

Warsaw University of Life Sciences – SGGW
Institute of Economics and Finance
Department of Econometrics and Statistics

**QUANTITATIVE METHODS
IN ECONOMICS**

**METODY ILOŚCIOWE W BADANIACH
EKONOMICZNYCH**

Volume XXVI, No. 2

Warsaw 2025

EDITORIAL BOARD

Editor-in-Chief: Bolesław Borkowski

Vice-Editor-in-Chief: Hanna Dudek

Managing Editors: Michał Gostkowski, Grzegorz Koszela

Theme Editors:

Econometrics: Bolesław Borkowski

Multidimensional Data Analysis: Wiesław Szczesny

Mathematical Economy: Zbigniew Binderman

Analysis of Labour Market: Joanna Landmesser-Rusek

Financial Engineering: Monika Krawiec

Data Science: Michał Gostkowski

Theory of Risk: Marek Andrzej Kociński

Statistical Editor: Wojciech Zieliński

Technical Editors: Jolanta Kotlarska, Elżbieta Saganowska

Language Editor: Agata Cienkusz

Native Speaker: Yochanan Shachmurove

Editorial Assistant: Łukasz Pietrych

SCIENTIFIC BOARD

Adnene Ajimi (University of Sousse, Tunisia)

Heni Boubaker (University of Sousse, Tunisia)

Peter Friedrich (University of Tartu, Estonia)

Paolo Gajo (University of Florence, Italy)

Agnieszka Gehringer (University of Göttingen, Germany)

Anna Maria Gil-Lafuente (University of Barcelona, Spain)

Jaime Gil-Lafuente (University of Barcelona, Spain)

Vasile Glavan (Moldova State University, Moldova)

Francesca Greselin (University of Milano-Bicocca, Italy)

Ana Kapaj (Agriculture University of Tirana, Albania)

Jirawan Kitchaicharoen (Chiang Mai University, Thailand)

Yuriy Kondratenko (Black Sea State University, Ukraine)

Vassilis Kostoglou (Alexander Technological Educational Institute of Thessaloniki, Greece)

Karol Kukula (University of Agriculture in Krakow, Poland)

Kesra Nermend (University of Szczecin, Poland)

Nikolas N. Olenov (Russian Academy of Sciences, Russia)

Alexander N. Prokopenya (Brest State Technical University, Belarus)

Yochanan Shachmurove (The City College of The City University of New York, USA)

Mirbulat B. Sikhov (al-Farabi Kazakh National University, Kazakhstan)

Marina Z. Solesvik (Nord University, Norway)

Ewa Syczewska (SGH Warsaw School of Economics, Poland)

Achille Vernizzi (University of Milan, Italy)

Andrzej Wiatrak (University of Warsaw, Poland)

Dorota Witkowska (University of Lodz, Poland)

ISSN 2082-792X

e-ISSN 2543-8565

© Copyright by Department of Econometrics and Statistics WULS – SGGW
(Katedra Ekonometrii i Statystyki SGGW w Warszawie)

Warsaw 2025, Volume XXVI, No. 2

Journal homepage: <https://qme.sggw.edu.pl>

Published by Warsaw University of Life Sciences Press

CONTENTS

Roman Kosmalski – Czynniki zróżnicowania wynagrodzeń oraz popytu na pracę w województwach.....	47
Krystian Jaworski, Nemanja Popović – Forecasting in Turbulent Times: How Artificial Intelligence and Machine Learning Are Reshaping Macroeconomic Prediction	60
Tomasz Woźniakowski – Machine Learning Based Predictions of Sales Leads: Proof of Quality from Polish Business-to-Business Company	70
Dorota Witkowska, Krzysztof Kompa – Capital Markets' Reaction to Covid-19 Pandemic and War in Ukraine	85

CZYNNIKI ZRÓŻNICOWANIA WYNAGRODZEŃ ORAZ POPYTU NA PRACĘ W WOJEWÓDZTWACH

Roman Kosmański  <https://orcid.org/0000-0002-9557-6357>

Instytut Gospodarki

Akademia Nauk Stosowanych im. Jana Amosa Komeńskiego w Lesznie

e-mail: roman.kosmalski@ansleszno.pl

Streszczenie: Praca ma na celu rozpoznanie empirycznych determinantów zróżnicowania popytu na pracę oraz wynagrodzeń w województwach w oparciu o teorię neoklasyczną. Cechą tej teorii jest przekonanie o doskonałej elastyczności cen i wynagrodzeń, zawsze sprowadzającej rynki do stanów równowagi. Uzyskane w pracy wyniki ujawniły, że neoklasyczne założenie o doskonałej elastyczności, charakterystyczne dla teoretycznych modeli neoklasycznych, nie jest gwarantowane. Województwa o wysokich wartościach współczynników elastyczności są jednocześnie województwami o wysokich stopach bezrobocia. Charakterystyczne dla teorii neoklasycznej automatyczne wyrównywanie się regionalnych dysproporcji, prawdopodobnie także nie wystąpi w przyszłości.

Słowa kluczowe: wzrost gospodarczy, wynagrodzenia, elastyczność, modele neoklasyczne

JEL classification: C02, C51, C53, O47

WPROWADZENIE

W Polsce występują znaczne dysproporcje rozwojowe pomiędzy województwami. Dzieje się tak, ponieważ w obszarach zurbanizowanych jest więcej szans na korzystanie ze zgromadzonego kapitału rzeczowego, wiedzy oraz sąsiedztwa innych inwestycji. Przekłada się to na wyższą efektywność realizowanych przedsięwzięć prowadząc do jeszcze większych dysproporcji regionalnych. Wspomnianą zależność można także odnieść do rynków pracy, wobec czego społeczeństwa stoją przed koniecznością dokonywania wyboru pomiędzy efektywnie działającym rynkiem pracy, lecz będącym źródłem większej rozpiętości wynagrodzeń, a ograniczeniem dysproporcji wynagrodzeń, jednak niekiedy kosztem

<https://doi.org/10.22630/MIBE.2025.26.2.5>

wyższego poziomu bezrobocia. Sprawne działanie neoklasycznych rynków uwarunkowane jest ich elastycznością, w kontekście rynków pracy przede wszystkim elastycznością wynagrodzeń. Skutkiem braku dostatecznej elastyczności mogą być zaburzenia w funkcjonowaniu rynku przekładające się na wzrost poziomu bezrobocia [Kwiatkowski 2002, s. 101], [Krugman 1994, s. 23-43].

W badaniach procesów gospodarczych realizowanych z wykorzystaniem metod ilościowych rozróżnia się szereg modeli matematycznych opisujących procesy wzrostu gospodarczego m.in. modele keynesowskie, neoklasyczne, modele wzrostu egzogenicznego, endogenicznego, jak również modele nowej geografii ekonomicznej. Przykładami modeli nawiązujących do teorii Keynesa są prace Harroda [1939] i Domara [1946]. Modele te zalicza się do grupy modeli krótkookresowych. Podwaliny pod nowoczesną, neoklasyczną teorię wzrostu gospodarczego stworzyły prace Solowa [1956] i Swana [1956], zmieniając sposób myślenia o ilościowych procesach wzrostu gospodarczego. Przykładami takich modeli mogą być prace Diamonda [1965], Uzawy [1965]. Modele te ukazują, że o długofalowej stopie wzrostu gospodarki decydują nie tylko nakłady kapitału, charakteryzujące jest malejąca krańcowa produktywność, lecz także siły determinujące ich produktywność. Stało się to przyczynkiem do rozwoju endogenicznych teorii, w ramach których wzrost produktywności czynników wytwórczych nie jest egzogeniczny, ale jest modelowany, czyli wyjaśniany przez model. W literaturze znane są liczne próby endogenizacji procesów wzrostu gospodarczego związane z pracami Arrowa [1962], Kaldora i in. [1962], Nelsona i Phelps'a [1966]. Jednak dynamiczny rozwój tej klasy modeli dały prace [Romer, 1990], [Grossman i in. 1991], [Aghion i in. 1992], [Mankiw i in. 1992].

Celem pracy jest rozpoznanie empirycznych determinantów zróżnicowania zatrudnienia oraz wynagrodzeń w województwach w oparciu o modele wzrostu gospodarczego. Zastosowanie w badaniu teorii neoklasycznej pozwoli empirycznie zweryfikować teoretyczne założenie o doskonałej elastyczności neoklasycznych mechanizmów sprowadzających zawsze gospodarkę do stanów równowagi. Dla realizacji celu badania wykorzystano neoklasyczną teorię wzrostu gospodarczego, którą rozbudowano o elementy endogeniczne w celu wyjaśnienia długofalowej dynamiki wynagrodzeń.

METODA BADANIA

Procesy produkcyjne w każdym z województw opisuje neoklasyczna funkcja produkcji typu Cobba-Douglasa z postępowem technologicznym w sensie Hicksa, dana wzorem¹:

¹ Dwuczynnikowa funkcja produkcji $F : R_+^2 \rightarrow R_+^1$ spełnia warunki neoklasycznej funkcji, jeżeli jest rosnąca, wklęsła, dodatnio jednorodna stopnia pierwszego i spełnia warunki Inady. Postęp technologiczny jest neutralny w sensie Hicksa, gdy funkcja produkcji spełnia warunek: $Y_i(t) = F_i(A_i(t), K_i(t), L_i(t)) = A_i K_i^{\alpha_i}(t) L_i^{1-\alpha_i}(t)$.

$$Y_i(t) = F_i(A_i(t), K_i(t), L_i(t)) = A_i(t)K_i^{\alpha_i}(t)L_i^{1-\alpha_i}(t), \quad \alpha_i \in (0, 1), \quad (1)$$

gdzie: $Y_i(t)$ – PKB wytwarzany w gospodarce i -tego województwa w momencie t ; $K_i(t)$ – nakłady kapitału rzeczowego zgromadzonego w i -tym województwie w momencie t ; $L_i(t)$ – zasoby pracy w i -tym województwie w momencie t ; $A_i(t)$ – zasoby dostępnej i wykorzystywanej w procesach produkcyjnych wiedzy w i -tym województwie; α_i – elastyczność produkcji ze względu na zgromadzone zasoby kapitału rzeczowego w i -tym województwie w momencie t ; $(1 - \alpha_i)$ – elastyczność PKB za względu na zasoby osób pracujących w i -tym województwie; t – oznacza dane liczbowe dla roku 2022 pochodzących z Banku Danych Lokalnych Głównego Urzędu Statystycznego.

Przedsiębiorstwa na neoklasycznym rynku kierują się zasadą maksymalizacji zysku oznaczającą opłacanie czynników produkcji według ich produktów krańcowych. Elastyczność produkcji ze względu na zasoby kapitału rzeczowego i pracy można wyznaczyć z warunków koniecznych maksymalizacji zysku przez typowe przedsiębiorstwo [Kosmański, 2023, s. 59]:

$$\Pi_i(K_i(t), L_i(t)) = \{A_i K_i^{\alpha_i}(t) L_i^{1-\alpha_i}(t) - r K_i(t) - w_i(t) L_i(t)\} \rightarrow \max, \quad (2)$$

gdzie: $w_i(t)$ – przeciętne wynagrodzenie w gospodarce i -tego województwa w momencie t ; r – stopa procentowa obowiązująca w Polsce.

Warunek maksymalizacji równania (2) ze względu na zmienną L_i określa:

$$\frac{d \Pi_i(A_i K_i^{\alpha_i}(t) L_i^{1-\alpha_i}(t))}{d L_i(t)} = (1 - \alpha_i) A_i K_i^{\alpha_i}(t) L_i^{-\alpha_i}(t) - w_i(t) = 0 \quad (3)$$

Współczynniki elastyczności PKB ze względu na liczbę osób pracujących równe są udziałowi wynagrodzenia pracy w PKB w i -tym województwie w momencie t :

$$(1 - \alpha_i) = \frac{w_i(t)}{A_i K_i^{\alpha_i}(t) L_i^{1-\alpha_i}(t)} = \frac{w_i(t) L_i(t)}{K_i^{\alpha_i}(t) A_i L_i^{1-\alpha_i}(t)} = \frac{w_i(t) L_i(t)}{Y_i(t)}, \quad (4)$$

po przekształceniach:

$$\alpha_i = 1 - \frac{w_i(t) L_i(t)}{Y_i(t)}. \quad (5)$$

Zasoby wiedzy zgromadzonej w i -tym województwie w momencie t wyznaczono ze wzoru (1) na podstawie zależności:

$$A_i(t) = \frac{Y_i(t)}{K_i^{\alpha_i}(t)}. \quad (6)$$

O zasobach wiedzy ukształtowanych w momencie początkowym $A_i(t) > 0$, zakłada się, że rosną według stopy postępu w sensie Hicksa [Hicks, 1963].

Popyt na pracę zgłaszany przez typowe przedsiębiorstwo działające na neoklasycznym rynku jest określany przez rozmiary zatrudnienia maksymalizujące zyski przedsiębiorstwa (Tokarski, 2011, s. 390), stąd na mocy równania (3) oraz z faktu, że:

$$\frac{d^2 \Pi_i(A_i(t)K_i(t)L_i(t))}{dL_i^2(t)} = -\alpha_i(1 - \alpha_i)A_i(t)K_i^{\alpha_i}(t)L_i^{-\alpha_i-1}(t) < 0 \quad (7)$$

wyznaczono funkcję popytu na pracę zgłaszanego przez typowe przedsiębiorstwo działające w neoklasycznej gospodarce. Ponieważ warunek (3) spełniony jest dla każdego $L_i > 0$, definiuje on rozmiary popytu na pracę w zależności od wysokości wynagrodzenia. Równanie (3) po przekształceniach zapisano w postaci:

$$L_i(t) = (1 - \alpha_i)A_i^{\frac{1}{\alpha_i}}(t)w_i^{-\frac{1}{\alpha_i}}(t)K_i(t). \quad (8)$$

Równanie (8) opisuje kształtowanie się popytu na pracę w zależności od wysokości wynagrodzeń, zasobów kapitału rzeczowego oraz zasobów wiedzy zgromadzonych w gospodarce i -tego województwa w danym momencie. Współczynniki elastyczności liczby osób pracujących ze względu na poziomy wynagrodzeń, zasoby kapitału rzeczowego oraz zasoby wiedzy są równe potęgde przy odpowiedniej zmiennej objaśniającej²: $\varepsilon_{i(w_i)}^{L_i}(t) = -\frac{1}{\alpha_i}$ elastyczność popytu na pracę ze względu na poziom wynagrodzeń, $\varepsilon_{i(A_i)}^{L_i}(t) = \frac{1}{\alpha_i}$ elastyczność popytu na pracę ze względu na zasób wiedzy, $\varepsilon_{i(K_i)}^{L_i}(t) = 1$ elastyczność popytu na pracę ze względu na zasób kapitału rzeczowego w gospodarce i -tego województwa.

Funkcja opisująca procesy produkcji w i -tym województwie w postaci intensywnej dana jest wzorem³:

$$y_i(t) = A_i(t)k_i^{\alpha_i}(t), \quad (9)$$

gdzie: $y_i(t)$ – PKB w przeliczeniu na osobę pracującą w i -tym województwie w momencie t ; $k_i(t)$ – kapitał rzeczowy w przeliczeniu na osobę pracującą w i -tym województwie w momencie t .

W wyniku przekształcenia równania (3) i (9) wyznaczono funkcję wynagrodzeń:

$$w_i(t) = (1 - \alpha_i)y_i(t). \quad (10)$$

Pracodawcy, kierując się zasadą maksymalizacji zysku, są skłonni zaoferować pracobiorcą wynagrodzenia w postaci części wydajności pracy mierzonej $y_i(t)$. Poziom wynagrodzeń będzie uzależniony od wydajności pracy oraz $(1 - \alpha_i)$ charakterystycznego dla każdego z województw.

² Elastyczność każdej funkcji można policzyć według wzoru: $E(Y/X) = f'(X) \cdot \frac{X}{f(X)}$. Dla

funkcji potęgowej: $E(Y/X) = (A \cdot X^b)' \cdot \frac{X}{A \cdot X^b} = \frac{Ab \cdot X^{b-1} \cdot X}{A \cdot X^b} = b$.

³ Równanie produkcji zapisane w postaci intensywnej otrzymano w wyniku przeliczenia zmiennych wyrażonych w równaniu (1) na osobę pracującą.

Tabela 1. Parametry modeli opisujących kształtowanie się popytu na pracę w województwach

Parametry	DŚ	KP	LB	LS	ŁD	MP	MZ	OP
α_i	0,589	0,505	0,407	0,518	0,526	0,508	0,639	0,511
$1/\alpha_i$	1,697	1,981	2,457	1,930	1,902	1,969	1,564	1,956
$1 - \alpha_i$	0,411	0,495	0,593	0,482	0,474	0,492	0,361	0,489
$A_i(t)$	115,5	278,3	773,7	221,1	225,4	302,9	68,8	237,4
Parametry	PK	PL	ŚL	PM	ŚK	WM	WP	ZP
α_i	0,502	0,465	0,550	0,576	0,433	0,477	0,598	0,503
$1/\alpha_i$	1,992	2,150	1,819	1,737	2,310	2,096	1,672	1,986
$1 - \alpha_i$	0,498	0,535	0,450	0,424	0,567	0,523	0,402	0,497
$A_i(t)$	260,7	409,0	186,0	136,5	597,7	346,6	99,0	258,8

Oznaczenia: DŚ – woj. dolnośląskie, KP – kujawsko-pomorskie, LB – lubelskie, LS – lubuskie, ŁD – łódzkie, MP – małopolskie, MZ – mazowieckie, OP – opolskie, PK – podkarpackie, PL – podlaskie, ŚL – śląskie, PM – pomorskie, ŚK – świętokrzyskie, WM – warmińsko-mazurskie, WP – wielkopolskie, ZP – zachodniopomorskie.

