.64 > 0.60$. This means that the source of the instability (volatility) should be identified in the naïve logic.

In particular, the standard deviation within country-specific time series had a substantial impact on correlation values, as it represents an entirely independent force field. Trend values are derived from raw annual Google Trends data, which is why the B-version differs less from the naïve version.

The differences between country-specific naïve and version-oriented optimized ranking (index) values become increasingly significant across the 62 countries. B-version: from +5 to -2 / T-version: from +10 to -15 / D-version: from +40 to -22 / A-version: from +37 to -32. This again confirms that the minimized standard deviation exerts a stronger influence than the maximized trend values.

Ranking differences between top and bottom groups should be evaluated differently: Low ranks (i.e., better positions) generally show smaller differences, while high ranks (i.e., worse positions) can exhibit much larger variations.

Ranking differences should always be interpreted relatively—that is, in proportion to the total number of objects (62). It is also important to note that even small differences in correlation values (e.g., 0.99 vs. 0.96) can result in substantial differences in rankings (e.g., +5/−2 vs. +10/−15).

The T–D–A versions contain progressively more over-norm objects compared to the B-version: 27 vs. 29, 30, and 34, respectively.

The D and A versions each produced one norm-like object—namely, Sweden (SE). While SE’s norm-like status is not technically valid in the A-version, norm-like objects always fall within a tolerance interval. Therefore, all objects in all model versions can be considered valid.

Country-specific details (based on Annex#10):

France (FR) and Brazil (BR) consistently rank at the top. In the A-version, Portugal (PT) is the next best EU country (FR > PT), and Poland (PL) is the highest-ranking country from the CEEC group. The United States (US) does not appear in the top group in the T-version—a result that may be interpreted as a warning signal. The US may need to take action in the near future regarding its population (cf. US-driven changes in global politics and/or domestic political dynamics – it could be said based on human intuition?!). This interpretation reflects a human-intuition-based effect, where LLMs could be used to generate associative textual descriptions based on numerical patterns in the results.

The lowest level of digitalisation consciousness, according to the A-version, is found in Hungary (HU), followed closely by Croatia (HR). Consequently, the A-version presents the most critical view of EU homogeneity, with France in first place and Hungary in last. Discussions at the EU level should be conducted using this objective framework to achieve meaningful progress toward greater homogeneity.

Some countries appear in both under-norm and over-norm groups across the four versions (especially from B-version to A-version). For example: Under-norm to over-norm: Mongolia (MN), Malta (MT), Cyprus (CY), Libya (LY), Cuba (CU), Slovenia (SI), Cambodia (KH) / Over-norm to under-norm: Mexico (MX), Germany (DE), China (CN).

The positions of countries relative to the norm value can be interpreted as optimized, objective, and aggregated fingerprints - or even as country profiles. A country’s neighbours can represent risk potential if they are associated with instability, and vice versa. This concept echoes the political principle “*Amicus meus, inimicus inimici mei*” - “The enemy of my enemy is my friend.”

This approach can be applied not only to geopolitical conflicts (cf. the IT-security project) but also to professional domains such as digitalisation

awareness. A similar effect was demonstrated in golden-age analyses (cf. <https://miau.my-x.hu/miau2009/index.php3?x=e0&string=golden>), where FAO statistics on food consumption from 1961 to 2013 were aggregated into index values that reflected good and bad years for each country—serving as a mirrored picture of political, environmental, and other challenges.

Due to space limitations, further layers of relationships between countries and model versions cannot be presented in this article.

Discussion:

The so-called digitalisation index must always be interpreted with clarity. The inputs - country-based developments in the form of time series—reflect a kind of consciousness regarding the keyword digitalisation. Other keywords with similar meanings (e.g., modelling, data, database, query, reporting, OLAP, archiving, etc.) could also be modelled. Parallel or synonymity-oriented calculations could be integrated into a broader “super-model.”

The naïve approach, represented by a simple average calculation, demonstrates that the human brain—without requiring complex mathematical capabilities - is capable of understanding reality, especially when patterns (as in this case) are long-term and suitable for pairwise comparisons.

