

**ECONOMICS
WORKING
PAPERS**

VOLUME 7

NUMBER 6

ISSN 1804-9516 (Online)

2023

ECONOMICS WORKING PAPERS

Volume 7 Number 6 2023

Publisher: University of South Bohemia in České Budějovice
Faculty of Economics

Reviewers: **doc. Ing. Eva Vávrová, Ph.D.**
Masaryk University
Faculty of Economics and Administration

doc. RNDr. Marta Žambochová, Ph.D.
Jan Evangelista Purkyně University in Ústí nad Labem
Faculty of Social and Economic Studies

Edition: 11, 2023

ISSN: 1804-9516

ECONOMICS WORKING PAPERS

EDITORIAL BOARD:

CHAIRMAN:

Ladislav Rolínek

University of South Bohemia in České Budějovice
Czech Republic

EDITORS:

Eva Cudlínová, University of South Bohemia
in České Budějovice, Czechia

Miloslav Lapka, University of South Bohemia
in České Budějovice, Czechia

Ivana Faltová Leitmanová, University of
South Bohemia in České Budějovice, Czechia

Tomáš Mrkvička, University of South
Bohemia in České Budějovice, Czechia

Darja Holátová, University of South Bohemia
in České Budějovice, Czechia

Ladislav Rolínek, University of South
Bohemia in České Budějovice, Czechia

Milan Jílek, University of South Bohemia in
České Budějovice, Czechia

ASSOCIATE EDITORS:

Věra Bečvářová, Mendel University in Brno,
Czechia

Věra Majerová, Czech University of Life
Sciences Prague, Czechia

Roberto Bergami, Victoria University,
Melbourne, Australia

Cynthia L. Miglietti, Bowling Green State
University, Huron, Ohio, United States

Ivana Boháčková, Czech University of Life
Sciences Prague, Czechia

Ľudmila Nagyová, Slovak University
of Agriculture in Nitra, Slovakia

Jaroslava Holečková, University
of Economics in Prague, Czechia

James Sanford Rikoon, University
of Missouri, United States

Lubor Lacina, Mendel University in Brno,
Czechia

Labros Sdrolas, School of Business
Administration and Economics Larissa, Greece

Daneil Stavárek, Silesian University in Opava,
Czechia

ECONOMICS WORKING PAPERS. Published by Faculty of Economics. University of South Bohemia in České Budějovice • The editor's office: Studentská 13, 370 05 České Budějovice, Czech Republic. Contact: tel: 00420387772493, Technical editor: Markéta Matějčková, e-mail: matejckova@ef.jcu.cz • ISSN1804-5618 (Print), 1804-9516 (Online)

Content

1	Introduction.....	6
2	Literature Review.....	10
2.1	Enlargement of the European Union to new countries in 2004 and 2007.....	10
2.2	Solvency frameworks.....	11
2.3	Financial integration and convergence towards a single European market for insurance.....	13
2.4	Functional clustering methods in big data analysis.....	14
2.5	Literature on Cluster analysis implementation.....	18
2.6	Cluster analysis features.....	20
2.7	Effects of Covid-19 on the European insurance industry.....	20
3	Objectives and hypotheses.....	23
4	Methodology.....	25
4.1	Data.....	25
4.2	Evaluation of other methodologies.....	25
4.3	Methodological framework.....	26
4.3.1	Dissimilarity Matrix.....	27
4.3.2	Combined Functional Ordering.....	27
4.3.3	Global envelopes and Functional Ordering with Intrinsic Graphical Interpretation.....	28
4.3.4	Dissimilarity Matrix Based on the Combined Ordering.....	32
4.3.5	Joint functional clustering for national insurance markets of Europe.....	33
5	Conclusions.....	33
6	References.....	35

European Insurance Market Analysis via a Joint Functional Clustering Method

Athanasiadis, S.

Abstract: The enlargement of the European Union (EU) to include Central and South-Eastern European countries in 2004 and 2007 launched an integration process that unifies the economies and financial markets of member states and enables the convergence of these two areas. This study focuses on analyzing the development and similarity of the European insurance sector after the EU enlargement. We study 34 European insurance markets from 2004 until 2021 based on a certain set of indicators that characterize insurance markets, such as Insurance Density, Insurance Penetration and Gross Written Premiums to name a few. With a functional clustering method applied to such indicators, we try to reveal whether there are similarities between the individual countries that could explain the European insurance market homogeneity and convergence via the EU integration process. The proposed method has also a practical importance since it provides visualization of the clustering results through the construction of global envelopes. This study supports the works of EU policy makers that have a major impact on the ability of further integration of the European insurance market.

Keywords: European insurance markets, enlargement EU, convergence, insurance integration, insurance market indicators, functional clustering, global rank envelope

JEL Classification: C38, F15, F36, G22

1. Introduction

Insurance companies manage risks that individuals are exposed to in various aspects of life - finance, healthcare, transportation, property, professional liability, etc. Insurance companies collect premiums from their customers and aggregate them into "risk pool" funds to manage these risks. From these funds, the insurers compensate those customers who sustain losses attributable to an insured event. One could also say that insurance is a bet against a risk, where the insurer provides financial cover on its occurrence (Franzetti, 2021; Jarrow, 2021). Thus, insurers are vital participants in the modern economy, and the very concept of risk mitigation through risk-pooling and collective contribution can significantly promote the general welfare of society, especially when looking at various forms of social and health insurance (Aussilloux et al., 2017; Busemeyer & Iversen, 2020).

One of the longstanding goals of the European Union (EU) is to create a single European market (SEM), as the removal of economic barriers should allow for better trade flow and a larger market with more competition. Studies have shown that the absence of barriers can increase the average GDP across participating markets by almost 10%. Albeit there is a "strong degree of heterogeneity across EU countries", which may make combining their economies less efficient (In 't Veld, 2019). Including the insurance sector in the SEM, we get the single European insurance market (SEIM). Efforts by European governments and the European Commission are in place to integrate insurance providers into the SEIM; however, this is not a straightforward process (McGee, 2020). Recent studies show that the EU insurance market remains divergent across member and non-member states, and its economics are not set significantly apart from insurance markets outside the EU (Jagric, Bojnec, & Jagric, 2018).

There are multiple potential barriers against the homogenization of the insurance market across Europe, and they might be connected to general economic challenges that the insurance industry faces and that can be quite specific to the national economies and governmental measures that insurers must navigate (Fournier *et al.*, 2015). Moreover, the insurance industry is facing systematic changes in its operational requirements and approaches to business functionality. For example, the allocation of reserve capital — a fundamental feature of every insurance company — is changing towards principles-based approaches. This is especially the case after the introduction of

the *Solvency II* (SII) framework that governs the solvency capital requirement (SCR) within the European Union (European Commission, 2015). SII is a replacement of *Solvency I* (SI), integrating lessons learnt from the initial framework. Specifically, SI provided too simplistic model to accurately estimate risks, and it did not lead to an efficient allocation of capital. To solve these shortcomings, SII is built on three major pillars, namely, quantitative requirements and protocols to accurately assess asset and liability values; specific risk management and governance protocols; and transparency and public disclosure to eliminate informational market asymmetries and thus better competition. In summary, the protocols and rulesets contained within SII are geared towards stochastic algorithms and simulations to better estimate required capital reserves; they are less of a 'one-size fits all' model than a nuanced approach that allows the insurers more flexibility in allocating their capital, and avoids having either too many or too few capital reserves, by making more accurate predictions about the funds needed (Clark, 2009; American Academy of Actuaries, 2016; Singh, 2021). Moreover, the insurance industry faces global challenges — changing economic systems, a diverging risk landscape, regulatory acts, and other trials. In the context of establishing a SEIM, national economies, governmental regulation, and the risk adversity of the national culture can pose additional problems (Gaganis *et al.*, 2019). Nevertheless, considering the notion that the European insurance market constitutes roughly one-third of the worldwide market share, the EU's attempts to move towards an SEIM model are significant in the global context (Jagric, Bojnec, & Jagric, 2018).