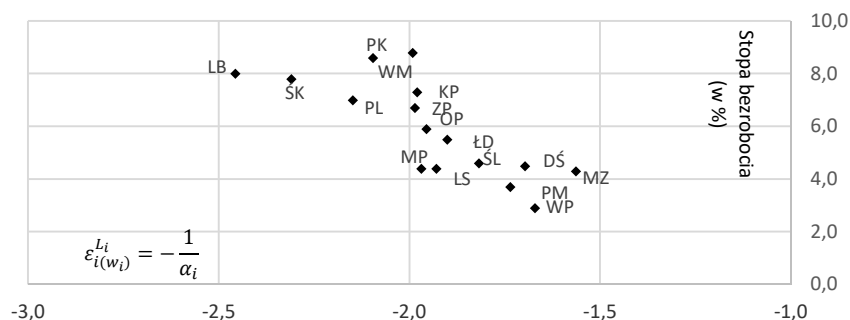
Źródło: opracowanie własne na podstawie wzorów (5) i (6)

Tabela 1 prezentuje estymowane parametry funkcji (8) do (10) opisujących kształtowanie się produkcji, wynagrodzeń oraz popytu na pracę w województwach. Parametry α_i oraz $A_i(t)$ zostały oszacowane na podstawie wzorów (5)-(6). Kształtowanie się popytu na pracę w zależności od poziomów wynagrodzeń, zasobów wiedzy, oraz zasobów kapitału rzeczowego zgromadzonych w województwach opisuje funkcja popytu na pracę (8), która jest funkcją potęgową, zatem jej parametry są jednocześnie elastycznościami liczby osób pracujących względem wskazanych zmiennych.

Wartości współczynników podanych w tabeli 1 uzupełniają naszą wiedzę o specyfice procesów gospodarczych i społecznych zachodzących w województwach, są także źródłem wiedzy o przyczynach nierówności regionalnych. Ponieważ miary elastyczności pokazują, o ile procent obniży się popyt na pracę w danym województwie, kiedy wynagrodzenia wzrosną o jeden procent, uwagę zwracają wysokie wartości współczynników $\varepsilon_{i(w_i)}^{L_i}(t)$ w województwach lubelskim, świętokrzyskim, podlaskim. W tych województwach jednoprocentowy wzrost wynagrodzeń będzie skutkował spadkiem liczby osób pracujących o więcej niż dwa procent. W województwie lubelskim jednoprocentowy wzrost wynagrodzeń będzie skutkował spadkiem popytu na pracę zgłaszanego przez przedsiębiorstwa o 2,457%. Podobnie należy interpretować wrażliwość (elastyczność) liczby osób pracujących ze względu na zasoby zgromadzonej wiedzy w województwach – z tą różnicą, że zależność ta ma charakter dodatni. Wzrost zasobów wiedzy w województwach przełoży się na silny wzrost popytu na pracę. Województwa te będą silnie wrażliwe na negatywne wstrząsy technologiczne.

Rysunek 1 pozwala zaobserwować zależności występujące pomiędzy stopami bezrobocia w województwach oraz wartościami współczynników elastyczności $\varepsilon_{i(w_i)}^{L_i}(t)$. Wartości współczynników przedstawiono w tabeli 1⁴. Wysokim ujemnym wartościom współczynników elastyczności liczby osób pracujących względem poziomów wynagrodzeń odpowiadają wysokie stopy bezrobocia. Z taką sytuacją mamy do czynienia m.in. w województwach lubelskim, świętokrzyskim, podlaskim, podkarpackim. W teoretycznej gospodarce neoklasycznej naturalną reakcją ze strony pracodawców jest spadek wynagrodzeń i wzrost popytu na pracę, zwłaszcza w gospodarkach województw o wysokich ujemnych wartościach współczynników elastyczności.

Rysunek 1. Stopa bezrobocia a elastyczność liczby osób pracujących względem wynagrodzeń



Źródło: opracowanie własne. Oznaczenia jak tab. 1

Prób tłumaczenia przyczyn występowania relatywnie wysokiego bezrobocia przy jednoczesnej sptywności wynagrodzeń w dół wielu badaczy doszukuje się w teorii płacy efektywnościowej, której fundamentalnym założeniem jest istnienie dodatniej zależności pomiędzy poziomami wynagrodzeń, a wydajnością pracy pracowników [Westley i in. 2006], [Bellante 1994, s. 21-33], [McConnell i in. 2006]. W sytuacji nadwyżki siły roboczej ponad popyt na pracę, pracodawcy w obawie przed pogorszeniem się wydajności pracy już zatrudnionych pracowników, nie obniżają wynagrodzeń. Skutkiem utrzymywania wynagrodzeń na poziomie wyższym od poziomu równowagi jest bezrobocie przymusowe.

Stany ustalone w neoklasycznych modelach

Neoklasyczne modele opisujące procesy gospodarcze zachodzące w długim horyzoncie czasu wykazują tendencję do zbieżności gospodarek przez nie opisywanych do stanów ustalonych, nazywanych także stanami stacjonarnymi. Są to stany gospodarek, w których inwestycje w kapitał rzeczowy zrównują się

⁴ Współczynnik elastyczności liczby osób pracujących ze względu na poziomy wynagrodzeń w województwach zgodnie z równaniem (8) przyjmują wartości ujemne.

z inwestycjami odtworzeniowymi niezbędnymi do utrzymania zasobów kapitału rzeczowego na stałym poziomie. Zmienne takie jak PKB, kapitał rzeczowy w przeliczeniu na osobę pracującą, ale także wynagrodzenia przyjmują określone wartości i dalej rozwijają się według stopy akumulacji wiedzy.

W długim horyzoncie czasu procesy akumulacji kapitału rzeczowego w przeliczeniu na osobę pracującą opisuje równanie [Malaga i in. 2007]:

$$\frac{dk_i(t)}{dt} = s_{k_i} A_i(t) k_i^{\alpha_i}(t) - (n_i + \rho) k_i(t) \quad (11)$$

gdzie: ρ – wskaźnik deprecjacji kapitału rzeczowego⁵; s_{k_i} – stopa inwestycji w i -tym województwie w momencie t , dana wzorem: $s_{k_i} = \frac{I_i(t)}{Y_i(t)}$; $I_i(t)$ – strumień inwestycji w kapitał rzeczowy w i -tym województwie w momencie t ; n_i – egzogeniczna stopa wzrostu liczby osób pracujących w gospodarce i -tego województwa, zależna od działania czynników demograficznych lub innych pozostających poza modelem; $(n_i + \rho)$ – suma wskaźników deprecjacji kapitału rzeczowego i stopy wzrostu liczby osób pracujących określana mianem inwestycji restytucyjnych lub odtworzeniowych; α_i – elastyczność produkcji ze względu na zgromadzone zasoby kapitału rzeczowego w i -tym województwie wyznaczono na podstawie wzoru (5), a jej oszacowania prezentuje tabela 1.

Tabela 2. Stopy inwestycji w kapitał rzeczowy oraz inwestycji odtworzeniowych

Parametry	DŚ	KP	LB	LS	ŁD	MP	MZ	OP
s_{k_i}	0,145	0,139	0,155	0,154	0,131	0,142	0,150	0,158
$(n_i + \rho)$	0,049	0,048	0,046	0,048	0,046	0,052	0,053	0,046
Parametry	PK	PL	SL	PM	ŚK	WM	WP	ZP
s_{k_i}	0,163	0,186	0,136	0,117	0,131	0,164	0,140	0,178
$(n_i + \rho)$	0,049	0,047	0,053	0,047	0,045	0,047	0,051	0,047

Źródło: opracowanie własne. Oznaczenia jak tab. 1

Tabela 3 przedstawia parametry równania (11). Stopy inwestycji w kapitał rzeczowy zdefiniowano jako udział inwestycji w i -tym województwie w PKB tego województwa. Najwyższe poziomy stóp inwestycji w roku 2022 wystąpiły w województwie podlaskim oraz zachodniopomorski. Najmniej w akumulację kapitału rzeczowego inwestowano w województwie pomorskim, dla którego parametr ten ukształtował się na poziomie 0,117, inwestycje stanowiły zaledwie 11,7% PKB tego województwa. Województwa, w których realna stopa deprecjacji kapitału rzeczowego $(n_i + \rho)$ była wyższa niż 5%, to województwa o dodatniej stopie wzrostu liczby osób pracujących i zaliczymy do nich mazowieckie, śląskie, małopolskie oraz wielkopolskie. W pozostałych województwach stopa liczby osób pracujących była ujemna.

⁵ Z uwagi na trudności z szacowaniem współczynników deprecjacji kapitału rzeczowego, w badaniu przyjęto arbitralnie ich wartość na poziomie 5%.

Rozwiązując równanie (11) ze względu na zmienną $k_i(t)$ otrzymano wartości kapitału rzeczowego w przeliczeniu na osobę pracującą w stanach ustalonych:

$$k_i^*(t) = \left(\frac{A_i s_{k_i}}{(n_i + \rho)} \right)^{\frac{1}{1-\alpha_i}}, \quad (12)$$

gdzie: k_i^* – wartość kapitału rzeczowego w stanie ustalonym w i -tym województwie.

W wyniku wstawienia równania (12) do równania (10) otrzymano formułę określającą poziomy wynagrodzeń w stanach ustalonych:

$$w_i^* = (1 - \alpha_i) A_i^{\frac{1}{1-\alpha_i}} s_{k_i}^{\frac{\alpha_i}{1-\alpha_i}} (n_i + \rho)^{-\frac{\alpha_i}{1-\alpha_i}}, \quad (13)$$

gdzie: w_i^* – przeciętne wynagrodzenie w gospodarce i -tego województwa znajdujące się w stanie ustalonym.

Tabela 3. Rozkład wynagrodzeń w zł dla roku 2022 oraz w stanach ustalonych

Parametr y	DŚ	KP	LB	LS	ŁD	MP	MZ	OP
w_i	6945,0	5888,6	5909,6	6014,4	6210,7	6824,7	7913,1	6134,2
w_i^*	16182,6	22365,6	23438,3	16732,6	27424,3	18229,7	31847,6	15092,8
w_i^*/w_i	2,33	3,80	3,97	2,78	4,42	2,67	4,02	2,46
Parametr y	PK	PL	ŚL	PM	ŚK	WM	WP	ZP
w_i	5662,5	6012,9	6696,8	6727,6	5782,6	5674,9	6019,7	6169,9
w_i^*	7654,6	19521,3	8523,2	15136,0	21758,9	18410,8	22661,3	33396,7
w_i^*/w_i	1,35	3,25	1,27	2,25	3,76	3,24	3,76	5,41

Źródło: opracowanie własne na podstawie równania (13). Oznaczenia jak tab. 1

W roku 2022 najwyższe poziomy wynagrodzeń zanotowano w województwie mazowieckim. Relatywnie wysokie były wynagrodzenia w dolnośląskim, pomorskim i śląskim, natomiast najniższe w podkarpackim i warmińsko-mazurskim. Relacja wynagrodzeń w roku 2022 do wysokości wynagrodzeń jakie ukształtują się w odległej przyszłości, pozwala sformułować przypuszczenie o braku konwergencji wynagrodzeń w przyszłości.

W literaturze podejmowane są liczne dyskusje o możliwych przyczynach różnicowania wynagrodzeń pomiędzy krajami lub regionami. Wyniki badań w tym zakresie są niejednoznaczne, a stosowane narzędzia i metody badania np. wpływu wymiany handlowej na procesy gospodarcze wywołują wiele dyskusji wśród badaczy [Cornia i in. 2001, s. 17-18]. Procesy takie często prowadzą do wzrostu poziomów wynagrodzeń wśród specjalistów w gospodarkach krajów wyżej rozwiniętych gospodarczo, przyczyniając się do wzrostu rozpiętości dochodowych [Kremer i in. 2006, s. 6-11]. Efektem może być też wzrost bezrobocia w niektórych sektorach gospodarki. Zróżnicowanie poziomów wynagrodzeń w województwach

pozostaje również przedmiotem badań w Polsce, czego przykładem jest praca [Mowczan i in. 2022].

Po zlogarytmowaniu równania (13) wyznaczono ścieżkę czasową zmian wynagrodzeń w gospodarce i -tego województwa znajdującego się w stanie ustalonym:

$$\ln w_i^*(t) = (1 - \alpha_i) \left(\frac{1}{1 - \alpha_i} \ln A_i(t) + \frac{\alpha_i}{1 - \alpha_i} \ln(s_{k_i}) - \frac{\alpha_i}{1 - \alpha_i} \ln(n_i + \rho) \right) \quad (14)$$

Ponieważ w stanie ustalonym stopa inwestycji w kapitał rzeczowy zrównuje się z inwestycjami odtworzeniowymi, wynagrodzenia rosną w tempie wyznaczonym przez zmienną $A_i(t)$:

$$\frac{\dot{w}^*(t)}{w^*(t)} = A_i(t), \quad (15)$$

gdzie: $\frac{\dot{w}^*(t)}{w^*(t)}$ – stopy wzrostu wynagrodzeń w gospodarkach województw znajdujących się w stanie ustalonym⁶.

W modelu neoklasycznym w długim horyzoncie czasu dynamika wynagrodzeń za pracę jest opisywana przez równanie (13) lub w postaci zlogarytmowanej (14). Ponieważ w modelach neoklasycznych standardowo zakłada się stałość parametrów w długim okresie, wyjaśnienie długofalowego tempa zmian wynagrodzeń w województwach, sprowadza się do określenia czynników determinujących dynamikę zmian zasobów wiedzy (15).

W prowadzonym badaniu dla określenia dynamiki zmian zasobów wiedzy w gospodarkach województw wykorzystano model Karla Shella [Shell 1967] oraz jego modyfikację przedstawioną w książce [Chiang 1999, s. 267]. Model czyni akumulację wiedzy zależną od ilości środków przeznaczanych na działalność badawczą oraz edukację. Zmienna $A_i(t)$, oznaczająca zasób wiedzy w gospodarce i -tego województwa, ma określony schemat zmian:

$$\dot{A}_i(t) = \sigma s_{A_i}(t) A_i(t) - \rho A_i(t), \quad (16)$$

gdzie: σ – współczynnik skuteczności wydatków w tworzeniu nowej wiedzy, arbitralnie ustalono na 0,9; $s_{A_i}(t)$ – część produkcji przeznaczana na tworzenie nowej wiedzy $s_{A_i}(t) = \frac{I_{A_i}(t)}{Y_i(t)}$ w i -tym województwie w momencie t ; $I_{A_i}(t)$ – suma wydatków na szkolnictwo wyższe, działalność badawczo-rozwojową oraz wydatki na oświatę z budżetów gmin i miast na prawach powiatu w i -tym województwie w momencie t ; ρ – jest stopą zaniku wiedzy⁷.

⁶ Zapis w ogólnej postaci: $\frac{\dot{x}(t)}{x(t)} \equiv \frac{dx}{dt} \equiv \frac{dx(t)}{dt}$ oznacza pierwszą pochodną zmiennej x ze względu na czas t , i interpretujemy jako przyrost wartości tej zmiennej w momencie t dla $t \in [0; +\infty)$.

⁷ Ze względu na trudności z oszacowaniem stopy zaniku wiedzy, w przykładzie przyjęto arbitralnie jej wartość na poziomie 2%.

Dzieląc równanie (16) przez zmienną $A_i(t)$ otrzymano stopy wzrostu zasobów wiedzy:

$$\frac{\dot{A}_i(t)}{A_i(t)} = \sigma s_{A_i}(t) - \rho. \quad (17)$$

Dopóki w gospodarce i -tego województwa $\sigma s_{A_i}(t) - \rho > 0$, gospodarka będzie rozwijała się, a wynagrodzenia będą rosły.

Tabela 4. Parametry równania dynamiki akumulacji zasobów wiedzy w województwach

Parametry	DŚ	KP	LB	LS	ŁD	MP	MZ	OP
$\dot{A}_i(t)/A_i(t)$	0,023	0,025	0,034	0,021	0,022	0,040	0,026	0,025
Parametry	PK	PL	ŚL	PM	ŚK	WM	WP	ZP
$\dot{A}_i(t)/A_i(t)$	0,034	0,031	0,031	0,022	0,023	0,030	0,021	0,022

Źródło: opracowanie własne na podstawie równania (17)

W tabeli 6. podano stopy wzrostu zasobów wiedzy w województwach wyznaczone według równania (17). Najwyższe wartości stóp wzrostu wiedzy, w konsekwencji także wynagrodzeń, występują w województwie małopolskim na poziomie 4% rocznie, następnie w województwach lubelskim i podkarpackim.