The static interpretation presented here has been extended with dynamic effects - for example, the trend of raw or ranking values as a new attribute, where the direction “the higher, the better” applies. In this dynamic context, a country with the same average but a decreasing trend in Google Trends data may not yield the same aggregated evaluation.

While the naïve solution appears to integrate trend and standard deviation effects, this integration suffers from arbitrary distortions in ranking ratios. For instance, Excel-based rankings can introduce unnecessary shifts in the calculation. Therefore, anti-discrimination-oriented optimization produces increasingly distinct and refined solutions (cf. correlation values).

Evaluation, historically and culturally, has often/always been a subjective domain in human history. But it does not have to remain that way - and with AI advancements, it will not. Subjective interpretations (e.g., intuitions, associations) are inherently risky. While they may inspire scientific inquiry, purely journalistic interpretations (including social media comment sections) are likely more counterproductive than beneficial for individuals or groups.

Anti-discrimination-based optimization enables relatively fast and efficient evaluations. At the same time, it must be acknowledged that human intuition always incorporates all interpretable signals—not just DIGITALIZED data.

Consciousness-index regarding digitalisation must be fine-tuned in the future based on infrastructural data specific to each country. In some countries, digitalisation may already be relatively resolved; in others, it may never have been relevant. This means that Google Trends-based awareness of digitalisation must be interpreted in light of the technical and mental prerequisites still expected in each context.

CONCLUSION

Digitalisation can be interpreted through human intuition in an almost unlimited way. Words do not - and arguably never did - have fixed meanings. It would be possible to create an entirely different kind of digitalisation index, one that is abstract and complex, derived from measurable components (attributes), such as the number of computers, software licenses, individuals with ECDL certification, etc. The number of measurable attributes is theoretically unlimited.

On the other hand, existing statistics are always limited. Questionnaires form a special category of numerical input: they can always be created, but the responses (e.g., “What do you think about your own digitalisation level? $1 < 5$ ”) originate from the domain of human intuition.

The parallel models (versions) presented here demonstrate that not only human intuition is capable of fine-tuning terms and phenomena—AI is already prepared to handle such nuances.

The concurrent process - human intuition - is never deep enough to define the true meaning of a focused keyword (in this case, digitalisation). As a result, human discussions often lack direction and clarity. In contrast, computers (AIs) operate at more precise levels of abstraction. The higher the level of abstraction, complexity, and consistency, the stronger the argument or interpretation.

Science faces one overarching challenge - especially for the future: to become capable of measuring the level of abstraction, complexity, and consistency.

REFERENCES

PITLIK, L. (2014). My-X Team, i.e. an innovative Idea-Breeding Farm. Regional Innovation Agency of Central Hungary. URL: [My-X Team_A5_fuzet_EN_borito.qxd](#)

PITLIK, L. (2025). Regional and dynamical effects of the "Taylor-Swift" phenomenon under an mAIcroscope based on Google Trends data. Taylor-Swift-Course of the Kodolanyi University (2024/2025: Summer Semester), Hungary + *MIAÚ* ID_25076 ([https://miau.my-x.hu/miau2009/adatlap.php3?where\[azonosito\]=25076&mod=l2003](https://miau.my-x.hu/miau2009/adatlap.php3?where[azonosito]=25076&mod=l2003))

PITLIK, L., RIKK, J. (2025). Regional and time series patterns in the concept of IT security. International IT-security Conference of the Kodolanyi University (International Week March 20, 2025: March), Hungary + *MIAÚ* ID_25087 ([https://miau.my-x.hu/miau2009/adatlap.php3?where\[azonosito\]=25087&mod=l2003](https://miau.my-x.hu/miau2009/adatlap.php3?where[azonosito]=25087&mod=l2003))