In 2004 and beyond, the EU was expanded from 15 countries ("Old Union") to 25 countries by including Eastern European nations like Hungary, Poland, and Cyprus ("New EU"; Toshkov, 2017). As the new member states often have a lower insurance market penetration than older ones, there is a rapid growth potential for the insurance market across the EU. At the same time, there are calls for the EU to both facilitate and regulate insurance market progression to homogenize growth across Europe (Müller-Reichart, 2005). This is not a trivial task. For example, between 2009 and 2016, insurance penetration, i.e., insurance premiums as % of GDP, fluctuated significantly, suggesting an inhomogeneous insurance market across the EU (Jagric, Bojnec, & Jagric, 2018). Other parameters that can be used to determine how homogenous the insurance market is across several nations, i.e., the EU, are the "number of insurance companies, the market concentration indicator, gross premiums written, total assets, insurance density ratio, life- and no in-life

insurance share of the total insurance market" (Denkowska & Wanat, 2020, p. 1). The analysis of these parameters will likely reflect that macroeconomic developments also play a role in insurance market effectiveness: economic downturns are associated with a decline in insurance demand, while interest rates and other financial factors can also add stress to the insurance industry and capital allocation (European Commission, 2014).

There is a growing use of artificial intelligence algorithms and applications in financial services (Hilpisch, 2020). Several methods can be used to analyze the degree of homogeneity among the European insurance industry. Jagric, Bojnec, & Jagric (2018) utilize optimized spiral spherical self-organizing maps to analyze parameters like average insurance premiums, total premiums to GDP ratio, and the number of insurance companies per capita (Jagric & Zunko, 2013). Utilizing an artificial intelligence/machine learning (AI/ML) method to analyze the data has the advantage that it can deal with highly multidimensional data spaces and variables that exhibit complex multicollinearities. Others, such as Denkowska & Wanat (2020), utilize Hellwig's development index, together with several unsupervised classification methods, such as k-means clustering, partitioning among medoids (PAM), and Ward's method, to analyze the similarity of the insurance markets across European countries. In both studies, the authors find that the European insurance market is rather heterogeneous and divergent, contrary to the goals of a unified SEIM.

Dai, Athanasiadis, & Mrkvička, (2021) have developed a functional clustering (FC) method that utilizes combined dissimilarity sources and allows for graphical interpretation of large volumes of multidimensional input data. It can be used for several purposes, for example, population growth or, pertaining to the context of this study, insurance penetration of markets. The advantage of this FC method is that it is not bound to any model and that it is nonparametric; this means that it is robust and not vulnerable towards skewed distributions or outliers in the data (Dai, Athanasiadis, & Mrkvička, 2021). This is an important feature, as economic data — especially when averaged on the national level — is complex and does not necessarily follow a specific pattern or distribution (Kirman, 2006).

In this context, the main goal of this study is to contribute to the existing literature regarding the challenge of the SEIM, by proposing a methodological procedure that will enable a more de-

tailed description and understanding of the mechanisms towards the European insurance development unification. To this end, we introduce a novel FC method that is able to provide evidence for accepting or rejecting the hypothesis on the homogeneity and convergence of the European insurance market. This FC method captures the similarities and differences between national insurance development levels and paths, as it divides 34 European countries into clusters with shared characteristics in terms of several insurance indicators that describe insurance market in these countries. The clustering solution may reveal the impact of the EU enlargement with countries of different economic level on the current EU insurance market status. It is likely also to reveal the changes in insurance development in the post-pandemic period compared to the time before the pandemic outbreak that may attributed to the Covid-19 pandemic. If there is enough evidence to reject that homogeneity and convergence hypothesis, we suggest analyzing the significance of this and the underlying reasons, and propose economic, regulatory, and policy measures that may help make the European insurance market more homogenous. In an effort to explain the case of an inhomogeneous European insurance dynamics, we also suggest exploring the significance of some economic, financial, sociodemographic factors affecting insurance market development and demand within each cluster produced.

Following this introduction, the rest of this study is organized on the basis that it serves as a guide on how to approach our main goal and on the methodology employed without presenting research results. These results along with their interpretation will be found in future work projects. This limitation is not intended to restrict the goals of this study but rather helps to make them more concise. First, in Section 2 we give an overview of the European Union enlargement, the EU financial integration and convergence to SEIM, the solvency II framework and the functional clustering methods as reported in existing literature. In Section 3, we present the objectives and hypotheses setting, while in Section 4, we reveal the data basis used and the rationale for the chosen methodology of the research. In the end, Section 5 concludes this study by summarizing our goals, contributions, and methodology framework, while discussing future research directions.

2. Literature Review

2.1 Enlargement of the European Union to new countries in 2004 and 2007

The European Union went through one of its largest extensions between 2004 and 2007. It admitted ten new countries, namely the Czech Republic, Cyprus, Estonia, Latvia, Lithuania, Hungary, Malta, Poland, Slovenia, and Slovakia, and recognized Bulgaria, Romania, and Turkey as candidates or potential candidates for a further expansion of the EU; Bulgaria and Romania were indeed added in 2007 (Heidbreder, 2011). These expansions added a dozen new nations with vastly different economic environments to an alliance of fifteen nations that had already established a common market for decades. Naturally, one would expect the new, expanded EU to be economically quite heterogenous, at least in the immediate years after the expansion. In addition, this new expanded economy added ca. 73 million people to a total of 454 million citizens, an increase of 16% (Wilson Center, 2004). Naturally, one of the most important goals for the EU government was to implement rules and regulations so that there could be homogeneity and balance across the European markets. Importantly, a recent study found that the EU expansion from 2004 to 2007 resulted in aggregate gains for welfare, albeit these gains were not always realized for all skill groups (Caliendo *et al.*, 2021).

Overall, the main impacts of Eastward enlargement on the EU – in terms of the decision-making process functionality and compliance with EU law – have not been negative, although concerns about corruption issues among new members have certainly become much more visible since the 2004 enlargement, especially in countries like Bulgaria or Romania (Sedelmeier, 2014). Despite previous concerns, the increase in the number of workers from Central and Eastern Europe was not reflected in a significant westward movement, and there was no displacement of local workers or a reduction of their wages (Drew & Sriskandarajah, 2007). However, such concerns may remain when looking at a potential further expansion of the EU to include Turkey or even such countries as Ukraine (Tsependa & Hurak, 2021). These factors are important to consider, as a country's population, and to a certain extent, its labor force can impact a countries insurance market as well.

Originally, the legal framework of the SEIM was enacted in June of 1994 (Ennsfellner & Dorfman, 1998). It included provisions for a single insurance license across EU member states, an

overarching regulatory system, and home state control. The latter is an important principle of European single market regulation, whereby cross-country services and products — generated in one country and applied in another one — are legally treated under the law of the *originating* country (Polański, 2018). This would mean that a German insurance company could sell its products to France without the need to adhere to French regulations. Of course, such a distribution of insurance products is predicated either on the freedom of services or freedom of establishment (Köhne & Brömmelmeyer, 2018). Specifically, this is reflected in the *insurance distribution directive* (IDD). Importantly, it should be noted that despite the welcome regulatory clarity of the IDD, questions remain as to the potential overregulation of the market (Köhne & Brömmelmeyer, 2018).

In this context, it is also prudent to mention the newly created *pan-European Personal Pension Product* (PEPP), a pension platform that is thought to complement national pension products for individuals. As pensions can be seen as a subcategory of life insurances, the PEPP is another sign of steps towards SEIM. The PEPP can be offered by insurance providers, credit institutions, and other financial organizations. Interestingly, the PEPP is a testament to how the European Union not only aims at regulating the market by providing a legal framework that enables a SEIM; it also actively provides resources that promote and facilitate the establishment of such a unified market. The legal provisions of the PEPP have only recently gone into effect, more specifically, on the 22nd of March 2022, so this is a very recent development, and it shows that the European Union is indeed actively involved in the unification of the European insurance market (EIOPA, 2022).

2.2 Solvency frameworks

To implement a common regulatory framework for a single European insurance market, the European Union has generated several guidelines as part of the solvency I and II frameworks.