W wielu badaniach czynnikiem przyczyniającym się do powstawiania różnicowania poziomów wynagrodzeń w społeczeństwach pozostaje postęp technologiczny, sprzyja on bardziej dynamicznemu wzrostowi wynagrodzeń pracowników o wysokich kwalifikacjach, prowadząc do większej rozpiętości wynagrodzeń [Gottschalk i in. 1997], [Katz 1999], [Acemoglu 2002], [Green i in. 2003], [Acemoglu 2003]. Wyniki niektórych badań dowodzą, że wzrost zasobów wiedzy zgromadzonej w kraju lub regionie może ograniczyć rozpiętość wynagrodzeń [Cornia i in. 2001].

PODSUMOWANIE

Przeprowadzone w artykule badania ujawniły, że w niektórych województwach występują wysokie stopy bezrobocia przy jednocześniej wysokiej ujemnej elastyczności zatrudnienia ze względu na poziomy wynagrodzeń $\varepsilon_{i(w_i)}^{L_i}(t)$. Występujące nierównowagi na rynku pracy w teorii neoklasycznej mogą być wynikiem niedostatecznej elastyczności wynagrodzeń, braku mobilności siły roboczej lub działania związków zawodowych oraz instytucji [Kwiatkowski, 2002], [Parente i in. 2000], [Acemoglu i in. 2007] zaburzających sprawne funkcjonowanie rynków.

W modelach neoklasycznych, ze względu na malejącą krańcową produktywność czynników wytwórczych trwały wzrost gospodarczy i wzrost wynagrodzeń możliwy jest tak długo, jak długo gromadzona jest w gospodarce nowa wiedza tworząca innowacje zwiększające produktywność czynników wytwórczych. Dla zrozumienia mechanizmu oddziaływania innowacji na tempo wzrostu produkcji

i wynagrodzeń konieczne jest sięgnięcie do modeli endogenicznych, wtedy możliwe staje się oddziaływanie na procesy gospodarcze nie tylko za pomocą stopy inwestycji w kapitał rzeczowy ale także na długofalowe tempo wzrostu gospodarek [Temple, 2003]. W stanie ustalonym gospodarki rozwijają się według stopy wzrostu wiedzy wykorzystywanej w procesach produkcyjnych. W małych i średnich gospodarkach podstawowym źródłem wiedzy jest jej transfer z zagranicy. W wielu gospodarkach transfer wiedzy wyjaśnia większą część postępu technologicznego [Keller, 2004]. Ważną rolę także odgrywa własna działalność badawczo-rozwojowa, która może okazać się niezbędna do udanego transferu wiedzy z zagranicy, a przynajmniej w nim pomocna [Griffith i in. 2004]. Wyniki badań nad kształtowaniem się stopy wzrostu zasobów wiedzy w województwach, podane w tabeli 6, dowodzą o dużym znaczeniu inwestycji w szeroko rozumianą naukę dla długofalowego wzrostu gospodarek, którego tempo zdeterminowane jest wysokością wydatków na tworzenie nowej wiedzy. Otrzymane wartości wynagrodzeń, jakie ukształtują się w stanach ustalonych wyznaczonych za pomocą modeli neoklasycznych, uprawniają do przyjęcia założenia o małym prawdopodobieństwie wyrównywania się poziomów wynagrodzeń pomiędzy województwami w przyszłości. Czynnikiem zmniejszającym różnicowanie wynagrodzeń może być migracja pracowników do dużych, efektywnych ośrodków, gdzie większa część pracowników znajduje zatrudnienie w nowoczesnych sektorach, gdzie wstępują efekty skali i pojawia się szansa na zmniejszanie się dyspersji wynagrodzeń [Piketty 2006], [Kuznets 1955].

BIBLIOGRAFIA

- Acemoglu D. (2002) Directed Technical Change. *Review of Economic Studies*, (69)4, 781-809. <https://doi.org/10.1111/1467-937X.00226>
- Acemoglu D. (2003) Cross-country Inequality Trends. *Economic Journal*, 113(485), F121-F149. <https://doi.org/10.1111/1468-0297.00100>
- Acemoglu D., Antràs P., Helpman E. (2007) Contracts and Technology Adoption. *American Economic Review*, 97(3), 916-943.
- Adamczyk A., Tokarski T., Włodarczyk R. (2009) Przestrzenne różnicowanie płac w Polsce. *Gospodarka Narodowa*, 9, 87-108.
- Aghion P., Howitt P. (1992) A Model of Growth through Creative Destruction. *Econometrica*, (60)2, 323-351.
- Arrow K. (1962) The Economic Implications of Learning by Doing. *Review of Economic Studies*, 29(3), 155-173.
- Bellante D. (1994) Sticky Wages, Efficiency Wages, and Market Processes. *Review of Austrian Economics*, 8, 21-33. <https://doi.org/10.1007/BF01102315>
- Chiang A. (1999) *Elements of Dynamic Optimization*. Waveland Press.
- Cornia G. A., Court J. (2001) *Inequality, Growth and Poverty in the Era of Liberalization and Globalization*. UNU World Institute for Development Economics Research, United Nations University, Helsinki.

- Domar E. D. (1946) Capital Expansion, Rate of Growth, and Employment. *Econometrica*, 14(2), 137-147.
- Diamond P. (1965) National Debt in a Neoclassical Growth Model. *American Economic Review*, 55(5), 1126-1150.
- Gottschalk P., Smeeding T. M. (1997) Cross-National Comparison of Earnings and Income Inequality. *Journal of Economic Literature*, (35)2, 633-687.
- Green F., Felstead A., Gallie D. (2003) Computers and the Changing Skill-Intensity of Jobs. *Applied Economics*, 35(14), 1561-1576.
- Griffith R., Redding S., Reenen J. (2004) Mapping the Two Faces of R&D: Productivity Growth in a Panel of OECD Industries. *The Review of Economics and Statistics*, Massachusetts Institute of Technology Press, 86(4), 883-895.
- Grossman G. M., Helpman E. (1991) Quality Ladders in the Theory of Growth. *The Review of Economic Studies*, (58)1, 43-61.
- Harrod R. F. (1939) An Essay in Dynamic Theory. *Economic Journal*, 49(193), 14-33.
- Kaldor N., Mirrlees J. A. (1962).. A New Model of Economic Growth. *The Review of Economic Studies*, 29(3), 174-192.
- Katz L. D. (1999) Changes in the Wage Structure and Earnings Inequality. [w:] Ashenfelter O., Card D. *Handbook of Labor Economics*. Elsevier, (26)3A, Amsterdam.
- Keller W. (2004) International Technology Diffusion. *Journal of Economic Literature*, (42)3, 752-782.
- Kosmalski R. (2023) Równość czy efektywność. Wpływ polityki spójności na procesy wzrostu gospodarczego w województwach. *Gospodarka Narodowa*. *The Polish Journal of Economics*, 316(4), 54-71. <https://doi.org/10.33119/GN/174278>
- Kremer M., Masking E. (2006) Globalization and Inequality. Working Paper 2008-0087, Weatherhead Center for International Affairs, Harvard University, Cambridge MA.
- Krugman P. (1991) Increasing Returns and Economic Geography. *Journal of Political Economy*, 99(3), s. 483-499.
- Krugman P. (1994) Past and Prospective Causes of High Unemployment. *Economic Review*, 79(4), 23-43.
- Kwiatkowski E. (2002) Bezrobocie. Podstawy teoretyczne. Wydawnictwo Naukowe PWN.
- Kuznets S. (1955) Economic Growth and Income Inequality. *The American Economic Review*, 45(1), 1-28.
- Malaga K., Kliber P. (2007) Konwergencja i nierówności regionalne w Polsce w świetle neoklasycznych modeli wzrostu. Wydawnictwo AE, Poznań.
- Mankiw N. G., Romer D., Weil D. N. (1992) A Contribution to the Empirics of Economic Growth, *The Quarterly Journal of Economics*, 107(2), 407-437.
- McConnell, C.R., Brue, S.L., Macpherson, D.A. (2006) *Contemporary Labor Economics*, Seventh Edition. McGraw-Hill Irwin, New York.
- Mowczan D., Sztudynger J. (2022) The Relationship between Wage Differentiation and Productivity in Polish Regions. *Gospodarka Narodowa*, 312(4), 60-73. <https://doi.org/10.33119/GN/154835>
- Nelson R., Phelps E. (1966) Investment in Humans, Technological Diffusion, and Economic Growth. *The American Economic Review*, 56(1/2), 69-75.
- Parente S., Prescott E. (2000) *Barriers to Riches*. The Massachusetts Institute of Technology Press, Cambridge.

- Piketty T. (2006) The Kuznets' Curve, Yesterday and Tomorrow. [w:] Banerjee A., Benabou R., Mookherjee D. (red.). Understanding Poverty. Oxford University Press, 21. <https://doi.org/10.1093/0195305191.003.0004>
- Romer P. M. (1990) Endogenous Technological Change. *Journal of Political Economy*, 98(5), 71-102. <https://doi.org/10.1086/261725>
- Shell K. (1966) Toward A Theory of Inventive Activity and Capital Accumulation. *The American Economic Review*, 56(1/2), 62-68.
- Solow R. M. (1956) A Contribution to the Theory of Economic Growth. *Quarterly Journal of Economic*, 70(1), 65-94.
- Swan T. W. (1956) Economic Growth and Capital Accumulation. *Economic Record*, (32)2, 334-361.
- Temple J. (2003) The Long-Run Implications of Growth Theories. *Journal of Economic Surveys*, 17(3), 497-510. <https://doi.org/10.1111/1467-6419.00202>
- Tokarski T. (2011) *Ekonomia matematyczna. Matematyczne modele wzrostu gospodarczego*. Polskie Wydawnictwo Ekonomiczne S. A.
- Uzawa H. (1965) Optimal Technical Change in an Aggregative Model of Economic Growth. *International Economic Review*, 6, 18-31. <https://doi.org/10.2307/2525621>
- Westley Ch., Schmidt B. H. (2006) Efficiency Wages: A Critical Assessment. *Journal of Social, Political, and Economic Studies*, 31(2), 163-174.

FACTORS DIVERSITY OF WAGES AND LABOUR DEMAND IN VOIVODESHIPS

Abstract: The aim of the paper is to identify the empirical determinants of the differentiation of labor demand and wages in voivodeships based on neoclassical theory. Its characteristic feature is the belief in perfect flexibility of prices and wages, which always brings markets to equilibrium states. The results obtained, revealed that the neoclassical assumption of perfect elasticity, characteristic of theoretical neoclassical models, is not guaranteed. Voivodeships with high values of elasticity coefficients are also voivodeships with high unemployment rates. The automatic leveling of regional disparities, characteristic of neoclassical theory, will probably not occur in the future either.

Keywords: economic growth, wages, elasticity, neoclassical models

JEL classification: C02, C51, C53, O47

FORECASTING IN TURBULENT TIMES: HOW ARTIFICIAL INTELLIGENCE AND MACHINE LEARNING ARE RESHAPING MACROECONOMIC PREDICTION

Krystian Jaworski  <https://orcid.org/0000-0001-5186-9732>

Collegium of World Economy, Department of Economics II
SGH Warsaw School of Economics
e-mail: kjawor@sgh.waw.pl

Nemanja Popović  <https://orcid.org/0009-0003-5686-9362>

Faculty of Economics Podgorica
University of Montenegro
e-mail: nemanja.p@ucg.ac.me

Abstract: This paper examines how artificial intelligence and machine learning have reshaped macroeconomic forecasting in the volatile post-COVID era. Highlighting the use of ensemble methods, neural networks, and large language models, it illustrates their advantages in capturing nonlinear dynamics and processing complex data. Drawing on central bank case studies, the paper shows that AI enhances predictive power, though interpretability and robustness remain challenges. AI is best seen as a complement to, not a replacement for, traditional economic models and human judgment.

Keywords: macroeconomic forecasting; artificial intelligence; machine learning; large language models; inflation forecasting

JEL classification: C53, E37

INTRODUCTION

Macroeconomic forecasting has long been a cornerstone of economic policy planning. Before 2020, forecasters operated in a relatively stable environment by historical standards – often characterized by moderate fluctuations and well-understood business cycle dynamics. In the decades leading up to the COVID-19 pandemic, especially during the so-called “Great Moderation”, the mid-1980s to mid-2000s, the volatility of key macroeconomic variables like output and inflation

<https://doi.org/10.22630/MIBE.2025.26.2.6>

was subdued. This stability enabled forecasters to rely on historical patterns and established models with some confidence. Standard forecasting tools included structural models grounded in economic theory (e.g. Dynamic Stochastic General Equilibrium models) and time series models (such as vector autoregressions and statistical trend projections), typically assuming that the future would resemble the recent past with only modest shocks. Indeed, large forecast errors were relatively infrequent in the absence of major disturbances, although events like the 2008–09 Global Financial Crisis had already highlighted that traditional models could falter in the face of unprecedented turmoil [Tchoketch-Kebir, Madouri 2024]. By the late 2010s, however, a prolonged period of low inflation and steady growth had perhaps lulled many forecasters into a false sense of security regarding the stability of economic relationships (a point underscored by the widespread failure to foresee the financial crisis or the sluggish recovery that followed). The period since 2020 has confronted forecasters with an unusually turbulent sequence of global shocks, each compounding the uncertainty of the previous one. First and foremost was the COVID-19 pandemic itself, which triggered the deepest global recession in living memory almost overnight (the “Great Lockdown”). The COVID-19 pandemic in 2020 marked a dramatic turning point for macroeconomic forecasting, exposing the fragility of standard forecasting approaches when faced with a truly novel disturbance. Given this context, the purpose of this paper is to examine how the COVID-19 shock and subsequent global upheavals have increased the complexity and uncertainty of macroeconomic forecasting, and how forecasting models and practices have been forced to adapt in response. The pandemic set off a chain reaction of successive global shocks – including supply chain breakdowns, energy price spikes, geopolitical conflicts, and an inflation surge – which together created a volatile and highly uncertain post-2020 environment for forecasters.

THE RISE OF AI IN FORECASTING

One of the most frequently noted innovations in macroeconomic forecasting in response to recent challenges has been the rapid integration of Artificial Intelligence and Machine Learning (AI/ML) techniques. Over the past few years, AI/ML tools have moved from the periphery of economic research toward the mainstream of forecasting practice, as traditional methods proved insufficient in the face of unusual shocks [Jouilil, Iaousse 2023]. The appeal of these methods lies in their ability to automatically detect complex patterns, utilize vast amounts of data (including novel data sources), and model nonlinear relationships that classical models might miss. In an environment where past equations broke down, machine learning offered a data-driven way to “let the data speak,” potentially uncovering new predictive signals.

One family of AI techniques making inroads is machine learning ensemble methods, such as random forests and gradient boosting machines (GBMs). These are powerful nonlinear regression algorithms that combine many simple prediction rules

to capture complex interactions in the data. A random forest, for example, builds numerous decision trees on random subsets of data and averages them, which tends to improve accuracy and robustness compared to a single tree. Boosting methods like XGBoost and LightGBM iteratively build an ensemble of trees, correcting errors at each step, and often achieve very high predictive accuracy. Central bank researchers have found these methods particularly useful for forecasting in volatile conditions. For instance, the European Central Bank (ECB) has begun using tree-based ensemble models as supplementary forecasting tools alongside its core structural models [Lenza et al. 2023a]. These models can ingest a large number of input variables – commodity prices, financial indicators, dozens of sectoral confidence indices, etc. – and algorithmically select the most predictive combinations. Machine learning can find patterns without an a priori model specification. Moreover, ensemble methods naturally accommodate nonlinear relationships. This was especially important during the pandemic and subsequent shocks, as relationships like “oil prices → inflation” or “financial stress → credit availability” likely have threshold effects and interactions that linear models don’t capture. Indeed, a commentary from the Czech National Bank [CNB 2025] notes that the main advantage of methods like random forests and gradient-boosted trees is their ability to capture nonlinear relationships and complex inflation dynamics, which proved “especially important during sudden shocks” such as COVID-19 and the inflation spike that followed. In practical terms, tree-based models were better able to fit the unusual post-pandemic inflation process (where, for example, inflation might respond nonlinearly once supply chain delays exceeded a certain length). Empirical tests have borne this out: researchers Lenza et al. [2023a] at the ECB find that an advanced ensemble method, *quantile random forests*, can not only forecast the path of inflation but also quantify the uncertainty by predicting the entire distribution (quantiles) of future inflation. This allows policymakers to see risks to the forecast (e.g. a certain probability of very high inflation) and plan accordingly, a critical capability in uncertain times.

Another set of AI tools making waves are neural networks, particularly advanced architectures like Recurrent Neural Networks (RNNs) and Long Short-Term Memory (LSTM) networks, as well as the more recent Transformer models. Neural networks are nonlinear function approximators inspired by the brain’s neurons; when given sufficient data, they can, in principle, approximate very complex mappings from inputs to outputs. Deep learning networks with many layers became famous for image and speech recognition in the 2010s, but they have also been applied to time series forecasting. LSTM networks, a type of RNN, are explicitly designed to handle sequence data and remember long-term dependencies. This makes them well-suited to economic time series that may have multi-quarter trends or delayed effects. For example, an LSTM can, in theory, learn that “a shock now might affect output more strongly three quarters ahead” if such patterns exist in the data, without the modeler needing to pre-specify lag lengths. During the pandemic, some institutions experimented with LSTMs to forecast indicators like

GDP and inflation [Li 2024]. A study at Banco de la República (Colombia's central bank) used LSTM models to forecast Colombian inflation one year ahead [Cárdenas-Cárdenas et al. 2023]. They tried two variants: one with only past inflation as input and another incorporating additional variables (like exchange rates, oil prices, etc.). They found that the LSTM with a richer input set significantly outperformed traditional ARIMA models, especially at longer horizons, with the gains becoming most pronounced beyond 6 months ahead. This suggests that the LSTM was able to capture nonlinear interactions between inflation and other indicators (and perhaps global factors) that a linear model could not, thus yielding more accurate medium-term forecasts.