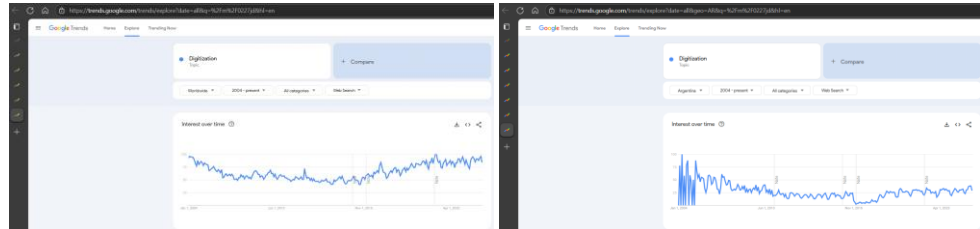
VÁRADI, D., PITLIK, L., (2025). Building statistical neurons in case of regional development projects. 13th International ZEUGMA - CONGRESS ON SCIENTIFIC RESEARCH / February 24-26, 2025 / Gaziantep, Turkiye + *MIAÚ* ID_25062 ([https://miau.my-x.hu/miau2009/adatlap.php3?where\[azonosito\]=25062&mod=l2003](https://miau.my-x.hu/miau2009/adatlap.php3?where[azonosito]=25062&mod=l2003))

VÁRADI, D., KULCSÁR, L., PITLIK, L., (2024). Automated detection of sustainability risks with AI support in the formation of regional entities. International Conference of the University Debrecen about Sustainable Economy and Sustainable Society, April 19, 2024 / Debrecen, Hungary + *MIAÚ* ID_25057 ([https://miau.my-x.hu/miau2009/adatlap.php3?where\[azonosito\]=25057&mod=l2003](https://miau.my-x.hu/miau2009/adatlap.php3?where[azonosito]=25057&mod=l2003))

VÁRADI, D., PITLIK, L., (2023). Measuring homogeneity of countries in the European Union based on similarity analyses. 5. International ANATOLIAN SCIENTIFIC RESEARCH CONGRESS July 21-23, 2023 Hakkari, Turkiye + *MIAÚ* ID_24830 ([https://miau.my-x.hu/miau2009/adatlap.php3?where\[azonosito\]=24830&mod=l2003](https://miau.my-x.hu/miau2009/adatlap.php3?where[azonosito]=24830&mod=l2003))

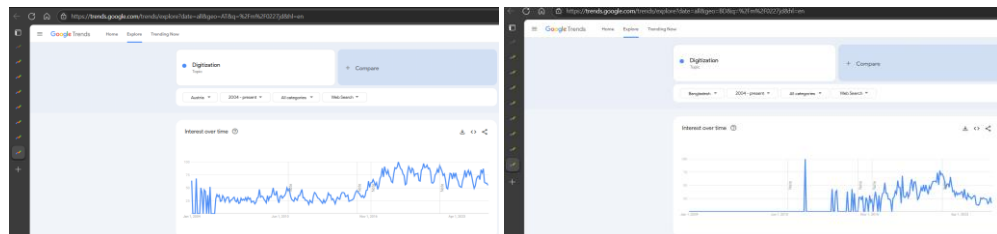
ANNEX

Annex#1: Highlighted Google Trends patterns:



Left: <https://trends.google.com/trends/explore?date=all&q=%2Fm%2F0227jd&hl=en>

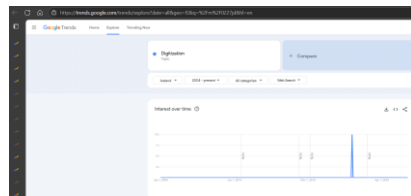
Right: <https://trends.google.com/trends/explore?date=all&geo=AR&q=%2Fm%2F0227jd&hl=en>



Left: <https://trends.google.com/trends/explore?date=all&geo=AT&q=%2Fm%2F0227jd&hl=en>

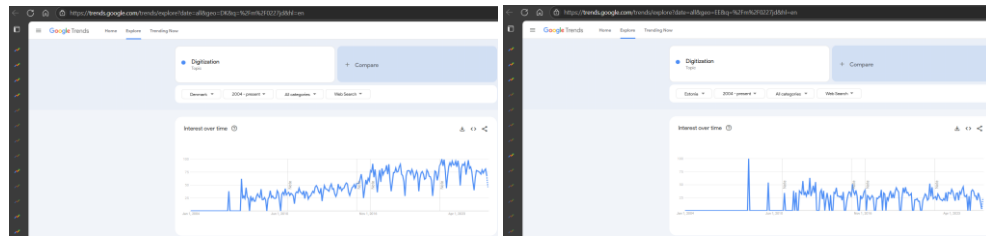
Right: <https://trends.google.com/trends/explore?date=all&geo=BD&q=%2Fm%2F0227jd&hl=en>