Solvency I (SI) was the first attempt by the European Union (EU) at implementing specific solvency requirements for insurers and came into force at the beginning of 2002. This framework has been developing since the 1970s and was the first major regulatory attempt at creating a SIEM; it was aimed at clarifying solvency rules for insurance and reinsurance providers across Europe. One of the major requirements of insurance providers is their ability to pay out larger sums of money to satisfy and honor contractual claims, and therefore, insurance companies must be able to provide enough liquid funds (Atlas Magazine, 2021). According to SI, with about 65% of assets

needing to conform to certain technical provisions, 25% of assets must be held in a liquid form, with 10% as a free surplus that can exist in any type of asset. However, different insurance types require different amounts of liquidity, which is solved by the solvency II framework explained below.

Like its predecessor, the solvency II framework primarily focuses on minimum capital requirements that insurance companies must hold going forward, so that they cannot become overwhelmed by a sudden increase in claims and stay solvent. Solvency II was implemented in the beginning of 2016 (Doff, 2016). In addition to default risk, defining a certain threshold above which capital must be held also provides an early warning mechanism for regulators and auditors that an insurance company may become unable to cover their clients' claims in the future, and thus also engenders confidence towards the ability of insurers to stay solvent. Thus, a large part of SEIM is a common reserve capital requirement. One major difference to solvency I is the adoption of a prudence margin on top of liquid holding, depending on how much liquidity is needed for a particular sector of the insurance economy (Zlateva, 2022).

In terms of policy updates, instead of replacing the Solvency II framework with Solvency III, the *European Insurance and Occupational Pensions Authority* (EIOPA) has undertaken a review of the Solvency II framework and provided the results in the form of a proposal to amend Solvency II, a “communication on the review of the Solvency II directive”, and a proposal for a new directive concerning insurance recovery and resolution (European Commission, 2021). By and large, the review aims to be in line with EU priorities, such as enabling funds to finance the recovery after the 2020 SARS-CoV-2 pandemic, completing the union of capital markets, and mobilizing assets for the *European Green Deal*. In addition, the review provides tools to help the European insurance sector become more resilient towards crises in the future. The fact that the European Commission decided to recommend modifications and updates of Solvency II rather than producing a Solvency III framework suggests that the core of Solvency II is already sufficient to promote the implementation of an SIEM.

2.3 Financial integration and convergence towards a single European market for insurance

There are several indicators for describing the status of the insurance industry, which can then be utilized to address whether the European insurance market is homogeneous and correlate it with measures of economic development across the EU.

However, while insurance market concentration may not correlate with overall economic homogeneity, one would assume that general economic factors do align with economic changes in the development of the insurance sector over time. For example, the part of the GDP that is due to the insurance market contribution is consistently growing and correlated with the overall economic development level in any given country (Ostrowska-Dankiewicz & Simionescu, 2020). However, economic growth and insurance market development are two processes the relationship of which can be ambiguous. For example, in alignment with the *supply-leading hypothesis* (SLH), the insurance market can impact economic growth, while other studies find that in alignment with the *demand-following hypothesis* (DFH), economic growth can positively influence insurance market development. Other hypotheses posit that there is no causal relationship between economic growth and the insurance market, or that there is feedback between both (Peleckienė *et al.*, 2019). Carefully controlled case studies will likely be needed to analyze the relationship between the insurance sector and the overall economy; it may as well be that the relationship is time- and context-dependent, given that markets often behave in complex and not entirely predictable manners.

The *insurance penetration* (IP) index is another parameter that can be used to describe the development of the insurance market. This index is simply the percentage of gross written premiums as compared to the total GDP of a given country (Kwon & Wolfrom, 2016). Becmer (2015) found that the index has developed similarly between the original 15 European countries and the countries of Central and Eastern Europe during and after the EU expansion. Thus, any differences between old and new EU countries in the IP index are insignificant, suggesting that there is some homogeneity across Europe indeed. Importantly, the analysis by Becmer (2015) also suggests that one should look at the dynamics or trends of the insurance sector over time rather than deriving an aggregate score, as individual differences over time do exist. Moreover, it may also be instructive to look at the life- and non-life insurance sector separately, as the latter shows more differences

between the old and new EU countries than the life insurance industry (Becmer, 2015). Interestingly, when applying clustering analysis to IP across European countries over time, strong economic development of insurances is correlated with preceding unstable market behavior, suggesting that the economic status of the insurance industry could be an indicator of market health (Athanasiadis & Mrkvička, 2018).

The *insurance density* (ID) index is the ratio of total insurance premiums to the overall population; one could also describe it as premiums *per capita* (Kwon & Wolfrom, 2016). Unlike the IP index, the ID ratio does not depend on a nation's economic productivity and could therefore be used as more of an independent factor when correlating insurance market development to overall economic factors. When comparing old vs. new EU states, the ID index shows similar behavior to the IP ratio - it is overall similar, with some local differences based on the analyzed year (Becmer, 2015). Han *et al.* (2010) found that an increase in the ID ratio was causally connected to economic growth, supporting the supply-leading hypothesis mentioned above. Given the complex nature of market behavior, it is likely that further studies might be required to confirm this causal connection. However, it shows that the density ratio can be used as an economic indicator to describe the insurance industry and its market.

2.4 Functional clustering methods in big data analysis

One of the central problems with large data volumes is that they are often multidimensional and cannot be easily described with static stochastic methods, such as linear or logistic regression. Beyond a certain number of parameters and data points, traditional mathematical analysis techniques will fail to adequately find trends and characteristic behaviors out of a subset of data, which is also often noisy; the signal-to-noise ratio is too low to extract any significant information, even if it is present. Thus, AI/ML techniques can be used to dynamically analyze the data and to have the computer recognize certain patterns without applying a set of relationships and laws *a priori* (Grout, 2021).

Clustering is a useful data analysis technique when classifying and organizing data and falls under the category of unsupervised learning (Léger & Mazzuco, 2021). There are many variations of this technique, and at their core, data is classified and grouped based on some similar property

or correlation between the data; the greater the similarity between the data, the better the corresponding effect of the cluster analysis. Unlike many other statistical methods, yet in line with the premise on unsupervised learning, cluster analysis is often used when no *a priori* assumptions can be made about possible relationships between data (Ma, 2021). It is generally not known how many clusters exist in the data until the model is run. The clustering or grouping can be based on the size and shape of the data points, or any other parameter, and the resulting clusters can often provide a reasonable graphical interpretation of the data at large.

Clustering, usually applied following a PCA analysis, is often used in unsupervised learning. The reason for this association is that the observed data are organized in uniform groups, and in clusters, without knowing the labels for them beforehand. As typical exploratory analyses, cluster analyses have been involved in large numbers of studies ranging from biology, meteorology, and climate science. In addition, in both meteorology and climate science, clustering algorithms may be performed over data of daily weather forecasts or historic measured climate records, to identify likely regions of similar weather patterns, or track changes of the climate over a wide region divided into smaller sub regions. There are two types of clustering algorithms, hierarchical and non-hierarchical clustering. The major difference between the two is that, in hierarchical clustering, after the observations are assigned to one cluster, they cannot be moved to any other clusters as the assigned groups are combined over iterations; in non-hierarchical clustering, on the other hand, an observation may be assigned to several clusters prior to a final decision on the clustering. Hierarchical clustering is a kind of classic clustering technique which defines cluster members by comparing differences or distances among observations. In many setups, N observations are used, and a distance matrix $N \times N$ is formed, which has the distance measurements for all pairs of observations. Then, the clustering algorithm is applied on the distance matrix, such that observations are clustered together in sequence according to their order of their distance measurements.

On the other hand, the non-hierarchical clustering algorithms consist of two subtypes of clustering algorithms, namely, partitioned clustering and model-based clustering. The best-known partitioning clustering algorithm is K-means. With a fixed number of clusters, K , the K-means algorithm seeks to find K clusters in which variance within each cluster is minimized, and variance across clusters is maximized. In each iteration of the algorithm, the results for the clusters are re-

estimated on every observation until all observations are within a cluster of whose mean they are closest. For model-based clustering, mixture distributions are fitted to observations, and the parameters estimation on probability density functions (PDFs) and clustering procedures is obtained by maximizing a likelihood function.