Similarly, other researchers have applied deep learning to GDP nowcasting, sometimes feeding in an array of high-frequency indicators as a multivariate input to an LSTM or a Transformer model. The Transformer architecture, which underlies modern Large Language Models (LLMs), has also been adapted for time series forecasting in experimental studies. Transformers excel at capturing long-range dependencies in sequences using attention mechanisms. Early results indicate that large models with well-designed learning can sometimes provide more accurate predictions than traditional parametric models in diverse scenarios. In other words, an LLM that has learned to predict patterns (even patterns in text) can be surprisingly powerful when that prediction ability is transferred to economic data – though this is still a cutting-edge area of research.

A notable example of LLM use in forecasting comes from experiments in inflation forecasting using GPT-style models. Analysts at the St. Louis Fed [Fariae-Castro, Leibovici 2024] tested an LLM (Google's PaLM 2, a large language model) on the task of projecting U.S. inflation and compared its performance to professional forecasters. They found that the LLM, after appropriate training/prompting, could estimate inflation trends more accurately than human forecasters in many cases over the 2019–2023 period. The PaLM-based model was *more accurate in most years and across nearly all forecast horizons* than the median professional forecast. This is a striking result, suggesting that LLMs (which incorporate vast textual knowledge, potentially including real-time news and narratives) picked up on the brewing inflation pressures earlier or interpreted the pandemic/economic data in a way that gave them an edge. However, the study also highlights a key issue: the LLM is essentially a black box, and it is “not entirely clear how [the] AI arrives at its predictions”. Complementary evidence comes from Lopez-Lira and Tang (2024), who show that LLMs can forecast indirectly through sentiment analysis of news headlines, with LLM-based sentiment scores outperforming traditional methods and accuracy improving with model size. We will revisit this black-box problem in the risk section, but it is noteworthy that AI models have begun to rival expert judgment in forecasting tasks. The Czech National Bank [2025] undertook a similar exercise with two proprietary AI models (nicknamed “OpenAI o1” and “Grok 2”) to forecast Czech inflation, comparing them to each other and to human analysts. During the stable pre-pandemic period, both the AI and

human forecasts were similar (all expecting inflation to stay near target). But when inflation surged in 2021–2022, neither humans nor AI predicted the magnitude of the spike accurately – confirming that this shock fooled everyone. Interestingly, one AI model (OpenAI o1) did start signaling persistent inflationary pressure earlier than the consensus; its forecasts in 2021 indicated inflation would keep rising, whereas many human analysts thought the rise was temporary. In retrospect, the AI’s direction was right (inflation was not as transitory as hoped), though it still underestimated the degree of the increase, as did essentially all forecasters. The second AI model (Grok 2) behaved more like the humans, treating the initial surge as largely temporary. By late 2023, Grok 2 expected inflation to fall quickly, correctly matching the actual decline, whereas OpenAI o1 expected more persistence. The differences between AI models highlight that not all AI is the same – their architectures and training can yield different biases. The CNB [2025] analysis concluded that these AI models can offer useful “alternative perspectives” – for example, an early warning that inflation might stick around – so they could serve as indicators or supplementary forecasts alongside traditional ones. However, their black-box nature means they aren’t replacing human forecasters; instead, they act as an additional input, much like a survey of models.

Aside from neural networks and LLMs, regularization and dimensionality-reduction techniques from machine learning have also proven valuable. A common challenge in macro forecasting is having too many potential predictors relative to the number of observations (since quarterly data records are limited). ML offers solutions like LASSO (Least Absolute Shrinkage and Selection Operator) and ridge regression, which automatically select or shrink coefficients on less useful variables, effectively performing variable selection. During the pandemic, when forecasters threw dozens of new indicators into the mix, these techniques helped avoid overfitting. In fact, an IMF [2024] paper titled “*Mending the Crystal Ball: Enhanced Inflation Forecasts with Machine Learning*” found that a simple LASSO regression outperformed more complex models and benchmarks for short-term inflation forecasting in the volatile 2022–23 period, in a study on Japan. The flexibility to incorporate many predictors (oil prices, output gap, global prices, etc.) and automatically shrink irrelevant ones gave LASSO an edge over both a small-scale AR model and even nonlinear ML like random forests in that case. The authors note that ML models’ flexibility and focus on pure forecasting (rather than structural explanation) were key advantages in capturing the evolving inflation dynamics. Essentially, ML could adapt more quickly to the fact that, for example, import prices and supply shocks had become the dominant drivers of inflation during 2021–22, whereas traditional models anchored on past low-inflation regimes struggled.

Another striking example comes from the Bank of England: researchers there used a high-dimensional approach with disaggregated price data. Joseph et al. [2024] built models to forecast UK CPI inflation using hundreds of sub-component price series (essentially breaking inflation into all its item categories) along with various ML algorithms. They found that when they exploit this “big data” of item-level

prices, they achieve strong improvements in forecast accuracy – up to 70% reduction in error at a 1-year horizon – compared to a standard aggregate model. In particular, shrinkage methods like ridge and LASSO performed best in handling the item-level data, indicating that the combination of a large, granular information set and appropriate regularization is key to good performance. This result underscores how ML enables forecasters to use the richness of micro data (which items are driving inflation? Are some prices spiking while others are stable?) to better predict the macro outcome. Traditional models could not handle hundreds of inputs or would suffer from overfitting, whereas ML made it feasible and even provided ways (like Shapley values and variable importance measures) to interpret which categories mattered most. In effect, this approach blends economic domain knowledge (knowing that disaggregates contain signal) with ML’s data-crunching power.

INSTITUTIONAL ADOPTION & REAL-WORLD APPLICATIONS

There have been numerous other case studies of AI in forecasting across various institutions. To highlight a few: the IMF has applied deep learning to nowcast GDP by feeding satellite imagery and other unconventional data into neural networks (expanding coverage of countries where data is sparse). The BIS (Bank for International Settlements) has explored ML for financial cycle prediction and early warning of crises (finding, for example, that tree-based models can improve predictions of banking stress by capturing nonlinear interactions among credit, asset prices, and global factors [Aldasoro 2025]). The Bank of Canada published guidance on when and how to use ML in economics, showing cases like housing price prediction where ML (e.g. gradient boosting) outperforms because of nonlinear effects of income, interest rates, and regional factors [Desai 2023]. At the Federal Reserve, researchers have incorporated ML into their forecasting toolkits on an experimental basis – for instance, using random forests to improve recession probability models and using textual analysis (with NLP algorithms) on the Fed’s Beige Book and news to augment economic forecasts. These examples illustrate a broad trend: AI/ML methods are no longer theoretical exercises but are being deployed in real forecasting contexts, especially to tackle problems of large data and structural change [Bareith et al. 2024].

How do these AI-driven forecasts compare to traditional approaches in the uncertain environment of recent years? The evidence so far is that AI/ML can substantially improve predictive performance in certain dimensions, but they are not uniformly superior in all cases and still work best in conjunction with human insight. In “normal” times, simpler models often hold their own – indeed, a well-known result from the forecasting literature is that with stable relationships, a simple linear model or even a random walk can be hard to beat. But in the chaotic period of 2020–2022, there have been clear instances where ML models captured turning points or complex drivers better. For example, as noted, a gradient-boosted tree was able to account for time-varying, nonlinear relationships during the U.S. pandemic

recovery and beat linear models in forecasting unemployment. Similarly, the use of large information sets combined with ML (factor models with penalization) yielded improved accuracy for inflation when old models failed to anticipate the surge. However, ML is no panacea: some central bank evaluations found that when structural breaks are extreme, no model – not even ML – had enough training data to predict the shift. For instance, during the initial COVID shock, many ML models trained on pre-2020 data also failed until they could be retrained with pandemic data included. Moreover, certain ML models like standard neural networks require enormous data to train effectively, and macroeconomic sample sizes are limited. This is why methods like tree ensembles and regularized linear models (which are more data-efficient) often performed better than very complex deep learning in macro forecasting competitions (the IMF study where simple LASSO beat newer nonlinear models in 2022 inflation forecasting is telling).

The consensus emerging is that AI/ML methods are powerful complements to, rather than outright replacements for, traditional models and economist judgment. They excel at pattern recognition and handling complexity – providing new “eyes” on the data. They might pick up an incipient trend or correlation that a human or simple model misses (for example, an ML model might detect that an uptick in online job postings plus rising used car prices is an early predictor of broader inflation). But human forecasters still play a crucial role in guiding these models, choosing sensible input features, and interpreting outputs (especially when the outputs seem counterintuitive). In practice, many institutions now run hybrid forecasting processes: a structural model might produce a baseline, and an ML model produces an alternative forecast, and experts will compare the two, understand why they differ, and often come up with a reasoned synthesis. In uncertain times, having these diverse model outputs enriches the information set for decision-makers. The rise of AI thus should be seen as an expansion of the forecasting arsenal. As Mullainathan and Spiess [2017] pointed out, machine learning excels at pure prediction tasks and can discover complex predictive signals, whereas traditional econometric models excel at interpretation and theory-consistent structure – blending the two can yield the best of both worlds. Recent experience validates this: the best forecasting performance has often come from combining a large and relevant information set (where ML helps sift signals) with sound judgment and parsimony (where human economists ensure interpretability and avoid nonsense).

RISKS AND LIMITATIONS OF AI

Although AI/ML models are a powerful tool, especially in terms of data processing, one should not forget the limitations that accompany their use, such as interpretability and transparency, overfitting and robustness, data quality and availability, big data dependency, and transparency, accountability, and ethical concerns [IMF 2024]. A central challenge lies in the interpretability of complex AI models, which often function as “black boxes”. Unlike traditional econometric

approaches, where the influence of each variable can be clearly identified and explained, many AI models provide little insight into the mechanisms behind their predictions, what is problematic in policy contexts where decision-makers must justify forecasts to stakeholders and base decisions on traceable reasoning. For example, if an AI model is showing a rise in inflation, then policymakers should know whether this reflects wages, commodity prices, or other factors, and how these drive overall prices.

Efforts to improve interpretability, such as the use of Shapley values or feature attribution methods, can be employed to decompose model outputs into input contributions. However, these approaches tend to detect correlation more than causation, limiting their explanatory power, and thus most institutions have responded by embracing hybrid approaches, combining AI insights with both structural models and expert judgment, in order to maintain both performance and intelligibility [Desai 2023]. Robustness of the model is also a critical issue, as highlighted above, given that ML models are highly susceptible to overfitting, especially in data-poor environments like macroeconomics. Model predictions that perform well on historical data can be disastrous when faced with structural breaks or regime shifts, as was seen in the COVID-19 pandemic. Data quality and availability is another concern, especially if we know that the quality and comparability of data is often questionable, and AI relies on the use of large databases. Many high-frequency or alternative datasets are noisy, short in duration, or subject to revision.

Finally, concerns about accountability, ethical data use, and model management remain. That is why institutions are increasingly emphasizing explainability, rigorous validation, and in particular human oversight to ensure the responsible application of AI in forecasting.

CONCLUSIONS

In conclusion, the upheavals of the past few years have ushered in a new paradigm for macroeconomic forecasting – one that is more data-intensive, probabilistic, and cognizant of uncertainty. Forecasting in the post-COVID world is undeniably more challenging; unprecedented shocks taught us to expect the unexpected and to be humble about our models' limitations. But at the same time, those challenges acted as a catalyst for innovation. Forecasters have broadened their toolkit to include high-frequency indicators, non-traditional data, and powerful AI/ML algorithms, enabling them to adapt in real time and capture complex dynamics better than before. Traditional models like DSGEs and VARs have not been thrown out – they still provide valuable theoretical consistency and a baseline for thinking – but they are now augmented by adaptive methods that can handle regime changes and nonlinearities [Liu 2024]. The integration of AI and machine learning has been a central development: as we have detailed, these methods have improved forecasting of nowcasted GDP, inflation, and other variables in volatile

conditions, often outperforming legacy approaches and providing new insights (for example, through analysis of text or disaggregate data). Major institutions from the ECB to the Fed and IMF have not only acknowledged this shift but actively incorporated it, blending human expertise with machine computations in their forecasting process.

The net effect is that macroeconomic forecasts today are arguably more robust and information-rich than those prior to 2020. A central bank forecasting round now might include input from an AI model flagging an uptick in online job postings that presages hiring, a textual analysis summarizing thousands of business reports into a sentiment index, and scenario drills of various tail risks – all alongside the conventional outlook. Forecasters have transformed from simply extrapolating trends to managing an ensemble of models and data streams, navigating through uncertainty with a combination of technological assistance and seasoned judgment. This evolution bodes well for dealing with future shocks: whether it is a climate-related event, a technological disruption, or another pandemic, the forecasting community is now better equipped to respond quickly and flexibly.

Of course, challenges remain. As we've argued, care must be taken to ensure models remain transparent, generalizable, and anchored in economic reality [Almosova, Andresen 2022]. Ongoing research will be needed to improve explainability and to prevent the misapplication of AI (for example, avoiding false confidence in predictions). Yet, the trajectory is clear – macroeconomic forecasting is becoming more of a high-tech endeavor, continuously updated with real-time data and refined with intelligent algorithms, all under the guidance of human expertise. AI/ML outputs should complement, not replace, structural models and expert judgment. Such a hybrid approach combines flexibility with theoretical grounding. Far from making human forecasters obsolete, the AI revolution in forecasting is reshaping their role – from solitary practitioners of “the dismal science” into collaborative teams of economists and machines working together to decipher an ever more complex economic landscape.

REFERENCES

- Aldasoro I., Hördahl P., Schrimpf A., Zhu X. S. (2025) Predicting Financial Market Stress with Machine Learning. BIS Working Papers 1250, BIS.
- Almosova K. Andresen M. (2022) Non-linear Inflation Forecasting with Recurrent Neural Networks. *Journal of Forecasting*, 41(8), 1250-1269. <https://doi.org/10.1002/for.2901>
- Bareith H., Adjei E., Hassan M. Kokot W. (2024) Navigating Inflation Challenges: AI-Based Portfolio-Management Insights. *Risks*, 12(3), Art. 46. <https://doi.org/10.3390/risks12030046>
- Cárdenas-Cárdenas J. A., Cristiano-Botía D. J., Martínez-Cortés N. (2023) Colombian Inflation Forecast using Long Short-Term Memory Approach. *Borradores de Economía*, 1241.

- Czech National Bank (2025) First Use of AI in Inflation Forecasting at the CNB. CNB Research Note.
- Desai A. (2023) Machine Learning for Economics Research: When What and How? Staff Analytical Note 2023-16. Bank of Canada.
- Faria-e-Castro M., & Leibovici F. (2024) Artificial Intelligence and Inflation Forecasts. Federal Reserve Bank of St. Louis Review, 106(12), 1-14.
- IMF (2024) Mending the Crystal Ball: Enhanced Inflation Forecasts with Machine Learning. IMF Working Paper.
- Joseph A., Potjagailo G., Chakraborty C., Kapetanios G. (2024) Forecasting UK Inflation bottom up. International Journal of Forecasting, 40(4), 1521-1538.
- Jouilil M., Iaousse A. (2023) Comparing the Accuracy of Classical and Machine-Learning Methods in Time-Series Forecasting: a Case Study of USA Inflation. Statistics, Optimisation & Information Computing, 11(3), 1767-1784. <https://doi.org/10.19139/soic-2310-5070-1767>
- Lenza M., Moutachaker I., Paredes J. (2023a) Density Forecasts of Inflation: a Quantile Regression Forest Approach. European Central Bank, Working Paper Series, 2830.
- Lenza M., Moutachaker I., Paredes J. (2023b) Forecasting Euro Area Inflation with Machine-Learning Models. European Central Bank, Research Bulletin, 112.
- Li J. (2024) Inflation Forecasting Using a Hybrid LSTM-SARIMA Model Based on Discrete Wavelet Transform. Advances in Economics, Management and Political Sciences, 73, 1419-1434. <https://doi.org/10.54254/2754-1169/73/20231419>
- Liu Y. (2024) Forecasting China's Consumer Price Index (CPI) based on Combined ARIMA-LSTM Models. Frontiers in Business Economics and Management, 3(2), 49-67. <https://doi.org/10.54097/v49dwv67>
- Lopez-Lira A., Tang Y. (2024) Can ChatGPT Forecast Stock Price Movements? Return Predictability and Large Language Models, arXiv preprint, arXiv:2304.07619.
- Mullainathan S., Spiess J. (2017) Machine Learning: An Applied Econometric Approach. Journal of Economic Perspectives, 31(2), 87-106.
- Tchoketch-Kebir A. Madouri A. (2024) Research Leadership and High Standards in Economic Forecasting: Neural-Network Models Compared with Etalon ARIMA Models. Business Ethics and Leadership, 8(1), 220-233. [https://doi.org/10.61093/bel.8\(1\).220-233.2024](https://doi.org/10.61093/bel.8(1).220-233.2024)

MACHINE LEARNING BASED PREDICTIONS OF SALES LEADS: PROOF OF QUALITY FROM POLISH BUSINESS-TO-BUSINESS COMPANY

Tomasz Woźniakowski  <https://orcid.org/0000-0002-0779-4769>

Institute of Economics and Finance
Warsaw University of Life Sciences – SGGW
e-mail: Tomasz_Wozniakowski@sggw.edu.pl

Abstract: Most sales managers struggle with achieving high lead conversion, key to lowering marketing costs and improving sales efficiency. Existing research emphasizes costly large-scale methods, often inaccessible to SMEs. Meanwhile, IT SMEs in B2B face numerous low-value leads without predictive support. This study proves that AI (AutoML on Google Cloud) can cost-effectively predict sales opportunities. Using 1000 historical leads, it demonstrates accurate predictions, offering SMEs a practical tool and paving the way for further research.