Legends: 4 patterns of 62 observed time series in order to be compared – Unit for the Y-axis is %, the X-axis represents the years / The robot-eye has to interpret the figures without any further instructions – only based on the direction rule (quasi prompt / strategical prompt): the higher (the percentual value pro year) the higher (the digitalisation index) / The seemingly lacking values in the first years are special inputs for the low-levelled digitalisation because the focus of digitalisation as such is not existing or even other aspects of the digitalisation (it means more exactly: the level of Internet-penetration and/or the level of the freedom using Internet at all) is alone or parallel not given ! The exploring of the relationships between pattern_ALL and pattern_country(i) will lead to a new article in the future. Hypothesis: Can we interpret each difference-structure as the same? (methodology - again: anti-discrimination-oriented modelling)



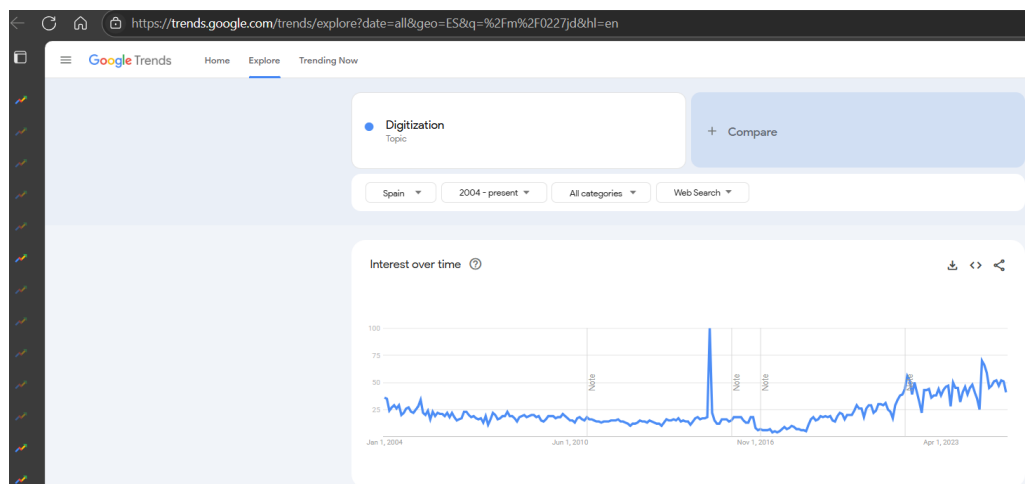
Iceland is not part of the 62 highlighted countries. But the characteristics are interesting from didactical point of view: one single peak (cover each other information unit) or quasi no data = lacking focus on digitalisation?!

Annex#2: Forecast potential of Google Trends



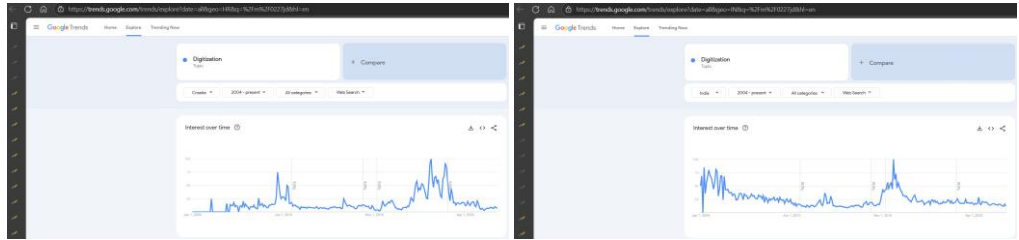
Legends: The dotted line (on the right edge) is a kind of forecast. A new article in the future can be written based (again) on similarity analyses using staircase functions for “future-production”. The different directions (DK decreasing phase, EE increasing phase) can also be used e.g. in case of short-term political discussions. The question is however: how good are these forecasts? Is Google-Trends a consolidated player (delivering forecasts only in secured cases? – why are no forecast in the Annex#1?)