Time-series data is a kind of multidimensional, high-dimensional data, as its size can be considered as either the duration of observed time, or the number of time points observed. The number of time points may be exceptionally large for longer-term observations or for higher-frequency records, such as inventory data, daily temperature records for years, and machine-based long-term tracking. In the cluster analyses for non-supervisory studies, the two procedures are typically used on the time-series data. One is to perform clustering method on raw (or smoothed) discrete time series data directly, while the other is to perform the clustering method after the transformation of the time series data to function data. Clustering methods for discrete time series data are classified as raw data methods and distance-based methods.

The classic idea is to treat every observation as a multivariate case, so each time point is a variable, and to apply a multivariate clustering technique. Clustering on time series data is more complex and challenging because of high dimensionality of the input space. Two major difficulties are the accuracy and the speed of the clustering. For clustering accuracy, since data noise has sensitive effects on time-series data clustering, some statisticians recommend that data should first be smoothed prior to clustering, so that the noise is reduced and that underlying patterns are captured over time.

The second clustering procedure treats the time series data as functional data. In functional data, every observation x is defined as a function of time t , that is, $x(t)$; in other words, functional data is considered a collection of curves of infinite dimensions. Data smoothing is the first step in all functional data fusion methods, performed by fitting each time-series data set to linear combinations of certain basis functions, for example, spline functions or polynomials.

Clustering procedures can then be performed on these smoothed data using functional data clustering methods, which can be broadly classified in two subgroups: distance-based methods and filter-based methods. Like the multivariate situation, distance-based methods in the functional case

perform clustering according to distances between functional objects. Distance-based methods entail constructing a distance matrix with a particular measure; the results of the clustering may be obtained using a hierarchical clustering or centroids-based tool. Filtering methods are methods which use the post-dimensionality information from functional data to cluster observed objects. In filtering methods, rather than directly clustering by fit functions after the time-series data are smoothed, certain information about features at lower dimensions of the data is extracted to depict the initial data and to decrease dimensions for further clustering.

In other words, the filtering-based methods involve approximating a curve by linear combinations of finite-basis functions, such as spline functions and functional principal components, while clustering analyses are performed on coefficients or scores with finite dimensions. Filtering methods exploit the functional principal component scores (FPC scores). This method is completed by using the other well-known dimension reduction method, Functional Principal Component Analysis (FPCA). FPCA is the extension of principal components analysis (PCA) for functional data. Like PCA, FPCA identifies the directions which account for most variation in data. These directions are equivalent to the eigenfunctions. Then, FPC estimates for each observation can be obtained from the eigenfunctions. Besides representing data via the coefficients of B-spline, the FPC scores are also recommended to reduce dimensions (Lin, 2019).

The focus of recent studies is on distance-based methods, and chosen functional orders (Dai, Athanasiadis, & Mrkvička, 2021). This involves constructing a similarity matrix by a chosen functional ordering applied to a set of differences between all pairs of the functional data being studied. Various functional orderings may be chosen, but the focus is often on those that have an internal graph interpretation. This allows the resultant clusters, as well as a center region achieving the natural interpretation, to be shown. This means that all functions contained within a central region do not exit the plot of a central region, and all functions that are not contained within a central region exit the plot of a central region in at least one spot. To get a certain balancing among different sources, one can standardize the curves prior to applying existing methods, such as the k-means or model-based methods. The term "standardizing" refers to the empirical marginal distributions to have a zero mean and a unity of variance. If function sorting is applied, each function portion is treated uniformly, while different sources of variance are combined uniformly.

For univariate cases, where only univariate functions are explored, original curves and the derivatives can be combined in the same way to measure both magnitude and form variations at once. For multivariate cases, it may be better to combine the marginal curves and covariance functions equally to capture both marginal and joint variations between the curves. Furthermore, this approach provides reasonable visual explanations for clustering results and inherits the robustness of the functional sorting, and it is capable of recovering clusters in case that anomalous observations corrupt data.

2.5 Literature on Cluster analysis implementation

In an illustrative case study, functional cluster analysis of administrative regions within communities is used to determine whether boundaries between neighborhoods are established in an optimal way (Martí *et al.*, 2021). The analysis of functional characteristics showed that the spatial distribution of activities of several dozen functional groups was more uniform than that of the administrative regions. Thus, cluster analysis can be used to manage a city more efficiently by dividing the municipality into significant functional units, rather than arbitrarily selecting borders based on administrative responsibilities. Such an analysis may even be useful to counteract the partisan method of gerrymandering and replace it with a subdivision of counties that work more efficiently, resulting in bigger contentment of citizens with their Government (Vagnozzi, 2020). Moreover, a functional division of city governments that better reflects the spatial distribution of economic activities, services, and structures of important micro-regional units can increase the overall efficiency of the administration in charge (On the one hand, the functional diversity and the number of economic activities, services, and structures of each cluster compared to other clusters were analyzed using the refined dataset; on the other hand, the functional specialization of each cluster was performed using the Google Places API The secondary category is determined (Martí *et al.*, 2021).

To solve challenges for economic policymaking, Warsame (2021) reported on a hierarchical cluster analysis based on a distance-based approach with a semi-metric measure using principal functional components. In this case, a Bayesian Information Criterion (BIC) with positive log probability was defined to model complexity as the number of parameters (Warsame, 2021). The algorithm was able to reveal socioeconomic structures within the analyzed clusters that would not

have been able to determine using marginal distributions and dependency structure of the data alone. In this analysis, the trend of BIC values and the stability of cluster sizes were tested by initializing the algorithm classes with the k-means function and different initial values (Warsame, 2021).

Similar approaches can be applied to variables of mixed types that are often found in clustering analysis applications in economic analyses (Hennig & Liao, 2013). Often, common elements can be found in the algorithm's input or clustering variables. Using a k-means clustering algorithm is an optimization process, as the objective function exceeds the global minimum in some local range. It should therefore be noted that k-means clustering can only provide local optimization to process continuous variable data. The k-means clustering algorithm must specify the number of clusters before clustering, and the clustering results are easily dependent on the initial clustering. On the other hand, the k-means cluster analysis method can effectively avoid the subjective negative impact of artificial thresholds, to more accurately and objectively distinguish the state range of various financial risks. Clustering of k-means is often used in functional and structural analysis, but the limitation of the k-means algorithm is that it is mainly used for scalar and not vectorial data (Zhu & Liu, 2021).

The goal of cluster analysis is to find groups of similar subjects, where "similarity" between each pair of subjects means a global measure over the complete set of characteristics. Therefore, a suitable application for cluster analysis is data segmentation strategies. The most important task of choosing an adequate clustering methodology by the underlying philosophy of the dataset and context is the translation of the desired interpretation value of clustering into analytical characteristics of the clustering method data (Martí *et al.*, 2021).

Often, data volume can be reduced before applying cluster analysis through means of factor analysis, which reduces the overall variables towards significantly correlated groups (Cheung, 2020). After standardizing the data, clustering can be performed using various libraries that are widely available in most computing languages, such as R package or Python. For high-density genetic data, a two-way clustering algorithm can better extract local information. Using functional data methods there is a greater variety of clustered arrays to choose from (Liu, Zhang, & Feng, 2021).

2.6 Cluster analysis features

When it comes to real-time data, due to the high speed and continuous characteristics of the underlying data streams, the use of traditional cluster analysis algorithms can be difficult.

Instead, bidirectional analysis methods can be used (Liu, Zhang, & Feng, 2021). In this context, it is important to mention that cluster analysis methods are considered static procedures because the inclusion of new observations or variables can change clusters, making it mandatory to develop a new analysis. Cluster analysis is a set of extremely useful research methods that can be applied whenever it is necessary to test for similar behavior between observations (individuals, companies, municipalities, countries, etc.) concerning certain variables. It is therefore also well suited to test for homogeneity of data, a parameter that is difficult to establish using traditional stochastic means. Moreover, the development of cluster analysis does not require extensive knowledge of matrix algebra or statistics, unlike methods such as factor analysis and correspondence analysis (Fávero & Belfiore, 2019).