Keywords: sales leads, sales leads conversion, sales lead value, artificial intelligence, automated machine learning

JEL Classification: A11; B16; B21; C45; D22; M31

INTRODUCTION

The digital transformation in business is driven, among other factors, by a huge increase in the amount of data and computing power [Małkowska et al. 2021]. The problem of sales procedures optimization becomes very intense, competitive and complex process. Sales win-propensity prediction is fundamental to effective sales management [Yan et al. 2015].

Large actors often use very sophisticated and expensive methods including big data analysis often utilizing artificial intelligence. [Yan et al. 2015] This seems to be the reason why most of researchers are focused on large scale operations through the use of large datasets [Plawgo et al. 2021], therefore describe time and

<https://doi.org/10.22630/MIBE.2025.26.2.7>

cost consuming methods of sales optimization which results in high entry threshold, very often unreachable for small and medium enterprises (SMEs).

The SME sector plays the leading role in labor creation, sales, and production of value-added goods and services in most countries [Ključnikov et al. 2022]. SMEs significantly contribute to the growth GDP dynamics. Small and medium-sized enterprises create a more intense competitive environment which translates into prices and quality of products and services [Oleksiuk et al. 2022]. Moreover, the high-tech sector was affected by globalization at a much faster pace than other sectors [Soniewicki et al. 2022]. Therefore, an author will argue that SMEs must reduce this technology gap [Hyder et al. 2022].

The article emphasizes the real-life experiment of the Business To Business (B2B) market model in which small enterprises often compete with large ones, which appears to be very strong in IT sector [Mang'unyi et al. 2019]. The SMEs can offer a very narrow but at the same time very deep range of competencies, effectively gaining serious positions and even an advantage in selected niches. Therefore they have to provide very intense marketing activities which results in high numbers of sales leads (SL), in most cases without any kind of sophisticated value prediction support. To maintain a competitive advantage on the market, companies must efficiently adapt their operations [Nogalski et al. 2020].

The aim of the paper is to prove: there is a cost-effective method of using Artificial Intelligence (Machine Learning) for successful predictions of sales in SME operating in B2B model.

The study is focused on complex data taken from Polish business-to-business company in IT sector. Although the standard European Union classification divides SMEs into three categories regarding their size (micro, small and medium-sized enterprises) [Civelek et al. 2021], the researcher classified the company as small since it has less around 30 employees.

Researched sample spans last 5 years of its more than 30 years existence on the market. Data analysis are empowered by couple real-time experiments conducted in last 6 months of year 2022 when results of sales predictions done by models, AI and more traditional approach have been compared.

The paper examines the complexity of the sales activities: 1) financial value of sales leads, 2) success probabilities defined by status, 3) intensity of contacts, 4) geographical factors, 5) interpersonal factors. The research is finally focused on data from 1000 historical sales leads (SL). The data has been classified by proposed AI model. The automated machine learning (AutoML) binary classification model on Google Cloud Platform has been created for experiment, trained by prepared data and used for sales lead future value prediction and confronted to historical outcomes of conversion rates and values.

The rest of the paper is structured as follows: Literature review will outline this research's theoretical background. Research methodology section will develop the research hypotheses, the data, used models, and the paper's research methods. The results from the analyses are illustrated and explained in another section. Then

the paper will discuss this study's results and arguments for them, including some existing and possible implementations. Finally, the conclusion part of the research paper will summarize main points.

LITERATURE REVIEW

Several queries in scientific resources have been done in order to define current knowledge in areas similar to the topic of the paper. They addressed subjects of:

- machine learning in predictive sales pipeline or sales leads analytics,
- machine learning models in sales predictions,
- business analytics for sales pipeline management in the software industry,
- sales forecasting machine-learning techniques,
- B2B sales predictive modeling: machine-learning approach.

Contrary to the author's expectations, few work has been done or released for involving predictive modeling in sales pipeline analytics, presenting similar focus on sales predictions, sales leads in B2B model. However, some research articles have been found, covering similar areas to the paper and share similar attitude to major problem.

The paper by Yan et al. [2015] presents profile-specific, two-dimensional Hawkes processes model, developed to capture the influence from seller's activities on their sales leads to the win outcome, coupled with lead's personalized profiles.

Another work by Bohanec et al. [2017] presents state-of-the-art black-box prediction model which is a novel use of this methodology inside an intelligent system in a real-world case of business-to-business (B2B) sales forecasting.

In general, many of literature found in topic shares the same attitude to an author of this paper:

1. focus on sales leads as a process reflecting whole history and outcome of sales process [Kaplan et al. 1996],
2. focus on B2B model as most suitable for SLs management [Zalocco et al. 2009],
3. use of Machine Learning as most appropriate method for sales forecasting [Huo 2021]
4. use of experiment in real-world case [Rohaani et al. 2022].

However, the most important differences are:

1. in every case, presented solution is large scale solution for big enterprises,
2. none of researchers analyses financial value of SL as a function of time and probability to support the AI model predictions,
3. they are mostly focused on binary result of the whole process (successful / unsuccessful lead).

It needs to be underlined that the first from list above is the most important obstacle to consider proposed by mentioned authors methods in a scale of SME. Yan

et al. even state that their model is deployed and in continual use to a large, global, B2B multinational technology enterprise listed at “Fortune 500”. It took three years and work of eight researchers to establish, train and deploy model according to their findings. The cost of such endeavor is almost equal to the whole year income of the SME company that has been described in given paper. It is obviously impossible to make use of such solutions by SMEs.

Also the key components of given subject has been studied in literature separately in order to observe current findings and understanding of terms: sales leads, sales pipeline, pipeline funnel, artificial intelligence models, automated machine learning and others.

Sales lead can be described as an entity (business or a person) that is not yet a client, however, there are strong signs that they can be. What’s more, sales leads can also be in a form of data, which allows to identify those potential future consumers.

The ideas of a "consumer funnel" or "sales funnel" were created by Peterson [1959], which was first visual image and description of mentioned phenomenon. In order to achieve a result (sales), it is necessary to get results at each of the stages. Peterson [1959] identified four phases of the sale:

- The initial phase is to attract attention;
- The phase of adaptation of the (offer to the needs of the client) – to arouse interest;
- Illustration /explanation phase – to form a desire;
- The selection phase is to achieve action [Stoop 2009].

SL ‘travels’ in time changing its parameters (mostly probability of success as it goes through different stages towards deal) and the average conversion rate of SLs to sales deals is around 10% in the B2B sector [Eitle et al. 2019].

In recent years the topic of Artificial Intelligence (AI) has been widely spoken about, across all different fields. Machine Learning works as a connecting bridge between AI and Data science allowing to automatically draw conclusions from the raw data provided and test hypothesis.

According to D. Guiterrez, we can distinguish two forms of machine learning used in data science: supervised and unsupervised learning [Gutierrez 2015]. Automated machine learning, also called automated ML or Auto ML, is the process of automating time-consuming and repetitive tasks of developing machine learning models. Preprocessing, training and evaluation is an experimental and iterative process that requires several trials before satisfactory results are achieved. Since these tasks tend to be repetitive, AutoML can help automate these steps. In addition to automation, optimization methods are used in the learning and evaluation process to search for and select algorithms and hyperparameters [Shen et al. 2021].

Automated machine learning is the process of automating tasks related to applying machine learning to real-world problems. AutoML potentially covers every stage, from a raw dataset to building a ready-to-deploy machine learning model. It is the next step in the development of artificial intelligence. The work uses the

AutoML service on Google Cloud Platform. AutoML is using Combined Algorithm Selection and Hyperparameter Optimization. According to Feurer et al. the goal of CASH is to find a common algorithm and hyperparameter settings that minimize the training dataset given a set of algorithms and hyperparameters of those algorithms [Feurer et al. 2015]. The authors define AutoML as the problem of automatically (without human input) producing test set predictions for a new dataset within a fixed computational budget. Therefore, the arguments for AutoML use in case of this paper are:

- minimization of the training data set,
- developer must set limits on the resources used in the AutoML optimization process. This budget usually consists of one or a combination of CPU/GPU usage, uptime, and memory usage,
- model access via simple REST-API and JSON.

AutoML limitations:

- AutoML algorithms rely on data to be clean and relevant. Data cleansing and feature engineering are not yet supported by any AutoML approach,
- developing a well-functioning solution can take quite a long time and in extreme cases it may not be possible.

In given example the AutoML runs on Google Cloud Platform, which is a suite of cloud computing services that runs on the same infrastructure that Google uses internally in its end-user products, such as Google Search, Gmail, Google Drive and YouTube.

RESEARCH METHODOLOGY

The company, chosen as a real-life research object has been established in 1989 in Poland. It is an IT company that produces and implements its own software in area of finance and budgeting for other companies and organizations. It has been chosen by certain criteria:

- Operates in B2B model,
- Belongs to SME (small enterprises less than 50 employees, net revenues less than or equal to EUR 10 million),
- Has minimum thousand historical sales leads including full record of its variability and final outcome.

The material selected from the company is historical database of 1496 records containing information on initial state of SL and its result.

Table 1. Example of a record including SL data

Voivodship	Income Planned	Date Planned	Pipeline Funnel	Company Employee	Success
Kujawsko-Pomorskie	12500	2	P4	JK	false

Source: own work based on company's MySQL database with sales leads (SL) data

The data will be used to train model of automated machine learning using binary classification which predicts a binary outcome. In simplest words it is going to predict if new SL is most likely a success or not. After several surveys done among sales persons, several features has been selected. Features are how the model identifies patterns to make predictions. In general they need to be relevant to the problem. Another goal was to minimize its number to reduce the cost.

As a result a dataset has been constructed. It contains transaction details shown in Table1. In general, the more training examples you have, the better your outcome is. The amount of records also scales with the complexity of the problem. Usually for classification problems we need to have 50 x the number of features, but since Google Cloud requires minimum amount of 1000 records, an author decided to use this amount. It could leave the room for more features, but it has to be stressed that the general idea is to predict outcome at initial stage of SL where amount of information on features is limited. It has been decided that finally a record would contain:

- Voivodship - geographical region of a client,
- Income Planned – planned financial flow at the end of process,
- Date Planned - denoted in months prediction of process duration,
- Pipeline Funnel - success probability connected to stage at which the SL is (from P5 - lowest to 0 - certain)
- Company Employee – initials of sales person in charge of SL process,
- Product Name – product acronym from available class,
- How Many Products – how many products certain client has implemented in the past.

THE HYPOTHESES:

1. There is an affordable method of using automated machine learning (binary classification) for sales predictions in small company within the budget less than 1% of its income yearly and with accuracy more than 80%.
2. There is an algorithm which utilizes predictions from above mentioned AutoML binary classification model to calculate the predicted income based on current sales leads in different stages of its life-time. The prediction accuracy would also be more than 80%.

RESEARCH METHODS:

Ad 1. The data have been implemented to the model. The parameters of binary classifications have been set and model was tested. In evaluation phase model metrics have been reviewed with special focus on confidence score and features relevance. Then the model was deployed and made available for use. For 6 months each new sales lead has been tested and received prediction result from the model. It was decided to be invisible for sales persons in order not to interfere with normal procedures. After completion of its life-cycle the prediction result has been confronted with real-life result to measure accuracy. The computation and resources cost has been automatically calculated by Google Cloud Platform and compared to company income in given period.

Ad 2. The method is to find a formula that may calculate predicted income of the company based on the pile of SLs in different stages. Usually such methods are based on every SL probability and predicted value and deliver very poor accuracy. The conceptual work around algorithms and composite measures have been done in order to find improved formula which would return value closest to real.

RESULTS

The data prepared for the model has one thousand records and has a structure shown in the Table 2. It is being exported to GCP in form of CSV file. After importing it had to be verified if the columns contained proper data format numerical or categorical and determine which field is meant to be target output – in our case “success”. Another check has to be made on each feature if the distinct values number is correct, ex. Pipeline Funnel is a probability measure from P5 to 0 then distinct value supposed to be 6. The parameters of binary classifications have been set and model was tested.

Table 2. Example of a record including SL data

AUC PR	AUC ROC	Accuracy	Log Loss	Score Threshold
0.97	0.99	89.35%	0.016	0.5

Source: own work based on company’s MySQL database with sales leads (SL) data

Table 3. Accuracy of implemented binary classification. Model metrics are generated based on the less common being the positive class

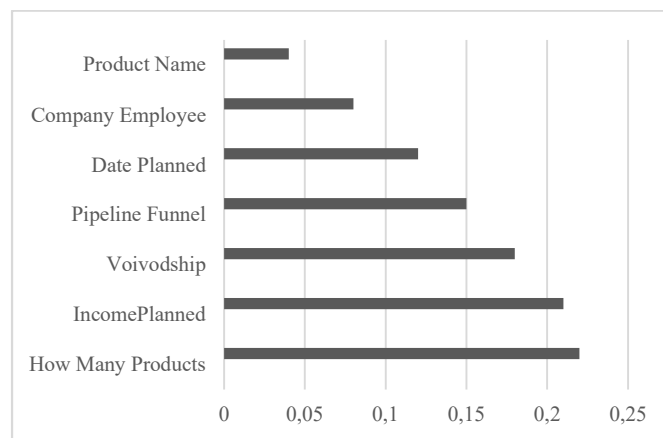
Voivodship	Income Planned	Date Planned	Pipeline Funnel	Company Employee	Product Name	How Many Products	Success
Kujawsko-Pomorskie	12500	2	P4	JK	PZP	2	false

Source: own work based on results shown in Google Cloud Platform model:

aut_ml_example_20190420105640

The model accuracy appeared to be reasonably high: 89,35% and feature importance distribution shows that there is no ‘noise’ in data as shown at the Figure 1. However small importance of product category was a surprise. After deploying the model to GCP user has to feed it with new data for prediction. It has been done by JSON format. Therefore it has been easily connected with existing system for registering SL data status by REST-API web service by http request type; GET. The data of each new sales lead has been stored to database.

Figure 1. Feature importance distribution of implemented binary classification model



Source: own work based on results shown in Google Cloud Platform model:
aut_ml_example_20190420105640

For 6 months each new sales lead has been tested and received prediction result based on described, automated machine learning binary classification model.

Table 4. Accuracy of predictions in given timeframe

	SL created	SL closed	False prediction	True prediction
July 2022	12	7	1	6
August 2022	26	18	2	16
September 2022	34	45	0	45
October 2022	91	57	11	46
November 2022	45	82	17	65
December 2022	22	21	0	21
	Totals:	230	31	199
	Accuracy:	86.52%		

Source: own work based on results taken from Google Cloud Platform model:
aut_ml_example_20190420105640 and stored in company's MySQL database

It was decided to be invisible and inaccessible for sales persons in order not to interfere with normal procedures and achieve highest purity of comparison. After

completion of its life-cycle the prediction result has been confronted with real-life result to measure accuracy as shown in Table 4.

The computation and resources cost has been automatically calculated by Google Cloud Platform and compared to company income in given period.

The average cost of quarterly use of the system is around 480 PLN which meets the criterion of total cost less then 1% of an income, which (together with accuracy) proves the hypothesis number one.

However such prediction even though it is very accurate brings only limited values to the company (selects most promising leads to focus efforts). Very important measure from stakeholder perspective is accurate sales value prediction in order to plan the budget and targets. The method is to find a formula that may calculate predicted income of the company based on the pile of SLs data in different stages. Among the most commonly used features of SLs we find (some of them have not been used in binary classification AutoML model):

- expected sales value V_{ps}
- real sales value V_{rs}
- expected completion date T_{ps}
- real completion date T_{rs}
- customer details
- product details
- salesperson details
- input probability of execution P_k
- output probability of execution P_0

Among the less analyzed features, there are some characteristics of the SL life cycle expressed in probability variability. You can also specify SL monetary value parameters, e.g. expected margin. In our case, for simplicity, we assume a relatively constant average margin. Thus, the amount of the revenue is of stakeholders interest.