Annex#3: Peaks as potential risk factors



Legends: The Google-Trend value are visualized in percentages. The 100% percent value is always the highest (not published) raw value. These 100-%-peaks can be seen as suspicious: see case BD and/or ES. Peaks can be eliminated as a kind of sensitivity analysis parallel to the general analytical process and peaks are partially eliminated based on average-building for each year and country. There is a further risky aspect in the Google Trends database: time-series of the “little” countries are more volatile, therefore peaks must be handled at any rate – like here and now based on calculations of annual averages.

Annex#4: Political waves?



Legends: The human intuitions make possible to have such associations like different political/professional periods in a country (c.f. HR, IN)?!

Annex#5: Ranked input values

ALL	2004	2005	2006	2007	2008	2009	2010	2011	2012	2013	2014	2015	2016	2017	2018	2019	2020	2021	2022	2023	2024	2025	YO	naive average	naive ranking	optimized index	optimized ranking	differences	validation	
AR	11	2	2	2	2	1	1	4	2	3	3	5	6	7	8	8	3	1	1	2	1	1	1000000	3	2	1000579	2	0	0 K.	
AT	17	10	6	6	11	16	17	20	23	24	23	25	25	42	35	30	31	35	37	36	37	31	1000000	24	20	1000112	20	0	0 K.	
AU	10	9	7	7	6	7	14	17	20	18	17	20	18	11	17	18	24	30	21	20	22	21	1000000	16	14	1000289	15	-1	0 K.	
BE	19	13	11	11	12	12	11	6	5	6	4	6	2	4	7	5	4	6	5	2	8	1000000	7	4	1000480	6	-2	0 K.		
BD	27	31	36	39	43	43	43	51	35	51	35	36	26	24	23	20	21	16	24	31	41	43	1000000	34	35	999904	35	0	0 K.	
BG	27	31	36	39	43	43	43	48	45	40	33	32	27	29	32	27	17	19	12	18	16	4	1000000	30	30	999982	30	0	0 K.	
BR	3	6	5	5	4	4	2	2	3	2	2	4	5	8	9	10	8	5	7	6	5	2	1000000	5	3	1000537	3	0	0 K.	
CA	6	4	4	4	5	6	8	8	9	8	7	8	10	14	13	14	12	14	9	8	10	12	1000000	9	6	1000450	8	-2	0 K.	
CH	18	17	19	18	19	17	16	16	18	19	16	14	9	3	5	7	9	10	18	22	23	24	1000000	15	13	1000306	14	-1	0 K.	
CN	14	14	16	22	26	26	30	26	29	29	32	47	31	43	51	23	19	7	3	4	19	28	1000000	25	21	1000105	21	0	0 K.	
CU	27	31	36	39	32	43	43	38	53	51	53	53	53	54	56	53	58	58	59	57	57	57	1000000	48	57	999578	57	0	0 K.	
CY	27	31	36	33	43	43	43	51	53	51	53	53	53	54	55	55	58	57	58	56	60	1000000	49	58	999560	58	0	0 K.		
CZ	27	28	26	20	7	10	23	22	33	33	38	38	38	39	50	48	55	55	52	51	48	51	1000000	36	42	999847	43	-1	0 K.	
DE	13	15	18	21	24	20	15	13	17	15	13	11	6	3	4	5	3	12	21	21	22	1000000	14	12	1000328	13	-1	0 K.		
DK	27	31	34	28	21	20	19	14	10	5	6	7	4	2	4	6	7	9	4	1	4	11	1000000	12	9	1000475	7	2	0 K.	
EE	27	31	36	39	32	43	37	40	15	14	21	21	20	25	25	25	36	37	40	34	36	41	1000000	31	31	999968	31	0	0 K.	
EG	27	31	36	39	39	43	43	49	53	47	50	51	48	40	34	24	27	20	27	37	40	39	1000000	38	46	999799	46	0	0 K.	
ES	15	18	22	25	27	27	29	25	28	30	26	22	29	41	39	38	34	32	31	28	26	25	1000000	28	25	1000026	25	0	0 K.	
ET	27	31	36	39	43	43	43	51	35	51	53	53	53	50	43	36	46	38	32	32	33	37	1000000	41	48	999738	48	0	0 K.	
FI	26	27	20	26	22	18	18	11	14	11	11	2	2	17	6	9	11	11	9	13	18	1000000	14	11	1000330	12	-1	0 K.		
FR	4	3	3	3	3	3	3	1	1	1	1	1	1	1	1	1	3	1	2	6	7	7	1000000	3	1	1000580	1	0	0 K.	