When researchers' primary goal is to classify and group observations, they may decide to develop a cluster analysis and then analyze the ideal number of clusters to form. If the groups of some data are known in advance, it is best to use discriminant function analysis to find the variables and matrices that best rank the remaining observations. Researchers often use the wrong random weighting procedure and qualitative variables (such as Likert scale variables) and then apply cluster analysis. Regardless of the purpose, clustering will always be an exploratory technique (Fávero & Belfiore, 2019).

Grouping changes detected in dynamic data flow analysis play a significant role in the analysis of the confluence of medical and healthcare economic data. Because of the enormous potential of gene expression analysis useful for disease diagnosis, drug development, and life science research, a two-way clustering algorithm is often widely used in gene expression data studies (Mardaneh, 2015).

2.7 Effects of Covid-19 on the European insurance industry

One of the most prevalent events of the years 2020 and 2021 was the worldwide SARS-CoV-2 or Covid-19 pandemic. As the virus caused significant disruptions to the economy and

global supply chains, affecting the entire range from small businesses to large corporations, the insurance industry was also impacted (Puławska, 2021); however, for insurance undertakings, one may assume that these disruptions were not only felt on the level of normal business operations, but also in the form of an increased volume of claims (Babuna *et al.*, 2020). Herein lies a specific challenge for the insurance sector, as this volume increase was brought about in a much more rapid and abrupt way than usual. This could be seen by a significant decrease in the return on assets among German and Italian insurers, and a similar decrease in solvency ratios among Belgian, German, and French insurance companies (Puławska, 2021). Of course, as such changes are not seen in insurance companies throughout all European nations — for example, Poland's insurers were not significantly affected by such variations —, they could pose an additional challenge to a unification of the European insurance market. On the other hand, these challenges may provide additional evidence to justify further managerial involvement of the European Union to regulate and eventually synchronize the European insurance market on the way to a unified SEIM (Puławska, 2021).

In general, the European Commission – and EIOPA in particular – have recognized that the pandemic crisis has led to serious adverse developments that threaten the orderly functioning and integrity of financial markets, and the stability of the EU financial system. In accordance with its responsibilities for monitoring and assessment of market developments, EIOPA conducted several targeted assessments in support of its understanding and decisions taken, considering the substantial impact that the crisis will have on financial markets and on the real economy. In addition, EIOPA assessed the situation and the consequences in a thorough manner, within the framework of its Crisis Prevention and Management Framework, as well as under article 18 of its founding resolution. In that context, it took a few facilitation and coordination actions within the NCAs to respond to adverse developments. This work also fed into regular public-available risk analysis products, which provided a detailed assessment of the situation and its impacts.

EIOPA publishes the Risk Dashboard quarterly, as well as the Financial Stability Report on two occasions per year. The December 2020 Financial Stability Report provides more detailed information about key risks facing the European insurance industry (EIOPA, 2020). Furthermore, the report provides an integrated framework which combines key market risk drivers - interest

rates, credit spreads, equity, and real assets - into one. The identified major risks shape a discussion for the continued surveillance and the sustainability assessment of insurance sectors affected by the Covid-19 crisis.

Adverse decompositions from the Covid-19 crisis, such as slower-than-expected economic recovery or the evolution of a double-dip recession, are reflected in adverse market scenarios developed with ESRB, which are tested as part of the 2021 Insurance Stress Test, launched in May of that year. The final narratives about the crisis and shocks will be developed in collaboration with the European Systemic Board (ESRB) considering recent macroeconomic and market developments. The ESRB-wide discussion of scenarios suggests that both EIOPA and the EBA stress tests will maintain the same narrative for 2021. This will enable the assessment of both sectors together and shed light on possible wider implications for Europe's financial sectors. Furthermore, considering recent experiences of Covid-19 outbreak, the EIOPA has initiated an exercise in liquidity surveillance, building on new data filings recommended by ESRB as well. In this vein, liquidity risks were considered for inclusion in the Insurance Stress Test 2021 (EIOPA, 2020).

Shevchuk, Kondrat, & Stanienda (2020) have interestingly argued that a large-scale disruptive event like a global pandemic may not only provide challenges and hardship, but also an opportunity to innovate. Specifically, companies need to transform themselves digitally and provide their services online, as Covid-19 has caused a general shift towards work-from-home (WFH) and communication via the internet. Digitalization will also allow insurance companies to automate a significant part of their operations, from streamlining management, customer relationship management (CRM), to sales funnels, products, policies, and claims (Weingarth *et al.*, 2019). Importantly, artificial intelligence and machine learning can be used as well, especially when it comes to estimating the best pricing and premium values depending on economic conditions, risk assessment, and the structure of the risk pool (Asimit, Kyriakou, & Nielsen, 2020). Another important change is that more, intermediaries are cut out of existing business structures, resulting in informed customers that directly interact with companies. If that is the case for the insurance industry as well, it would remove some of the cost exposure related to acquisition that these companies have had in the past. Given the growing influence of machine learning within the insurance industries, companies may increase the number of services that offer risk management to their customers, rather than

the simple fulfilment of claims. Of course, due to the impact of the pandemic, companies will likely be more eager in the future to demand additional insurance products, for example, business continuity insurance and related risk management, cover for the cancellation of events, protection for the temporary loss of workforce, and so on (Shevchuk, Kondrat, & Stanienda, 2020).

Sugimoto & Windsor (2020) point out that based on the experience from previous pandemics, health insurers are the companies that deal most with the increased volume of claims, whereas other insurers are less likely to be significantly affected. However, event insurance, travel insurance, and business interruption insurance will also incur an increase in claims, whereas motor insurance may see less claims, as people drive less with their cars during a pandemic (Sugimoto & Windsor, 2020). Importantly, many insurers have taken care that they stay within the confines of insurance regulation during the pandemic. For example, insurance capital must be kept above certain thresholds; to fulfill that requirement, policies must incorporate a certain level of flexibility regarding both the payment of premiums and claims. Importantly, as the EIOPA advocates for a more flexible approach to the *Solvency II* framework considering the Covid-19 pandemic, an opportunity could open regarding the integration of a pandemic or other events with negative economic outlook into the framework of a unified SEIM (Sugimoto & Windsor, 2020).

The above-mentioned decreases in solvency ratios and returns on assets are likely due to devaluations in the capital markets, specifically pertaining to bonds. In the European Union, around 40% of long-term bond investments come from insurance companies. Unfortunately, stock prices of such companies have declined more strongly than the overall value of the equity market, suggesting dampened expectations by investors (Sugimoto & Windsor, 2020).

3. Objectives and hypotheses

The central objective of the present study is to propose a new methodology for exploring whether insurance industry in Europe is becoming more homogeneous or more heterogeneous while, in parallel, for questioning its convergence process. The main rationale behind pursuing this objective is to show the current level of realization of SEIM and how far European integration has brought national insurance markets, when analyzing insurance data with a method from the functional clustering theory. If indeed, the SEIM is not yet realized, then insurance industry should focus on identifying the major challenges and obstacles remaining before SEIM can function as

planned. Therefore, this study needs to continue with an analysis of the development and similarity of the European insurance markets. On this account, future research work will incorporate the results of the proposed methodology into this study with the aim to provide the required analysis.

Further working steps are needed to embed and integrate this study that would lead to future work making a contribution to the literature: (1) we should conduct the development-similarity analysis on a country-wide level between current EU and non-EU member states as well as between old and new EU member states over a long span of time; (2) we should combine the different strands of the literature on insurance economics, insurance policy making and functional clustering theory; (3) we should provide the recent and most updated evidence on the comparison of the insurance markets in Europe and their dynamics in-depth; (4) we should employ a new methodology utilizing a functional clustering (FC) method, which contrary to the standard clustering methods, it offers the advantage of capturing the variability mechanisms of data by taking into account the whole curve values from the past.