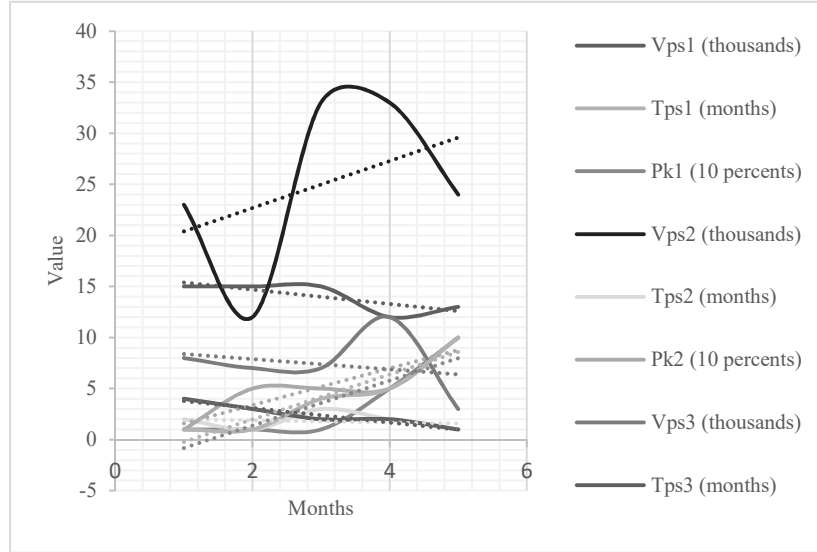
Table 5. Feature values of three chosen SLs over 5 months period

	month 1	month 2	month 3	month 4	month 5
V_{ps1} (thousands)	15	15	15	12	13
T_{ps1} (months)	4	3	2	2	1
P_{k1} (10 percents)	1	1	1	5	10
V_{ps2} (thousands)	23	12	33	33	24
T_{ps2} (months)	2	1	3	2	1
P_{k2} (10 percents)	1	5	5	5	10
V_{ps3} (thousands)	8	7	7	12	3
T_{ps3} (months)	4	3	2	2	1
P_{k3} (10 percents)	1	1	4	5	10

Source: own work based on results taken from company's MySQL database

There are usually around 6 levels of probability. In absolute terms, it is usually assumed that 5 to 10 percent of the SL turns into real sales. Further considerations are aimed at finding relationships between data that affect the predicted value of the sales funnel in the most effective way.

Figure 2. Feature values of three chosen SLs over 5 months period



Source: own work based on results taken from company's MySQL database

The gathered data (Table 5 and Figure 2) shows the great variability and diversity of SL features behavior over time. We clearly see that finding a proper formula for value prediction is extremely complicated without AI models. Traditional methods of calculating SLs value are based on probability and predicted value exclusively and treat probability as a weight for expected value. Sum of all elements gives predicted total value:

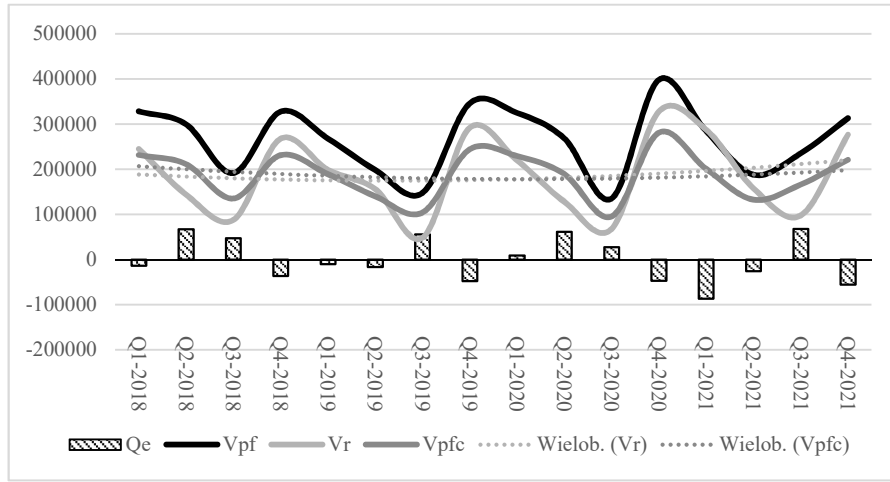
$$V_{pf} = \sum_{i=0}^n (V_{psi} \cdot P_{ki}) \quad (1)$$

Given formula has been tested on historical data from year 2018 to 2021. The assumption is to calculate the expected value of sales at specific points in time, which are the beginnings of the quarters, and the forecast is to apply to all sales for the coming quarter. Therefore, the formula (1) includes only those leads that at given moment have a parameter of expected execution time less than three months. The comparison to each quarter real sales value resulted in calculation on average error of prediction. As shown on Figure 3 the forecasts were quite inaccurate and mostly too optimistic. It seems to be general rule since as it was stated in literature review, usually 90% of SL end up with no sales which results in strong positive feedback on predictions from such a large base of SLs. Therefore we add coefficient which is calculated from average error (σ_M) in percents, updated at the end of each quarter :

$$(T_{ps} \leq 3) \rightarrow V_{pfc} = \sum_{i=0}^n (V_{psi} \cdot P_{ki}) \cdot (1 - \sigma_M) \quad (2)$$

As a result we see that overall trend of predictions and real values of sales in quarters are very similar, however in most periods the predictions have significant errors (Q_e) to real sales value (V_r) and overall accuracy is much less then requested 80%.

Figure 3. Comparison of predictions of sales against real value and ultimate errors in quarters



Source: own work based on results taken from company's MySQL database

Thus an idea appeared to implement the predictions results of AutoML binary classification model into the formula. It was modified in such a way that SLs that have a value of success $P_{aml} = 1$, so they were marked as successful by AutoML binary classification model where excluded from previous formula (2) and calculated differently, using the value of probability P_{true} , calculated as a percentage of successfully predicted sales leads by AutoML (updated after each quarter).

$$(T_{ps} \leq 3) \cap (P_{aml} \neq 1) \rightarrow V_{pfc1} = \sum_{i=0}^n (V_{psi} \cdot P_{ki}) \cdot (1 - \sigma_M) \quad (3)$$

$$(T_{ps} \leq 3) \cap (P_{aml} = 1) \rightarrow V_{pfc2} = \sum_{i=0}^n (V_{psi} \cdot P_{true}) \quad (4)$$

$$V_{pfr} = V_{pfc1} + V_{pfc2} \quad (5)$$

The value of V_{pfr} has been tested on quarters when the AutoML binary classification model has been implemented, so in last two quarters of year 2022. The results are shown in Table 6 and prove hypothesis number two with the reservations described in the discussion.

Table 6. Comparison of predictions of sales against real value with errors and accuracy in last two quarters of year 2022

	Q3-2022	Q4-2022
Vr	127567	321876
Vpfr	147235	289765
Qe	-19668	32111
Accuracy	84.58%	90.02%
Accuracy for both quarters:	87.30%	

Source: own work based on results taken from Google Cloud Platform model: aut_ml_example_20190420105640 and stored in company's MySQL database

DISCUSSION

The major two questions raised by stakeholders of the company were: how to identify sales leads (SLs) that promise better outcomes and how to calculate accurate income predictions at least one quarter ahead. Both hypotheses in the paper were formulated in a similar way and the research achieved its intended goals. However, several points need further clarification before an assessment of the overall outcome can be made.

During the observation of historical data and interviews with sales staff, dependencies were observed indicating that customers from certain regions (e.g., Wielkopolskie and Dolnośląskie voivodeships) had a higher likelihood of conversion. This may be connected to market saturation and the regional diffusion of product information. Some of these factors were not directly included in the predictive formula but were implicitly captured by the binary classification model.

It is worth noting that model validation procedures could be extended in future studies. Although the model achieved a high accuracy score (89.35%), further validation using a clear training/test set split and cross-validation would enhance the robustness of the results. Additionally, possible overfitting should be examined, as the dataset - while sufficient for an initial proof of concept - was relatively small and based on data from one company only.

Some limitations also relate to the specific nature of the case study. The results are drawn from one small IT company operating in Poland, which limits the generalizability of findings to other industries or regions. Furthermore, seasonal effects were not explicitly modeled, and their potential influence on quarterly sales volume remains to be verified statistically. In future, incorporating seasonality and additional classifiers (e.g., quarter identifiers) could improve predictive accuracy and model interpretability.

Overall, despite these limitations, the research provides strong evidence that cost-effective, automated machine learning can substantially enhance decision-making for SMEs in the B2B sector.

CONCLUSIONS

The research confirmed that an affordable and accurate application of automated machine learning (AutoML) can be used to predict sales lead outcomes and estimate company income within the SME sector. The study proved that a binary classification model trained on a dataset of 1000 records achieved high predictive accuracy at a cost below 1% of the company's annual income.

The process of model validation included training on historical data and evaluation based on accuracy metrics. Future extensions should apply explicit training/test data splits, use k-fold cross-validation, and monitor for possible overfitting. Such steps will strengthen the reliability and generalizability of the proposed method.

The paper also introduced an income prediction formula that integrates probabilistic weights with AI-based classification results. The combination of traditional probabilistic modeling and machine learning outcomes proved to improve forecasting accuracy, confirming the second hypothesis.

Nevertheless, the current solution is subject to several limitations, such as the narrow scope of one company, limited dataset size, and potential lack of generalizability to other business environments. Further research should include a wider range of companies, industries, and datasets to build a more universal framework.

In conclusion, the findings indicate that AutoML-based predictive modeling can serve as a practical, low-cost decision-support tool for SMEs. With proper validation and scaling, it holds promise for broader applications in sales management and beyond.

REFERENCES

- Bohanec M., Kljajić-Borštnar M., Robnik-Šikonja M. (2017) Explaining Machine Learning Models in Sales Predictions. *Expert Systems with Applications*, 71(1), 416-428. <https://doi.org/10.1016/j.eswa.2016.11.010>
- Civelek M., Ključnikov A., Fialova V., Folvarčná A., Stoch M. (2021) How Innovativeness of Family-Owned SMEs Differ Depending on Their Characteristics?. *Equilibrium. Quarterly Journal of Economics and Economic Policy*, 16(2), 413-428. <https://doi.org/10.24136/eq.2021.015>
- Eitle V., Buxmann P. (2019) Business Analytics for Sales Pipeline Management in the Software Industry: A Machine Learning Perspective. Retrieved from <http://hdl.handle.net/10125/59541> (19.01.2023)
- Feurer M., Klein A., Eggenberger K., Springenberg J., Blum M., Hutter F. (2015) Efficient and Robust Automated Machine Learning. Retrieved from <https://papers.nips.cc/paper/2015/file/11d0e6287202fcd83f79975ec59a3a6-Paper.pdf> (19.01.2023)
- Gutierrez D. D. (2015) Machine Learning and Data Science: An Introduction to Statistical Learning Methods with R. Retrieved from <https://github.com/linux08/machine-learning-books/blob/master/Machine%20Learning%20and%20Data%20Science%20>

- %20An%20Introduction%20to%20Statistical%20Learning%20Methods%20with%20R%202015.pdf (19.01.2023)
- Huo Z. (2021) Sales Prediction Based on Machine Learning. 2021 2nd International Conference on E-Commerce and Internet Technology (ECIT), 410-415. <https://doi.org/10.1109/ECIT52743.2021.00093>
- Hyder A. S., Sundström A., Chowdhury E. H. (2022) Knowledge of Network-Based Market Orientation for the Internationalization of Disruptive Innovation in SMEs. *Central European Management Journal*, 30(3), 36-60. <https://doi.org/10.7206/cemj.2658-0845.81>
- Kaplan R. S., Norton D. P. (1996) Linking the Balanced Scorecard to Strategy. *California Management Review*, 39(1), 53-79. <https://doi.org/10.2307/41165876>
- Ključnikov A., Civelek M., Klimeš C., Farana R. (2022) Export Risk Perceptions of SMEs in Selected Visegrad Countries. *Equilibrium. Quarterly Journal of Economics and Economic Policy*, 17(1), 173-190. <https://doi.org/10.24136/eq.2022.007>
- Małkowska, A., Urbaniec, M., & Kosała, M. (2021) The Impact of Digital Transformation on European Countries: Insights from a Comparative Analysis. *Equilibrium. Quarterly Journal of Economics and Economic Policy*, 16(2), 325-355. <https://doi.org/10.24136/eq.2021.012>
- Mang'unyi E., Govender K. (2019) Exploring the Development Strategies of Rural Small, Medium and Micro Enterprises. *Central European Management Journal*, 27(4), 95-123. <https://doi.org/10.7206/cemj.2658-0845.11>
- Nogalski B., Niewiadomski P., Szpitter A. (2020) Agility Versus Flexibility? The Perception of Business Model Maturity in Agricultural Machinery Sector Manufacturing Companies. *Central European Management Journal*, 28(3), 57-97. <https://doi.org/10.7206/cemj.2658-0845.27>
- Oleksiuk A., Pleśniak A. (2022) Environment Characteristics and Internationalization of SMEs: Insights from a Polish and Finnish Sample. *Central European Management Journal*, 30(3), 175-194. <https://doi.org/10.7206/cemj.2658-0845.86>
- Peterson A. F. (1959) *Pharmaceutical Selling, "Detailing," and Sales Training*. New York, Heathcote-Woodbridge.
- Plawgo B., Ertman A. (2021) Competency Needs of Industry 4.0 Companies. *Central European Management Journal*, 29(4), 172-195. <https://doi.org/10.7206/cemj.2658-0845.64>
- Rohaani D., Topan E., Groothuis-Oudshoorn C. G. (2022) Using Supervised Machine Learning for B2B Sales Forecasting: A Case Study of Spare Parts Sales Forecasting at an After-Sales Service Provider. *Expert Systems with Applications*, 188, <https://doi.org/10.1016/j.eswa.2021.115925>
- Shen Z., Zhang Y., Lu J., Xu J., Xiao G. (2020) A Novel Time Series Forecasting Model with Deep Learning. *Neurocomputing*, 396, 302-313. <https://doi.org/10.1016/j.neucom.2018.12.084>
- Sleep S., Dixon A. L., DeCarlo T., Lam S. K. (2020) The Business-to-Business Inside Sales Force: Roles, Configurations and Research Agenda. *European Journal of Marketing*, 54(5), 1025-1060. <https://doi.org/10.1108/EJM-06-2018-0416>
- Stoop J. (2009) Developing a Reference Model for KPI and Dashboard Reporting in Sales & Marketing. Retrieved from <http://essay.utwente.nl/60699/> (13.01.2023)

- Soniewicki M. (2022) Is the Importance of Market Orientation Growing? A Study of High-Tech Manufacturing Companies. *Central European Management Journal*, 30(4), 150-174. <https://doi.org/10.7206/cemj.2658-0845.93>
- Zallocco R., Bolman-Pullins E., Mallin M. L. (2009) A Re-Examination of B2B Sales Performance. *Journal of Business & Industrial Marketing*, 24(8), 598-610. <https://doi.org/10.1108/08858620910999466>
- Yan J., Zhang C., Zha H., Gong M., Sun C., Huang J., Chu S., Yang X. (2015) On Machine Learning towards Predictive Sales Pipeline Analytics. *Proceedings of the AAAI Conference on Artificial Intelligence*, 29(1), 1945-1951. <https://doi.org/10.1609/aaai.v29i1.9455>

CAPITAL MARKETS' REACTION TO COVID-19 PANDEMIC AND WAR IN UKRAINE

Dorota Witkowska  <https://orcid.org/0000-0001-9538-9589>

Faculty of Management, University of Lodz, Poland
IPPM College

of Business and Economics

University of Johannesburg, South Africa

e-mail: dorota.witkowska@uni.lodz.pl

Krzysztof Kompa  <https://orcid.org/0000-0002-2810-6654>

IPPM College of Business and Economics

University of Johannesburg, South Africa

e-mail: kkompa@uj.ac.za

Abstract: The paper aims to investigate how the capital markets in different parts of the world reacted on economic problems caused by COVID-19 pandemic and the war in Ukraine. We investigate the daily quotations of 28 stock indexes from October 2019 till the end of 2022. In our research we apply tests comparing expected rates of return and risk in 4 distinguished subperiods. We conclude that Russian-Ukrainian war causes the increase of capital markets volatility but the decreases of index values caused by pandemic was higher than after the Russian invasion in all countries but Russia and Greece.

Keywords: pandemic, war in Ukraine, capital markets, statistical tests

JEL classification: G15, C12

INTRODUCTION

In 2020, the entire world was hit by the COVID-19 pandemic, causing a significant economic slowdown worldwide as a result of lockdowns and disrupted supply chains. To mitigate the economic and social impact of the downturn, countries mobilized their reserves and in 2021 most economies recorded positive GDP growth, which, however, was still negative in many countries if 2019 GDP is taken as the basis. In February 2022 Russia attacked Ukraine, and the effects of the

<https://doi.org/10.22630/MIBE.2025.26.2.8>

war were felt in many countries because the economies of both countries produced in 2019, 3.5% of global GDP, which at the same time was 22.7% of the GDP produced by the European Union, including being significant exporters of food, and Russia additionally of oil and gas.

Effects of the COVID-19 pandemic and the armed conflict in Europe have been affecting capital markets. Our study aims to answer the questions how the selected capital markets reacted on economic problems caused by pandemic and the war in Ukraine together with comparison of the strength and direction of reaction on both events. The novelty of the research is to compare the reaction of capital markets world-wide on mentioned above two dramatic situations. In our research we analyze data concerning 28 blue-chip indexes of selected world-wide capital markets from the period 30.09.2019-31.12.2022, applying statistical tests.