GB	5	11	10	10	8	5	6	18	21	21	18	18	15	12	14	17	16	23	10	6	5	1000000	13	10	1000364	10	0	0 K.		
GR	27	24	25	31	37	29	31	29	32	32	28	33	32	32	33	42	43	43	45	39	38	34	1000000	34	35	999904	35	0	0 K.	
HR	27	30	32	29	31	30	12	31	38	36	30	34	36	33	38	28	15	13	39	45	51	52	1000000	32	32	999933	32	0	0 K.	
HU	9	5	8	15	23	23	25	24	26	28	29	31	36	51	52	54	56	49	50	50	49	50	1000000	34	37	999900	37	0	0 K.	
ID	27	31	36	39	43	39	41	41	43	42	47	44	46	37	35	34	38	27	16	23	12	10	1000000	34	39	999892	39	0	0 K.	
IE	27	31	36	39	32	43	40	33	31	23	25	23	19	16	19	16	25	24	17	15	11	6	1000000	25	22	1000092	22	0	0 K.	
IN	8	7	12	14	18	22	26	27	24	26	31	30	23	9	22	31	39	39	44	42	46	47	1000000	27	24	1000056	24	0	0 K.	
IQ	27	31	36	39	43	43	43	51	42	51	53	53	53	52	53	57	50	29	27	29	33	1000000	42	51	999709	51	0	0 K.		
IL	27	31	36	39	43	38	21	46	41	38	26	24	21	20	24	19	20	17	14	25	24	23	1000000	28	25	1000026	25	0	0 K.	
IT	23	22	27	30	36	34	33	39	39	39	41	39	35	30	20	40	42	41	41	41	39	36	1000000	35	41	999876	41	0	0 K.	
JP	22	23	24	24	15	24	27	22	19	27	36	37	40	38	40	44	44	45	48	46	45	45	1000000	33	33	999908	33	0	0 K.	
KH	27	31	36	39	43	43	34	53	51	53	53	53	53	54	56	52	52	54	54	54	54	56	1000000	48	55	999594	55	0	0 K.	
LA	27	31	36	39	43	43	43	51	53	51	42	53	53	54	56	57	48	60	55	48	30	40	1000000	46	54	999626	54	0	0 K.	
LK	24	31	36	39	43	43	43	51	52	50	52	52	52	53	54	55	57	59	60	58	58	62	1000000	49	60	999555	60	0	0 K.	
LT	27	31	36	39	43	43	34	51	53	51	53	48	49	54	48	41	26	28	22	16	15	14	1000000	37	45	999821	45	0	0 K.	
LU	27	31	36	39	43	33	43	42	53	43	37	29	28	19	11	11	13	21	23	26	27	20	1000000	30	28	999988	28	0	0 K.	
LV	27	31	36	39	43	37	43	51	47	35	39	46	43	31	27	29	30	26	28	14	18	15	1000000	33	33	999908	33	0	0 K.	
LY	27	31	36	39	43	43	43	51	53	51	53	53	53	54	56	57	58	50	53	59	59	61	1000000	49	59	999560	58	1	0 K.	
MN	27	31	36	39	43	43	43	51	53	51	53	53	53	54	56	57	58	53	61	59	59	59	1000000	50	62	999551	62	0	0 K.	
MT	27	31	36	33	43	43	43	51	48	51	53	53	53	54	56	57	58	60	61	59	59	55	1000000	49	60	999555	60	0	0 K.	
MX	7	8	9	12	17	20	21	21	25	25	24	26	30	36	37	33	32	34	35	35	35	30	1000000	25	23	1000091	23	0	0 K.	
MY	27	31	29	39	38	35	39	36	37	36	43	40	41	28	28	31	28	31	34	30	31	32	1000000	34	38	999899	38	0	0 K.	
NE	27	31	36	39	43	43	43	51	53	51	53	53	53	54	56	57	41	60	42	59	59	46	1000000	48	56	999589	56	0	0 K.	
NG	27	31	36	39	43	43	43	34	13	4	5	3	7	13	12	13	23	22	20	19	25	26	1000000	23	18	1000142	18	0	0 K.	
NL	12	12	13	9	10	7	9	9	7	9	9	9	8	10	10	12	10	12	12	15	17	14	19	1000000	11	7	1000401	9	-2	0 K.
NO	27	31	33	32	29	28	24	18	22	20	20	17	13	23	15	2	6	8	2	3	2	9	1000000	17	16	1000353	11	5	0 K.	
NZ	27	31	36	35	32	31	36	32	27	22	40	28	24	22	26	26	34	33	33	29	28	27	1000000	30	29	999984	29	0	0 K.	
PL	20	16	15	13	13	13	7	10	8	10	10	12	17	20	21	21	22	25	30	40	44	35	1000000	19	17	1000221	17	0	0 K.	
PT	16	21	21	19	14	15	15	12	12	12	15	16	18	18	18	22	14	17	25	24	17	16	1000000	17	15	1000272	16	-1	0 K.	
RO	27	31	36	39	43	42	43	45	46	46	46	43	45	44	45	44	32	29	26	13	20	17	1000000	36	44	999841	44	0	0 K.	
RS																														