Based on the objective, the available literature and the methodology, this study intends to form the basis to test the following hypotheses:

- Hypothesis I: European insurance industry is homogenous within EU.
- Hypothesis II: European insurance industry is inhomogeneous outside EU.
- Hypothesis III: European insurance industry is converging within EU.
- Hypothesis IV: European insurance industry is diverging outside EU.
- Hypothesis V: European insurance industry is diverging from non-EU markets.
- Hypothesis VI: The new EU insurance markets are homogeneous and converging towards old EU insurance markets.
- Hypothesis VII: Structural reforms associated with the EU integration and enlargement process are the only determinants of national insurance development position within the European insurance market.
- Hypothesis VIII: The Covid-19 pandemic has influenced the change in insurance development across European countries.

The results from these tests will provide us important insight into the economic integration process towards SEIM in the EU, and they will also inform us to which degree close economies are already converging by default, and how much central control over the process needs to be exerted. Thus, this study will serve as a tool to better allocate resources and make the insurance integration process within the EU more efficient.

4. Methodology

4.1 Data

We will use non-missing time series data at country level taken at fixed year-end points from Swiss Re Sigma database (2021). This cluster analysis concentrates on combining several insurance market indicators (IMI) of European countries during the period between 2004 and 2021, such as insurance penetration, insurance density, Gross Written premiums, to name few. The data available consisted of 18 time series percentage-valued samples of 34 European countries. We will follow clustering procedure for European insurance market segmentation. after converting the time series data on combined IMI into functional data.

The proposed FC method will be based on both the magnitude and shape of the IMI curves. Particularly, the shape information within data will be uncovered by initial data transformation. In line with Dai, Mrkvicka, Sun and Genton (2020), we will rescale each raw IMI curve to have a zero mean by subtracting its mean value from the curve itself, resulting in the centered IMI curves. This operation is going to mitigate the widely different magnitudes of the IMI curves. By shifting each curve towards its center and normalizing the centered curves by their L_2 norms, the overall shape of each IMI curve will be better identified without the complication of comparing IMI curves with different magnitudes. After this, the clustering will be performed in terms of both the magnitude - raw IMI curves, and shape – normalized centered IMI curves. Then the results will be open for their graphical presentation and analytical interpretation.

4.2 Evaluation of other methodologies

As demonstrated in Athanasiadis and Mrkvička (2019) the clustering methods based on specific functional data techniques were unable to produce meaningful clusters of countries that could reveal, at the same time, both the magnitude and the shape of specific curves, such as IP

curves. They also found that the proposed clustering methods suffer from certain difficulties. For distance-based time series clustering, the difficulty is caused by the uncertainty of feature-based representations of original data (time series) such as autocorrelations. For distance-based functional clustering, the complex shape mechanisms underlying the data may not be well understood by the developed distance measure. For filtering methods, the FPCA approach may form clusters that do not share the same number of common PCs. For adaptive methods, the lack of a definition for the distribution of a curve may also arise some difficulties.

4.3 Methodological framework

The methodology of this study is chosen based on the objectives and hypotheses of the undertaken research, which is to analyze a wide sample of national insurance markets. When including on top the development of IMI over time, this leads to an infinite-dimensional input space. In this study the input space consists of functions or curves, so that each curve represents data for a selected country during the whole selected period. The chosen method needs to be capable of dealing with infinite-dimensional data, which in our case is also a balanced panel. Functional data analysis satisfies this set conditions.

Functional data differ in several ways, making them difficult to analyze when clusters exist from multiple perspectives. The existing methods either focus on a single type of variation or pool the various sources of variation with weightings that rely on a delicate selection procedure. The proposed method combines the different sources of variation in an equal manner and can be used for both univariate and multivariate cases. The study defines the proposed procedure with an arbitrary functional ordering, reviews several functional orderings already defined in the literature (Dai, Athanasiadis, & Mrkvička, 2021).

This study analyzes the integration of Europe's insurance markets between 2004 and 2021 by studying the time-dependent development of IMI. We distinguish between the level of market integration – the degree to which the single European market is attained, or the developments of IMI across selected countries present a homogeneous pattern over time — and market convergence — the degree to which the development of IMI converges to a unique path. To capture both market integration and convergence, our line of argument consists of two steps. First, we demonstrate by

means of a joined functional cluster analysis that utilizes an ordering (combined functional ordering) of the IMI curves, whether curves are further apart or closer together forming a single cluster of curves. Second, using the computed central region of the cluster solution along with its graphical representation, we will be able to show if curves have moved towards a common steady state over the years.

In this study, cluster analysis involves developing a clustering method, which is sensitive, at the same time, to both magnitude and shape heterogeneities of functional data. To begin with, cluster analysis is an essential part of exploratory data analysis that aims to identify homogeneous subgroups in data. It is widely used in functional data analysis for tasks such as classifying electrocardiogram curves for diagnosing cardiovascular ischemic diseases and extracting representative wind behavior. This study focuses on distance-based methods and proposes a new family of clustering algorithms based on the chosen functional ordering. The dissimilarity matrix is constructed via the chosen functional ordering, which is applied to the set of differences of all pairs of functional data under investigation. This study concentrates on orderings with intrinsic graphical interpretation but any ordering that treats the sources equally can be used. The proposed method provides a reasonable global envelope that has an intrinsic graphical interpretation of the clustering results and inherits the robustness of functional orderings (Dai, Athanasiadis, & Mrkvička, 2021).

4.3.1 Dissimilarity Matrix

A dissimilarity matrix is a square matrix that quantifies the dissimilarity or distance between pairs of objects in a dataset. Each element in the matrix represents the dissimilarity between two objects. In the context of functional data analysis, the dissimilarity matrix captures the dissimilarity between pairs of functions or curves based on some distance measure, such as the Euclidean distance or a specific functional distance metric (Kumar and Bezdek, 2020).

4.3.2 Combined Functional Ordering

Combined functional ordering is a method that allows combining several different views of a function so that these views are equally represented in the distance construction. Instead of using a single distance, this specific method combines several different distances (e.g., volatility, trend, etc.) into one composite distance in which all components are represented equally (Hlásková &

Mrkvička, 2022). In our case, this approach aims to capture and address various aspects of the IMI curves simultaneously, such as their magnitude and shape, and combine them in equal manner to provide a more complete and comprehensive ordering scheme of the data.

Thus, assume we observe functions $\mathbf{T}_i(r), i = 1, 2, \dots, s$ in fixed set of points $r_1, r_2, \dots, r_d$. Hence, various transformations could be applied to the raw functions to obtain the transformed data sets of interest, such as $V_1, V_2, \dots, V_q$. We denote with $V_q(T_{ij})$ the resultant curves of $T_{ij}, i = 1, 2, \dots, s; j = 1, 2, \dots, d$ via the transformation V_q , the dissimilarity matrix is then constructed by applying to the long vectors

$$\mathbf{T}_i = (V_1(T_{i1}), \dots, V_1(T_{id}), \dots, V_q(T_{i1}), \dots, V_q(T_{id})), i = 1, 2, \dots, s$$

the corresponding orderings. The desired perspectives of the ordering are considered and treated equally in such a combined ordering due to the ranking nature of proposed ordering.

4.3.3 Global envelopes and Functional Ordering with Intrinsic Graphical Interpretation

The global envelope is constructed for a set of functions or d-dimensional vectors, $\mathbf{T}_i, i = 1, 2, \dots, s$. During past years, different global envelopes have been proposed (see Table 1 for references), which are based on a measure M that orders the functions from the most extreme to the least extreme one. Let $M_1, M_2, \dots, M_s$ be the values of the measures for $\mathbf{T}_1, \mathbf{T}_2, \dots, \mathbf{T}_s$. For the measure M , a critical value m_a can be determined as the largest of M_i such that the number of those i for which $M_i < m_a$ is less or equal to as , where the level $a \in (0, 1)$ is user-defined. If the vector $\mathbf{T}_i$ is among the as most extreme vectors, i.e., $M_i < m_a$, then it is regarded as extreme by the given measure M_i at level a (Myllymäki & Mrkvicka, 2020).