LITERATURE REVIEW

There is quite rich literature concerning world-wide crisis caused by the COVID-19 pandemic and Russia-Ukraine war. In the works [Aslam et al. 2020] and [Topcu, Gulal 2020] their authors find contagion transmission from the original source (i.e., China) to the global, emerging and the Asian stock markets. Contessi and De Pace (2021) identify periods of mildly explosive dynamics and collapses in the stock markets of 18 countries during the first wave of the COVID-19 pandemic in 2020. They find statistical evidence of instability transmission from the Chinese stock market to all other markets. In studies [Liu et al. 2020], [Liu et al. 2022b], [Zhang et al. 2020] and [Chowdhury et al. 2020], the authors examine the pandemic induced shock on the stock market volatility.

Abuzayed et al. [2021] look the contagion of systemic risk spillover of pandemic crises between the global stock market and stock markets hit hardest by the pandemic and found that systemic risk contagion intensifies more between the global stock market and individual markets between North American and Europe than Asian stock markets. Foo and Witkowska (2024) show that the COVID-19 pandemic financial crisis had a similar impact on 12 the analyzed capital markets (5 developed and 4 post-communist European and 3 non-European). Pappas et al. [2013] analyzing CEE markets in terms of their synchronization and contagion, find a significant strengthening of the correlation between markets (especially for young EU members). Whereas Żebrowska-Suchodolska et al. [2021] analyze situation in European post-communist countries during pandemic in comparison to the most developed European stock markets.

At the beginning of 2022, countries were still dealing with the effects of the pandemic when the world was hit by the news that Russia had invaded Ukraine. Therefore, the question arises about the impact of Russian-Ukrainian war on financial and commodity markets all over the world, and if this impact has global or local character. Balbaa et al. [2022] and Ozili [2022] claim that Russia's invasion

and sanctions imposed on Russia have a spillover effect on the global economy because it caused energy and commodity supply shocks, rising the prices of energy, food and commodities. Izzeldin et al. [2023] reveal that the intensity of the post-invasion crisis is noticeably smaller compared to both the Covid-19 and the 2007 global financial crisis.

Boungou and Yatie [2022] show a negative relationship between the Ukraine-Russia war and world stock market returns which was the largest during the first two weeks after the invasion of Ukraine. They also notice that these effects were most pronounced for countries bordering Ukraine and Russia, as well as for those UN member states that demanded an end of the Russian offensive in Ukraine.

Karamti and Jeribi [2023] claim that the Russo-Ukrainian conflict's influence on global equity markets is heterogeneous and based on countries' economic-political connection or proximity to the war zone. They investigate the G7 countries (i.e., Canada, France, Germany, Italy, Japan, UK and the USA) which sanctioning Russia and 3 non-sanctioning countries i.e., Russia, China and India (abr. RIC). Their results suggest that the war has a stronger influence on the G7 stock and the whole financial markets through commodity prices, with a greater impact of natural gas and wheat prices for this group of states than in the case of the RIC countries.

Literature review shows that the impact of pandemic and Russian aggression on capital markets is usually analyzed separately [Li et al. 2024; OECD 2022]. It is also worth mentioning that majority of publications consider the war in Ukraine in terms of oil, gas and food supply.

METHODOLOGY AND DATA

In our research we look at capital markets in some countries which were selected because of their position in world-wide economy, their role or situation caused by pandemic or Russia-Ukrainian war. For instance, USA is the biggest financial market whereas France, Germany and UK are the most developed capital markets in Europe. China is the biggest world economy but capital market there is still treated as an emerging one. Pandemic started in China in 2019, and it was spread all over the world. Israel was the leader in covid vaccination thus it was expected to recover after the lockdown fast. Russia invaded Ukraine on February 24, 2022, that affects economies in states which have borders with Russia or their economy depends on Russian oil and natural gas.

We investigate the daily quotations of stock indexes in 19 European and 9 non-European capital markets for the period from the third quarter of 2019 till the end of 2022. The indexes selected for analysis represent:

- (1) European post-communist countries (Poland, Czechia, Hungary, Romania, Bulgaria, Russia and Croatia),
- (2) other European countries (the United Kingdom, Germany, Holland, Finland, Norway, Denmark, Iceland, France, Italy, Spain, Portugal and Greece) and

- (3) non-European countries (the USA, Australia, New Zealand, China, India, Singapore, Turkey, South Africa and Israel).

In our research we distinguish four subperiods of investigated capital markets behavior. The first one (denoted hereafter as I) is the pre-pandemic period contains the time span from 30.09.2019 to the date when the stock index reached the maximum peak value in the first quarter of 2020. The second subperiod (II), called the pandemic shock is defined as the time span between the dates when the maximum and minimum values of stock indexes in 2020 were observed¹. The third subperiod (III) - pandemic stabilization and post-pandemic (pre-war) period covers the time span from the date of minimal index value to the beginning of war, and the last period (IV) from February 24 to the end of 2022 is a war period.

To investigate the situation in the selected capital markets in the subperiods distinguished above, we verify the following hypothesis to identify the market trend.

$$H_0: \mu = 0, \quad (1)$$

with the test statistics:

$$u = \frac{\bar{y}}{S} \sqrt{T}, \quad (2)$$

where for the distinguished subperiods, μ - the expected return, $\bar{y}$ - average rates of return μ_i obtained in the sample, S - standard deviation of returns in the sample, T - number of observations in the certain period of time.

We also test if there are significant changes in the diverse capital markets in the distinguished subperiods by verifying the following hypotheses:

$$H_0: \mu_1 = \mu_2, \quad (3)$$

$$H_0: \sigma_1^2 = \sigma_2^2, \quad (4)$$

The test statistics for both hypotheses are as following:

$$t = \frac{\bar{y}_1 - \bar{y}_2}{\sqrt{\frac{S_1^2}{T_1} + \frac{S_2^2}{T_2}}}, \quad (5)$$

$$F = \frac{S_1^2}{S_2^2}, \quad (6)$$

where, $\mu_1, \mu_2, \sigma_1^2, \sigma_2^2$ - the expected returns, and variances of returns in the two comparable subperiods, $\bar{y}_1, \bar{y}_2, S_1^2, S_2^2$ - average rates of return and variances in the samples, T_1, T_2 - number of observations in both compared samples.

The proposed research methodology (1)–(6) is used for analyses that aim to determine the strength and direction of reactions to specific events [Tarczyński, Witkowska, & Kompa 2013]. And since the goal of our research is to draw general conclusions, we decided that using simple analytical tools instead of event studies or Perron tests would suffice.

¹ Two first subperiods are distinguished following the idea presented in [Żebrowska-Suchodolska et al. 2021].

EMPIRICAL FINDINGS

Analysis of quotation of selected stock indexes was used to distinguished four subperiods which are used for further investigation. In all considered countries, but 9 states, the highest quotation appeared in the third period i.e., after pandemic shock i.e., before the Russian invasion in Ukraine. In Spain global maximum is observed before lockdown (19.02. 2020). While in the remaining countries, global maximum appeared during the war, i.e. in the year 2022 - South Africa (2.03.), Iceland (4.04.), Singapore (5.04.), Norway (13.04.), UK (8.04.), Portugal (7.06.), India (1.12.) and Turkey (30.12.). Minimal quotations during lockdown appeared in majority of markets between 12th and 23rd of March 2020 except Greece (26.03.), and during the war between 29th of September and 14th of October 2022 with exception of Turkey (28.02.), Bulgaria (7.03.), China (26.04.), India (20.06.), Greece, Australia and New Zealand (all three 20. 06.), Singapore (21.10.), Israel (25.12.) and USA (28.12). It is noticeable that during pandemic minimal values of indexes were smaller than during the war in all countries except Greece and Russia.

Table 1. Test statistics to verify hypotheses (1), (3) and (4)

Periods	I	II	III	IV	I	II	III	IV
	BUX Hungary				PX Czechia			
$E(R)=0$	1.852	-2.081	1.516	0.233	1.859	-3.398	2.965	-0.453
Compared periods	I&II	II&III	III&IV	I&IV	I&II	II&III	III&IV	I&IV
$E(R_t)=E(R_{t+1})$	2.539	-2.290	0.423	1.084	3.662	-3.738	1.796	1.462
$D^2(R_t)=D^2(R_{t+1})$	9.197	4.536	2.060	4.176	16.853	5.063	1.539	5.1219
	WIG20 Poland				CROBEX Croatia			
$E(R)=0$	0.234	-2.625	1.144	0.079	1.130	-2.160	2.404	0.000
Compared periods	I&II	II&III	III&IV	I&IV	I&II	II&III	III&IV	I&IV
$E(R_t)=E(R_{t+1})$	2.567	-2.786	0.482	0.111	2.282	-2.339	1.430	0.862
$D^2(R_t)=D^2(R_{t+1})$	7.398	3.080	1.440	3.4596	36.787	14.607	1.261	1.997
	BET Romania				SOFIX Bulgaria			
$E(R)=0$	1.645	-2.086	2.915	-0.368	0.568	-2.475	2.297	0.149
Compared periods	I&II	II&III	III&IV	I&IV	I&II	II&III	III&IV	I&IV
$E(R_t)=E(R_{t+1})$	2.283	-2.385	1.722	1.222	2.538	-2.696	0.924	0.308
$D^2(R_t)=D^2(R_{t+1})$	30.941	7.423	1.475	6.146	13.007	8.380	1.629	2.529

Note: Bolded values indicate rejection of the null hypothesis at the $\alpha=0.05$ level. Shading denotes events when the standard deviation in the first period compared is smaller than in the second period. The Roman marks I–IV denote the research periods defined in the text.

Source: own calculations on the basis of data in Appendix

Testifying results (Tables 1-4) show that in the covid shock period (i.e., the second one) there were bear markets in all considered 28 countries except Greece, China and the USA, where rates of return were insignificantly different from zero.

Table 2. Test statistics to verify hypotheses (1), (3) and (4)

Periods	I	II	III	IV	I	II	III	IV
OMX Helsinki Finland				AEX Holland				
$E(R)=0$	1.651	-3.309	1.897	0.105	1.089	-2.993	2.162	0.000
Compared periods	I&II	II&III	III&IV	I&IV	I&II	II&III	III&IV	I&IV
$E(R_t)=E(R_{t+1})$	3.533	-3.500	0.894	1.017	3.111	-3.174	1.110	0.724
$D^2(R_t)=D^2(R_{t+1})$	10.726	4.190	1.196	3.063	13.535	5.776	1.221	2.861
OMX Oslo Norway				ICEX Iceland				
$E(R)=0$	0.551	-2.296	1.980	0.000	1.985	-1.726	3.577	-0.238
Compared periods	I&II	II&III	III&IV	I&IV	I&II	II&III	III&IV	I&IV
$E(R_t)=E(R_{t+1})$	2.361	-2.551	1.044	0.392	2.197	-2.200	1.892	1.618
$D^2(R_t)=D^2(R_{t+1})$	10.355	4.164	1.134	2.821	7.955	5.040	1.550	2.446
OMX Copenhagen Denmark				DAX Germany				
$E(R)=0$	2.282	-2.600	2.271	0.674	1.253	-3.113	1.720	0.100
Compared periods	I&II	II&III	III&IV	I&IV	I&II	II&III	III&IV	I&IV
$E(R_t)=E(R_{t+1})$	2.929	-2.817	0.579	1.066	3.237	-3.247	0.841	0.751
$D^2(R_t)=D^2(R_{t+1})$	9.221	4.607	1.256	2.513	15.538	5.699	1.086	2.962
PSI Portugal				IBEX35 Spain				
$E(R)=0$	1.444	-2.617	0.382	0.537	1.119	-2.821	1.086	0.119
Compared periods	I&II	II&III	III&IV	I&IV	I&II	II&III	III&IV	I&IV
$E(R_t)=E(R_{t+1})$	2.731	-2.637	-0.220	0.515	2.906	-2.889	0.567	0.688
$D^2(R_t)=D^2(R_{t+1})$	28.502	8.004	1.131	3.148	25.000	7.716	1.349	2.403
FTSEMIB Italy				FTSEATHEX20 Greece				
$E(R)=0$	1.640	-2.777	1.545	-0.094	1.046	-1.626	0.582	0.688
Compared periods	I&II	II&III	III&IV	I&IV	I&II	II&III	III&IV	I&IV
$E(R_t)=E(R_{t+1})$	2.899	-2.868	0.936	1.139	1.863	-1.703	-0.499	-0.513
$D^2(R_t)=D^2(R_{t+1})$	25.331	8.297	1.026	2.977	3.0112	6.725	5.104	103.361
CAC40 France				FTSE United Kingdom				
$E(R)=0$	1.008	-2.899	1.957	0.106	0.484	-2.667	1.576	0.286
Compared periods	I&II	II&III	III&IV	I&IV	I&II	II&III	III&IV	I&IV
$E(R_t)=E(R_{t+1})$	2.984	-3.039	0.974	0.566	2.706	-2.888	0.775	0.185
$D^2(R_t)=D^2(R_{t+1})$	20.881	6.943	1.044	3.141	10.563	3.449	1.526	2.007

Note: Bold, shading and Roman marks as in Table 1.

Source: own calculations on the basis of data in Appendix

Also returns from stock indexes in the first and the third subperiods are significantly bigger than in the second one in all stock exchanges but SSE (China) and NASDAQ100 (USA) where differences of returns in comparable periods are insignificant. After Russia invasion, returns from all considered indexes insignificantly differ from zero with exception of XU100 (Turkey), which generated significantly higher returns during the war than in the third period. Comparison of returns in the pre-pandemic to the pandemic stabilization period shows insignificant differences between these both periods in all considered markets, whereas returns in the war period are significantly higher than during the covid shock period in all countries but Iceland, the USA, Greece and China. Chinese capital market was the only one which shows stable situation in the whole analyzed period.

Table 3. Test statistics to verify hypotheses (1), (3) and (4)

Periods	I	II	III	IV	I	II	III	IV
TA35 Israel				NASDAQ 100 USA				
$E(R)=0$	1.387	-2.741	2.085	-0.383	3.086	-1.232	2.150	-0.642
Compared periods	I&II	II&III	III&IV	I&IV	I&II	II&III	III&IV	I&IV
$E(R_t)=E(R_{t+1})$	2.899	-2.929	1.481	1.221	1.447	-1.366	1.533	2.059
$D^2(R_t)=D^2(R_{t+1})$	16.834	5.890	1.054	2.713	44.560	11.140	1.772	7.088
S&PNZX50 New Zealand				S&ASX20 Australia				
$E(R)=0$	1.565	-3.316	1.680	-0.225	0.924	-6.075	1.659	0.309
Compared periods	I&II	II&III	III&IV	I&IV	I&II	II&III	III&IV	I&IV
$E(R_t)=E(R_{t+1})$	3.488	-3.453	1.228	1.328	6.106	-6.279	0.828	0.502
$D^2(R_t)=D^2(R_{t+1})$	13.1406	5.572	1.467	1.608	3.240	1.2656	1.530	1.673
SSE China				STI Singapore				
$E(R)=0$	1.222	-1.117	1.297	-0.503	1.828	-2.535	1.760	0.000
Compared periods	I&II	II&III	III&IV	I&IV	I&II	II&III	III&IV	I&IV
$E(R_t)=E(R_{t+1})$	1.410	-1.310	1.087	1.227	3.122	-2.923	1.211	1.660
$D^2(R_t)=D^2(R_{t+1})$	7.807	3.610	1.323	2.860	7.739	2.341	2.041	1.620
SA40 TOP40 South Africa				XU100 Turkey				
$E(R)=0$	1.206	-2.318	2.394	0.100	1.514	-2.494	2.409	4.542
Compared periods	I&II	II&III	III&IV	I&IV	I&II	II&III	III&IV	I&IV
$E(R_t)=E(R_{t+1})$	2.526	-2.656	1.182	0.737	2.885	-2.948	-2.555	-0.348
$D^2(R_t)=D^2(R_{t+1})$	11.595	3.800	1.119	3.416	3.179	1.930	1.109	1.826
SENSEX India				MOEX Russia				
$E(R)=0$	1.477	-2.216	2.636	0.846	2.756	-2.854	0.214	0.283
Compared periods	I&II	II&III	III&IV	I&IV	I&II	II&III	III&IV	I&IV
$E(R_t)=E(R_{t+1})$	2.452	-2.564	1.071	0.556	3.366	-2.815	-0.150	0.833
$D^2(R_t)=D^2(R_{t+1})$	19.076	4.987	1.667	2.294	11.725	1.181	1.410	14.006

Note: Bold, shading and Roman marks as in Table 1.

Source: own calculations on the basis of data in Appendix

Risk, measured by standard deviation, is significantly bigger in the pandemic shock period than in the pre-pandemic one in all countries. Similar relation appears if risk in the pandemic shock period compares to the pandemic stabilization period with exception of Greece where situation is opposite and Russia together with Australia which show insignificant differences of risk in both compared periods.