Annex#6: Sorted results

countries	naive average	naive ranking	optimized index	optimized ranking	differences	validation	max	min
MN	50	62	999551	62	0	O.K.	1	-1
LK	49	60	999555	60	0	O.K.		
MT	49	60	999555	60	0	O.K.		
CY	49	58	999560	58	0	O.K.		
LY	49	59	999560	58	1	O.K.		
CU	48	57	999578	57	0	O.K.		
NE	48	56	999589	56	0	O.K.		
KH	48	55	999594	55	0	O.K.		
LA	46	54	999626	54	0	O.K.		
SI	45	53	999651	53	0	O.K.		
TR	43	52	999701	52	0	O.K.		
IQ	42	51	999709	51	0	O.K.		
RS	42	50	999712	50	0	O.K.		
UA	42	49	999714	49	0	O.K.		
ET	41	48	999738	48	0	O.K.		
SE	40	47	999754	47	0	O.K.		
EG	38	46	999799	46	0	O.K.		
LT	37	45	999821	45	0	O.K.		
RO	36	44	999841	44	0	O.K.		
CZ	36	42	999847	43	-1	O.K.		
RU	36	43	999848	42	1	O.K.		
IT	35	41	999876	41	0	O.K.		
ZA	34	40	999888	40	0	O.K.		
ID	34	39	999892	39	0	O.K.		
MY	34	38	999899	38	0	O.K.		
HU	34	37	999900	37	0	O.K.		
BD	34	35	999904	35	0	O.K.		
GR	34	35	999904	35	0	O.K.		
JP	33	33	999908	33	0	O.K.		
LV	33	33	999908	33	0	O.K.		
HR	32	32	999933	32	0	O.K.		
EE	31	31	999968	31	0	O.K.		
BG	30	30	999982	30	0	O.K.		
NZ	30	29	999984	29	0	O.K.		
LU	30	28	999988	28	0	O.K.		
VN	29	27	1000004	27	0	O.K.	5	-2
ES	28	25	1000026	25	0	O.K.		
IL	28	25	1000026	25	0	O.K.		
IN	27	24	1000056	24	0	O.K.		
MX	25	23	1000091	23	0	O.K.		
IE	25	22	1000092	22	0	O.K.		
CN	25	21	1000105	21	0	O.K.		
AR	24	20	1000112	20	0	O.K.		
NG	23	18	1000142	18	0	O.K.		
SK	23	18	1000142	18	0	O.K.		
PL	19	17	1000221	17	0	O.K.		
PT	17	15	1000272	16	-1	O.K.		
AU	16	14	1000289	15	-1	O.K.		
CH	15	13	1000306	14	-1	O.K.		
DE	14	12	1000328	13	-1	O.K.		
FI	14	11	1000330	12	-1	O.K.		
NO	17	16	1000353	11	5	O.K.		
GB	13	10	1000364	10	0	O.K.		
NL	11	7	1000401	9	-2	O.K.		
CA	9	6	1000450	8	-2	O.K.		
DK	12	9	1000475	7	2	O.K.		
BE	7	4	1000480	6	-2	O.K.		
AT	11	8	1000488	5	3	O.K.		
US	9	5	1000518	4	1	O.K.		
BR	5	3	1000537	3	0	O.K.		
ALL	3	2	1000579	2	0	O.K.		
FR	3	1	1000580	1	0	O.K.		