Table 1 lists five different measures along with their detailed references that are implemented in the R library GET. The extreme rank length (erl), continuous rank (cont) and area rank (area) measures are all refinements of the extreme rank measure (rank), which is simply defined as the minimum of pointwise ranks of $\mathbf{T}_i$ among each other. Since many $\mathbf{T}_i$ can reach the same minimal rank, the given ordering of $\mathbf{T}_i$ is only weak. The erl, cont, and area break the ties in the extreme

ranks and practically lead to more strict ordering (Myllymäki & Mrkvicka, 2020). For completeness, we included into this study the Studentized measure (student), which is based on the approximation of the distributions of $\mathbf{T}_i$ from a few characteristics.

For all the measures of Table 1, a global envelope can be constructed such that it has the intrinsic graphical interpretation (IGI) (Myllymäki & Mrkvicka 2023): the vector $\mathbf{T}_i$ is outside the global envelope if and only if $M_i < m_a$, and $\mathbf{T}_i$ is within the global envelope if and only if $M_i \geq m_a$. This holds for all $i = 1, 2, \dots, s$. The choice of the measure with IGI (see Table 1) depends on several aspects: 1) the number of vectors which are available, 2) the dimensionality of the vectors, 3) the amount of the dependence between the vector elements, and 4) the type of extremeness which is expected (Myllymäki & Mrkvicka, 2020).

Table 1 Different measures with global envelopes, their abbreviations and references

Measure	Abbreviation	Referenced in
Extreme rank	rank	Myllymäki et al. (2017)
Extreme rank length	erl	Myllymäki et al. (2017); Narisetty and Nair (2016); Mrkvicka et al. (2020)
Continuous rank	cont	Hahn (2015); Mrkvicka et al. (2019)
Area rank	area	Mrkvicka et al. (2019)
Studentized rank	student	Myllymäki & Mrkvicka (2023)

Source: Own creation based on Myllymäki and Mrkvicka (2020)

To sum up, a functional ordering with IGI is the ordering for which exists a global envelope with IGI with respect to this ordering. In other words, the functional ordering with IGI is a technique that combines functional ordering with visualization methods to interpret the ordering results. It provides graphical representations that enables us to understand the ordering structure of the functions. Several ordering measures fall under this category of orderings with IGI as shown in Table 1. For further information on the definitions of these orderings can be found in Dai, Athanasiadis, & Mrkvicka (2021). For the sake of completeness, we provide here a brief list of these definitions.

Extreme rank length

Let $r_{1j}, r_{2j}, \dots, r_{sj}$, $j = 1, 2, \dots, d$ be the raw ranks of $T_{1j}, T_{2j}, \dots, T_{sj}$ such that the smallest T_{ij} has rank 1. In the case of ties, the raw ranks are averaged. The two-sided pointwise ranks are then calculated as $R_{ij} = \min(r_{ij}, s + 1 - r_{ij})$. Consider now the vectors of pointwise ordered ranks, where $\mathbf{R}_i = (R_{i[1]}, R_{i[2]}, \dots, R_{i[d]})$ where $\{R_{i[1]}, R_{i[2]}, \dots, R_{i[d]}\} = \{R_{i1}, R_{i2}, \dots, R_{id}\}$ and $R_{i[k]} \leq R_{i[k']}$ whenever $k \leq k'$. The extreme rank length measure of the vectors $\mathbf{R}_i$ is equal to:

$$E_i = \frac{1}{s} \sum_{i'}^s (\mathbf{R}_{i'} < \mathbf{R}_i) \quad (1)$$

where

$$\mathbf{R}_{i'} < \mathbf{R}_i \Leftrightarrow \exists n \leq d : R_{i'[k]} = R_{i[k]} \forall k < n, R_{i'[n]} < R_{i[n]}$$

The division by s leads to normalized ranks that obtain values between 0 and 1. Consequently, the erl measure corresponds to the extremal depth as defined in Narisetty and Nair (2016).

Let e_a be a critical value and let $I_a = \{i \in 1, 2, \dots, s : E_i \geq e_{(a)}\}$ be the index set of vectors less extreme than or as extreme as e_a . Then, the $100(1 - a)\%$ global extreme rank length envelope (or global extreme rank length central region) induced by E_i is:

$$\mathbf{T}_{low\ k}^{(a)} = \min_{i \in I_a} T_{ik} \text{ and } \mathbf{T}_{upp\ k}^{(a)} = \max_{i \in I_a} T_{ik} \text{ for } k = 1, 2, \dots, d \quad (2)$$

Global continuous rank ordering

The continuous rank measure is:

$$C_i = \min_{j=1,2,\dots,d} c_{ij} / \lceil s/2 \rceil,$$

where c_{ij} are the pointwise continuous ranks defined as:

$$\begin{aligned}
 c_{ij} &= \sum_{i'} \mathbf{1}(T_{i'j} > T_{ij}) + \frac{T_{[i+1]j} - T_{ij}}{T_{[i+1]j} - T_{[i-1]j}} && \text{for } i : T_{ij} \neq \max_{i'} T_{i'j} \\
 &&& \text{and } T_{ij} > \text{median}(T_{ij}) \\
 c_{ij} &= \exp\left(-\frac{T_{ij} - T_{[i-1]j}}{T_{[i-1]j} - \min_i T_{ij}}\right) && \text{for } i : T_{ij} = \max_{i'} T_{i'j} \\
 c_{ij} &= \sum_{i'} \mathbf{1}(T_{i'j} < T_{ij}) + \frac{T_{ij} - T_{[i-1]j}}{T_{[i+1]j} - T_{[i-1]j}} && \text{for } i : T_{ij} \neq \min_{i'} T_{i'j} \\
 &&& \text{and } T_{ij} < \text{median}(T_{ij}) \\
 c_{ij} &= \exp\left(-\frac{T_{[i+1]j} - T_{ij}}{\max_i T_{ij} - T_{[i+1]j}}\right) && \text{for } i : T_{ij} = \min_{i'} T_{i'j} \\
 c_{ij} &= R_{ij} && \text{for } T_{ij} = \text{median}(T_{ij})
 \end{aligned}$$

Here, $T_{[i-1]j}$ and $T_{[i+1]j}$ denote the values of the functions, which are in a j -th element below and above T_{ij} , respectively (i.e., $T_{[i-1]j} = \max_{i': T_{i'j} < T_{ij}} T_{i'j}$ and $T_{[i+1]j} = \min_{i': T_{i'j} > T_{ij}} T_{i'j}$).

The 100 $(1 - a)\%$ global continuous rank envelope induced by C_i is constructed in the same manner as the global extreme rank length envelope.

Global area rank ordering

The area rank measure:

$$A_i = \frac{1}{[s/2]d} \sum_j \min(R_i, c_{ij})$$

where $R_i = \min_j \{R_{ij}\}$ and R_{ij} are two-sided pointwise ranks defined above. The 100 $(1 - a)\%$ global area rank envelope induced by A_i is constructed in a manner similar to that of the global extreme rank length envelope.

Studentized maximum ordering

Because we construct a symmetric set of functions to compute the dissimilarity matrix, here we use only the symmetric studentized ordering. The above orderings are based on the whole distributions of $T_j, j = 1, 2, \dots, d$. It is also possible to approximate the distribution from a few sample characteristics. The studentized maximum ordering approximates the distribution of $T_j, j = 1, 2, \dots, d$ by the sample mean T_{0j} and sample standard deviation $sd(T_j)$. The studentized measure is:

$$S_i = \max_j \left| \frac{T_{ij} - T_{0j}}{sd(T_j)} \right| \quad (3)$$

The $100(1 - \alpha)\%$ global studentized envelope induced by S_i is defined by:

$$\mathbf{T}_{low\ j}^{(l)} = T_{0j} - s_\alpha sd(T_j) \text{ and } \mathbf{T}_{upp\ j}^{(l)} = T_{0j} + s_\alpha sd(T_j) \text{ for } j = 1, 2, \dots, d \quad (4)$$

where s_α is a critical value.