Table 4. Test statistics to verify hypotheses (3) and (4)

Compared periods	Country	$E(R_t) = E(R_{t+1})$	$D^2(R_t) = D^2(R_{t+1})$	Country	$E(R_t) = E(R_{t+1})$	$D^2(R_t) = D^2(R_{t+1})$
I&III	BUX Hungary	0.979	19.94	BET Romania	-0.567	440.35
II&IV		-2.045	358.77		-1.977	2081.44
I&III	PX Czechia	-0.250	119.99	SOFIX Bulgaria	-0.509	63.76
II&IV		-3.207	672.34		-2.456	448.96
I&III	WIG20 Poland	-0.343	26.40	OMX Helsinki Finland	0.292	80.39
II&IV		-2.519	113.50		-3.272	164.63
I&III	CROBEX Croatia	-0.319	2152.90	AEX Holland	-0.301	122.95
II&IV		-2.149	850.64		-2.961	272.97
I&III	OMX Oslo Norway	-0.564	83.34	PSI Portugal	0.860	1039.72
II&IV		-2.238	137.97		-2.658	634.71
I&III	ICEX Iceland	0.198	26.35	IBEX35 Spain	0.194	1136.69
II&IV		-1.612	152.00		-2.821	343.65
I&III	OMX Copenh. Denmark	0.710	53.90	FTSEMIB Italy	0.345	674.97
II&IV		-2.677	134.12		-2.757	609.97
I&III	DAX Germany	0.000	204.60	FTSEATHEX20 Greece	0.000	2.87
II&IV		-3.100	284.90		-1.359	236.24
I&III	CAC40 France	-0.399	399.84	S&PNZX50 New Zealand	0.321	371.70
II&IV		-2.891	475.50		-3.270	80.22
I&III	FTSE United Kingdom	-0.498	259.79	S&ASX20 Australia	-0.215	24.59
II&IV		-2.670	47.91		-5.989	4.48
I&III	TA35 Israel	-0.112	314.99	SSE China	0.425	34.85
II&IV		-2.659	254.95		-0.942	106.60
I&III	NASDAQ 100 USA	0.861	632.35	STI Singapore	0.970	249.42
II&IV		-1.141	6234.79		-2.425	14.38
I&III	SA40 TOP40 South Africa	-0.361	107.30	SENSEX India	-0.394	1011.69
II&IV		-2.274	168.42		-2.324	130.90
I&III	XU100 Turkey	0.233	8.22	MOEX Russia	1.573	69.14
II&IV		-3.837	12.43		-2.685	273.40

Note: Bold, shading and Roman marks as in Table 1.

Source: own calculations on the basis of data in Appendix

Also risk in the pre-pandemic period is smaller than during the pandemic stabilization in all markets but Greek. Risk in the last period is significantly higher than in the pandemic shock period in all considered countries. Similarly, when it is compared to the pre-pandemic period but with the exception of Germany. Whereas

comparing risk during the war to the one in pandemic stabilization period, it is significantly higher in 12, lower in 7 and insignificant in 9 states.

CONCLUSION

Our research shows that the decreases of index values caused by COVID-19 pandemic were higher than after the Russian invasion in all countries but Russia and Greece. Although, the expected returns from indexes in the war period were significantly higher than in the pandemic shock period in all analyzed markets but Greece, Iceland, the USA and China. After Russia invasion, returns from the majority of considered indexes insignificantly differ from zero, except NASDAQ100 and XU100 (Turkey) which reported significantly positive expected returns in the war period. Expected rates of return during pandemic were much smaller than during the war. Investigating returns from the selected markets it is noticeable that the highest returns in all periods with the positive maximum values appear in non-European markets whereas the smallest returns in the post-communist European countries. In the covid shock period there were bear markets in all considered countries except Greece, China and the USA. During the war, stagnation of capital markets is visible with exceptions of Turkey. The USA and Singapore create the group of markets with significantly higher returns during the war than in the pre-pandemic period.

Risk measured by standard deviation is significantly smaller in all countries in the pre-pandemic period comparing to both the second and third periods with exception of Russia and Greece for the latter. In the war period risk was bigger than in the pre-pandemic period in all considered markets. Although risk in the pandemic shock period was significantly bigger than risk in the last period for all markets.

To sum up our research, it is visible that Russian-Ukrainian war causes perturbances and increase of volatility of capital markets but the impact of the war on selected markets is significantly smaller than Covid-19 pandemic what is in line with results obtained by Izzeldin et al. (2023). After invasion, rates of return from stock indexes observed in countries which were dependent on Russian oil sent via pipelines or neighboring Russia do not significantly differ from zero. Although indexes BUX (Hungary), PX (Czechia) together with OMX from Norway and Finland characterized negative asymmetry. Our results therefore do not indicate a significantly stronger response from countries that appear to be more affected by the Russia-Ukraine war due to their proximity to hostilities and dependence on Russian and Ukrainian supplies.

REFERENCES

- Abuzayed B., Bouri E., Nedal Al-Fayoumi N., Naji Jalkh N. (2021) Systemic Risk Spillover Across Global and Country Stock Markets During the COVID-19 Pandemic. *Economic Analysis and Policy*, 71, 180-197. <https://doi.org/10.1016/j.eap.2021.04.010>

- Aslam F., Mohmand Y., Ferreira P., Memon B.A., Khan M., Khan M. (2020) Network Analysis of Global Stock Markets at the Beginning of the Coronavirus Disease (Covid-19) Outbreak. *Borsa Istanbul Review*, 20-S1, 49-61
- Balbua ME., Eshov M., Ismailova N. (2022) The Impacts of the Russian-Ukrainian War on the Global Economy (Preprints). <https://doi.org/10.13140/RG.2.2.14965.24807>
- Boungou W., Yatie A. (2022) The Impact of the Ukraine- Russia War on World Stock Market Returns. *Economics Letters*, 215(3). <https://doi.org/10.1016/j.econlet.2022.110516>
- Chowdhury E. M., Khan I. I., Dhar B. K. (2020) Catastrophic Impact of Covid-19 on the Global Stock Markets and Economic Activities. *Business and Society Review*, 127(2), 437-460. <https://doi.org/10.1111/basr.12219>
- Contessi S., De Pace P. (2021) The International Spread of Covid-19 Stock Market Collapses. *Finance Research Letters* (42). <https://doi.org/10.1016/j.frl.2020.101894>
- Foo J., Witkowska D. (2024) The 2020 COVID-19 Financial Crisis Impact on the European Stock Markets and Economies. A Preliminary Analysis. *Folia Oeconomica Stetinensia*, 24(1), 22-40. <https://doi.org/10.2478/9oli-2024-0002>
- Izzeldin M., Muradoglu Y., G., Pappas V., Petropoulou A., Sivaprasad S. (2023) The Impact of the Russian-Ukrainian War on Global Financial Markets. *International Review of Financial Analysis*, 87, 102598. <https://doi.org/10.1016/j.irfa.2023.102598>
- Karamti Ch., Jeribi A. (2023) Stock Markets from COVID-19 to the Russia–Ukraine Crisis: Structural Breaks in Interactive Effects Panels. *Journal of Economic Asymmetries*, 28, e00340. <https://doi.org/10.1016/j.jeca.2023.e00340>
- Li P., Zhang P., Guo Y., Li J. (2024) How Has the Relationship Between Major Financial Markets Changed During the Russia–Ukraine Conflict?. *Humanit Soc Sci Commun*, 11, 1731 (2024). <https://doi.org/10.1057/s41599-024-04231-7>
- Liu H., Manzoor A., Wang C.Y., Zhang L., Manzoor, Z. (2020) The COVID-19 Outbreak and Affected Countries Stock Markets Response. *International Journal of Environmental Research and Public Health*, 17, 2800. <https://doi.org/10.3390/ijerph17082800>
- Liu Y., Wei Y., Wang Q., Liu Y. (2022) International Stock Market Risk Contagion During the COVID-19 Pandemic. *Finance Research Letters*, (45), 102145. <https://doi.org/10.1016/j.frl.2021.102145>
- OECD (2022) Impacts of the Russian Invasion of Ukraine on Financial Market Conditions and Resilience: Assessment of Global Financial Markets, OECD Publishing, Paris. <https://doi.org/10.1787/879c9322-en>
- Ozili P. K. (2022) The Global Economic Consequence of the Russian Invasion of Ukraine. *SSRN Electronic Journal*, 1-7. <https://doi.org/10.2139/ssrn.4064770>
- Pappas V., Ingham H., Izzeldin M., Steele G. (2013) Financial Markets Synchronization and Contagion: Evidence from CEE and Eurozone. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.2411978z>
- Tarczyński W., Witkowska D., Kompa K. (2013) Współczynnik Beta. *Teoria i Praktyka. Pielaszek Research*.
- Topcu M., Gulal O.S. (2020) The Impact of COVID-19 on Emerging Stock Markets. *Finance Research Letters*, 36, 101691. <https://doi.org/10.1016/j.frl.2020.101691>
- Zhang D., Hu M., Ji Q. (2020) Financial Markets Under Global Pandemic of COVID-19. *Finance Research Letters*, 36, 101528. <https://doi.org/10.1016/j.frl.2020.101528>

Żebrowska-Suchodolska D., Karpio A., Kompa K. (2021) COVID-19 Pandemic: Stock Markets Pandemic: Stock Markets Situation in European Ex-communist Countries. *European Research Studies Journal*, 24(3), 1106-1128. <https://doi.org/10.35808/ersj/2408>

APPENDIX

Table A1. Daily rates of return characteristics in the sample 30.09.2019 – 31.12.2020

Subperiods	I	II	III	IV	I	II	III	IV
Index	BUX (Hungary)				PX (Czechia)			
Average	0.0022	-0.0079	0.0009	0.0003	0.0012	-0.0129	0.0014	-0.0004
Stan. dev.	0.0092	0.0279	0.0131	0.0188	0.0057	0.0234	0.0104	0.0129
Asym.	0.0693	-1.3533	-0.5783	-0.7838	-0.4840	-1.3154	0.6616	-0.2275
Kurtosis	-0.5281	3.9067	7.1792	5.6244	1.7177	1.7943	7.5454	2.7268
N	60	54	487	213	78	38	485	213
Index	WIG20 (Poland)				CROBEX (Croatia)			
Average	0.0003	-0.0102	0.0008	0.0001	0.0006	-0.0093	0.0008	0.0000
Stan. dev.	0.0100	0.0272	0.0155	0.0186	0.0046	0.0279	0.0073	0.0065
Asym.	0.6862	-2.2085	-0.4428	0.5382	0.9483	-1.5016	-0.6205	0.3117
Kurtosis	1.7911	8.0242	6.0657	1.2197	4.6765	3.9814	18.4768	3.8858
N	61	49	491	213	75	42	481	209
Index	BET (Romania)				SOFIX (Bulgaria)			
Average	0.0009	-0.0087	0.0013	-0.0003	0.0004	-0.0084	0.0008	0.0001
Stan. dev.	0.0048	0.0267	0.0098	0.0119	0.0061	0.0220	0.0076	0.0097
Asym.	-0.4191	-0.8812	-0.0362	0.0876	1.4346	-2.5621	-0.0174	-0.6485
Kurtosis	0.7325	2.9913	5.5685	7.0379	4.3928	9.3294	3.1676	11.6569
N	77	41	483	213	75	42	476	208
Index	AEX (Holland)				FTSEMIB (Italy)			
Average	0.0009	-0.0186	0.0012	0.0000	0.0015	-0.0318	0.0021	-0.0002
Stan. dev.	0.0081	0.0298	0.0124	0.0137	0.0091	0.0458	0.0183	0.0200
Asym.	-0.5428	-1.3709	0.6209	0.0621	-0.5938	-2.0352	0.3486	-2.2742
Kurtosis	1.9667	2.1253	5.7059	1.1986	1.4366	4.3984	3.6187	19.8472
N	96	23	499	218	99	16	499	217
Index	DAX (Germany)				CAC40 (France)			
Average	0.0011	-0.0236	0.0011	0.0001	0.0008	-0.0234	0.0012	0.0001
Stan. dev.	0.0086	0.0339	0.0142	0.0148	0.0079	0.0361	0.0137	0.0140
Asym.	-0.5424	-1.3192	0.8297	0.6141	-1.1017	-1.0903	0.5035	0.5225
Kurtosis	2.2941	2.1669	8.4758	3.6001	3.2576	1.3993	5.8567	3.3736
N	96	20	493	218	99	20	499	218
Index	OMX Helsinki (Finland)				OMX Oslo (Norway)			
Average	0.0014	-0.0170	0.0011	0.0001	0.0005	-0.0090	0.0011	0.0000
Stan. dev.	0.0080	0.0262	0.0128	0.0140	0.0078	0.0251	0.0123	0.0131
Asym.	-0.3021	-1.5992	0.1489	-0.1249	-0.2434	-1.4098	0.1018	-0.2658
Kurtosis	0.6220	3.1254	3.8120	0.6918	0.5903	2.8851	2.8359	0.4432
N	89	26	487	214	74	41	490	214
Index	ICEX (Iceland)				OMX Copenhagen (Denmark)			
Average	0.0018	-0.0056	0.0016	-0.0002	0.0019	-0.0135	0.0012	0.0006
Stan. dev.	0.0078	0.0220	0.0098	0.0122	0.0082	0.0249	0.0116	0.0130
Asym.	0.2969	-0.7588	-0.0421	0.1041	-0.3735	-0.6263	-0.0452	-0.1009

Kurtosis	-0.0672	1.0529	3.0602	0.5455	1.0876	0.8949	0.4850	0.4037
N	74	46	480	211	97	23	482	213
Index	MOEX (Russia)				XU100 (Turkey)			
Average	0.0021	-0.0102	0.0002	0.0005	0.0021	-0.0084	0.0018	0.0053
Stan. dev.	0.0066	0.0226	0.0208	0.0247	0.0125	0.0224	0.0161	0.0170
Asym.	-0.0787	-1.7613	-8.5645	2.3574	-0.3063	-1.3360	-1.4611	-0.2163
Kurtosis	-0.0462	3.3485	134.375	22.503	4.8284	2.7359	7.8626	2.5552
N	75	40	493	196	79	44	482	213
Index	FTSE100 (United Kingdom)				PSI (Portugal)			
Average	0.0004	-0.0092	0.0009	0.0002	0.0009	-0.0189	0.0002	0.0004
Stan. dev.	0.0072	0.0234	0.0126	0.0102	0.0062	0.0331	0.0117	0.0110
Asym.	-1.0369	-1.5528	0.5396	-0.2865	0.0635	-0.6332	0.6623	-0.1704
Kurtosis	4.7843	3.5309	6.0469	1.4234	1.1136	0.6770	3.9610	0.4159
N	76	46	487	212	99	21	498	218
Index	IBEX35 (Spain)				FTSEATHEX20 (Greece)			
Average	0.0009	-0.0266	0.0007	0.0001	0.0012	-0.0060	0.0012	0.0049
Stan. dev.	0.0080	0.0400	0.0144	0.0124	0.0102	0.0177	0.0459	0.1037
Asym.	-0.4589	-1.2656	0.6864	0.0346	0.0814	-2.5049	3.7710	1.9275
Kurtosis	1.3076	2.4428	4.9674	1.3333	0.8880	9.1084	74.620	13.762
N	99	18	499	218	79	23	495	212
Index	S&PNZX50 (New Zealand)				S&PASX20 (Australia)			
Average	0.0008	-0.0132	0.0006	-0.0001	0.0007	-0.0171	0.0009	0.0002
Stan. dev.	0.0056	0.0203	0.0086	0.0071	0.0075	0.0115	0.0120	0.0097
Asym.	-0.8385	-0.4375	1.2349	0.0920	-0.9998	0.2534	0.4426	-0.2611
Kurtosis	2.5178	0.3253	9.2568	0.7875	1.6326	0.4214	6.9833	1.2487
N	120	26	580	255	98	23	489	224
Index	NASDAQ100 (USA)				SENSEX (India)			
Average	0.0024	-0.0135	0.0033	0.0018	0.0012	-0.0095	0.0016	0.0006
Stan. dev.	0.0077	0.0514	0.0183	0.0207	0.0068	0.0297	0.0133	0.0103
Asym.	-0.3861	0.2280	0.1642	-0.5904	0.0710	-1.8269	0.4368	0.0150
Kurtosis	0.5425	0.1838	2.8010	7.7712	0.7953	6.0087	7.1502	0.2524
N	98	22	487	214	70	48	480	211
Index	SSE (China)				STI (Singapore)			
Average	0.0010	-0.0032	0.0015	0.0007	0.0019	-0.0040	0.0008	0.0000
Stan. dev.	0.0068	0.0190	0.0113	0.0121	0.0055	0.0153	0.0100	0.0070
Asym.	-0.2633	-1.3950	0.3006	-0.8771	-0.9531	-1.9204	0.7951	-0.1497
Kurtosis	0.2648	3.9653	4.2698	7.6720	0.4921	7.3469	6.4841	-0.1028
N	69	44	467	209	28	94	484	213
Index	SA 40TOP 40 (South Africa)				TA35 (Israel)			
Average	0.0011	-0.0094	0.0015	0.0001	0.0010	-0.0150	0.0004	-0.0003
Stan. dev.	0.0079	0.0269	0.0138	0.0146	0.0068	0.0279	0.0126	0.0112
Asym.	-0.1063	-1.6978	1.0583	0.4016	-0.6714	-0.0278	-0.5039	-0.2712
Kurtosis	-0.2537	2.8948	7.1556	0.8359	3.3810	-0.4379	6.1161	-0.1099
N	75	44	485	211	89	26	475	204

Source: own calculations on basis of <https://stooq.pl/>