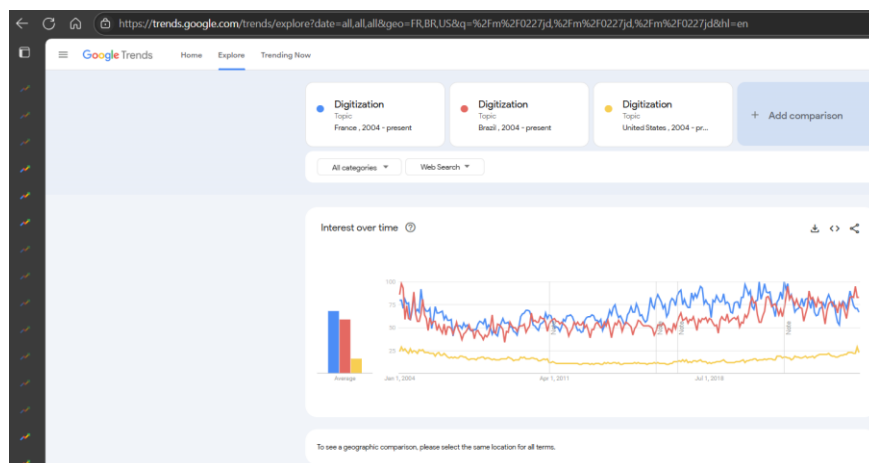
Annex#7: The over-norm-group

countries	naive average	naive ranking	optimized index	optimized ranking	differences	validation
NL	11	7	1000401	9	-2	O.K.
CA	9	6	1000450	8	-2	O.K.
DK	12	9	1000475	7	2	O.K.
BE	7	4	1000480	6	-2	O.K.
AT	11	8	1000488	5	3	O.K.
US	9	5	1000518	4	1	O.K.
BR	5	3	1000537	3	0	O.K.
ALL	3	2	1000579	2	0	O.K.
FR	3	1	1000580	1	0	O.K.

Annex#8: The under-norm-group

countries	naive average	naive ranking	optimized index	optimized ranking	differences	validation
MN	50	62	999551	62	0	O.K.
LK	49	60	999555	60	0	O.K.
MT	49	60	999555	60	0	O.K.
CY	49	58	999560	58	0	O.K.
LY	49	59	999560	58	1	O.K.
CU	48	57	999578	57	0	O.K.
NE	48	56	999589	56	0	O.K.
KH	48	55	999594	55	0	O.K.
LA	46	54	999626	54	0	O.K.
SI	45	53	999651	53	0	O.K.
TR	43	52	999701	52	0	O.K.
IQ	42	51	999709	51	0	O.K.
RS	42	50	999712	50	0	O.K.

Annex#9: Top 3 countries compared



Legends: The country-comparisons are also possible with Google Trends. In such cases it can be important involving population data to relativize Google Trends outputs...

Annex#10: Parallel analyses

[illegible]

Legends: B-version (left – Basic version) = only years, T-version (second on the left – Trend-version) = +trend-effects, D-version (second on the right – stDeviation-version) = +standard-deviation-effects, A-version (right – All-effect version)