4.3.4 Dissimilarity Matrix Based on the Combined Ordering

A dissimilarity matrix based on the combined ordering incorporates the combined functional ordering method to compute the dissimilarity between functions. The dissimilarity matrix designed through this functional ordering method is then used as input in clustering algorithms or other analysis techniques to group similar functions together based on their combined ordering characteristics. These concepts are often employed in functional data analysis to order and compare functions, facilitating the understanding and interpretation of complex datasets represented as curves or time series (Dai, Athanasiadis, & Mrkvička, 2021).

Assume that we observe functions assume $\mathbf{T}_i(r), i = 1, 2, \dots, s$ in fixed set of points $r_1, r_2, \dots, r_d$. We construct the new set of functional differences

$$D_f = \{\mathbf{T}_i(r) - \mathbf{T}_{i'}(r)\}, i, i' = 1, 2, \dots, s$$

Apply a chosen IGI ordering to D_f . Finally, the elements of the dissimilarity matrix $d_{ii'} = 1 - M_{ii'}$, where $M_{ii'}$ is measure, which induce the chosen ordering, of $\mathbf{T}_i(r) - \mathbf{T}_{i'}(r)$.

4.3.5 Joint functional clustering for national insurance markets of Europe

Our proposed functional clustering method then consists of the following steps:

- Choose the appropriate data sources (e.g., raw data, sequentially transformed data i.e., normalized centered IMI curves)
- Choose the functional ordering which allows for IGI and which gives the same weight to every chosen source (e.g., the Studentized maximum ordering, the area rank ordering).
- Compute the dissimilarity matrix from the set of differences.
- Apply partitioning around medoids (PAM) found in Kaufman and Rousseeuw (1990) using the dissimilarity matrix of the previous step.
- Plot the resulted clusters together with their central region with IGI.

5. Conclusions

This study aims to provide a methodology to explore the homogeneity and convergence of the European insurance market within SEIM. We start by exploring the establishment of the SEIM in 1994 and initiatives like the pan-European Personal Pension Product (PEPP) as steps toward unifying the European insurance market. Next, the regulatory frameworks of Solvency I and Solvency II are introduced that ensure solvency and liquidity requirements are met by insurance companies. The ongoing review of Solvency II by EIOPA suggests that the existing framework is sufficient for implementing the SEIM. However, before the SEIM implementation, the homogeneity and convergence towards a single European market for insurance via financial integration should be examined and proved through an empirical application. To this end, a comprehensive functional clustering method is introduced to analyze complex and time-dependent insurance data created by a composite IMI that aggregates insurance activity indicators such as penetration and density, while it becomes our proxy for insurance industry development. At this stage, the focus should be on analyzing economic, regulatory, and policy interventions to achieve homogenization or proposing measures to address possible observed heterogeneity in the market. In that case, further quantitative research will be needed to investigate whether there are other circumstances or factors affecting

the development of the insurance industry in each homogeneous group of European countries, where the GDP per capita ratio will be approached as an economic growth measure.

Starting with the enlargement of the EU in 2004 and 2007, the impact of adding new countries with different economic development levels on EU insurance industry should be examined. Emphasis will be placed on the need for rules and regulations to achieve balance and homogeneity across new and old EU countries' insurance markets. The effects on welfare and concerns about corruption in new EU member states will also be discussed.

The impact of the Covid-19 pandemic on the European insurance industry should be addressed, considering disruptions to the economy and increased claims volume. The role of the European Commission and EIOPA in assessing crisis events such as the pandemic and conducting stress tests will be examined. Then, their potential for regulatory involvement and the call for further unification of the European insurance market will also be discussed.

6. References

American Academy of Actuaries, 2016. *Principle-based reserving: A new way to insure for life: American Academy of Actuaries* [online], Available from: <https://www.actuary.org/node/13473>

ASIMIT, V., I. KYRIAKOU, and J.P. NIELSEN., 2020. Special issue “machine learning in insurance” *Risks* [online]. 2020, vol. 8, no. 2, p. 54. Retrieved from: doi:10.3390/risks8020054

ATHANASIADIS, S and MRKVIČKA, T, 2018, European Insurance Market Analysis: A Multivariate Clustering approach. *Proceedings of the 12th International Scientific Conference INPROFORUM* [online]. 2018. Available from: <http://ocs.ef.jcu.cz/index.php/inproforum/INP2018>

ATHANASIADIS, S and MRKVIČKA, T, 2019, European insurance market analysis via functional data clustering techniques. *37th International Conference on Mathematical Methods in Economics* [online]. 2019. Available from: <https://mme2019.ef.jcu.cz/>

Atlas Magazine, 2021. *Solvency I and II* [online], Available from: <https://www.atlas-mag.net/en/article/solvency-i-and-ii>

AUSSILLOUX, V., A. BÉNASSY-QUÉRÉ, C. FUEST, et al., 2017. Making the best of the European Single Market. Bruegel Policy Contribution Issue No. 3: 2017. *AEI Banner* [online]. 1. February 2017. Retrieved from: <http://aei.pitt.edu/83988/>

BABUNA, P., X. YANG, A. GYILBAG, et al., 2020. The impact of COVID-19 on the insurance industry. *International Journal of Environmental Research and Public Health* [online]. 2020, vol. 17, no. 16, p. 5766. Retrieved from: doi:10.3390/ijerph17165766

BECMER, Anna, 2015, The Insurance Market of Selected Countries in Central and Eastern Europe Against a Background of the EU Countries-15. In: Some Current Issues in Economics. p. 158–169.

BUSEMEYER, M.R. and T. IVERSEN, 2020. The Welfare State with private alternatives: The transformation of popular support for Social Insurance. *The Journal of Politics* [online]. 2020, vol. 82, no. 2, pp. 671–686. Retrieved from: doi:10.1086/706980

CALIENDO, L., L.D. OPROMOLLA, F. PARRO, et al., 2021. Goods and factor market integration: A quantitative assessment of the EU enlargement. *Journal of Political Economy* [online]. 2021, vol. 129, no. 12, pp. 3491–3545. Retrieved from: doi:10.1086/716560

CHEUNG, Zedric, 2020, Factor Analysis & Cluster Analysis on countries classification. *LaptrinhX* [online]. 2020. Available from: <https://laptrinhx.com/factor-analysis-cluster-analysis-on-countries-classification-2738054554/>

CLARK, Matthew, 2009, Uncertainties, Challenges and Opportunities of Global Insurance Regulatory Convergence. *Risk Management*. 2009. No. 15, p. 49–51.

DAI, Wenlin, ATHANASIADIS, Stavros and MRKVIČKA, Tomáš, 2021, A new functional clustering method with combined dissimilarity sources and graphical interpretation. *Computational Statistics and Applications*. 2021. DOI 10.5772/intechopen.100124.

DAI, Wenlin, MRKVIČKA, Tomáš, SUN, Ying and GENTON, Marc G., 2020, Functional outlier detection and taxonomy by Sequential Transformations. *Computational Statistics and Data Analysis*. 2020. Vol. 149, p. 106960. DOI 10.1016/j.csda.2020.106960.

DENKOWSKA, A. and S. WANAT, 2020. Development and similarity of insurance markets of European Union countries after the enlargement in 2004. *arXiv.org* [online]. 2020. Retrieved from: <https://arxiv.org/abs/2012.15078>

DREW, Catherine and SRISKANDARAJAH, Dhananjayan, 2007, EU enlargement in 2007: No warm welcome for Labor migrants. *migrationpolicy.org* [online]. 2007. Available from: <https://www.migrationpolicy.org/article/eu-enlargement-2007-no-warm-welcome-labor-migrants/>

DOFF, René, 2016, The final solvency II framework: Will it be effective? *The Geneva Papers on Risk and Insurance - Issues and Practice*. 2016. Vol. 41, no. 4, p. 587–607. DOI 10.1057/gpp.2016.4.

EIOPA, 2020. *Annual report 2020* [online], Available from: https://www.eiopa.europa.eu/publications/annual-report-2020_en

EIOPA, 2022. *Pan-European personal pension product (PEPP)* [online], Available from: https://www.eiopa.europa.eu/browse/regulation-and-policy/pan-european-personal-pension-product-pepp_en#:~:text=The%20pan%2DEuropean%20Personal%20Pension,to%20existing%20national%20pension%20regimes.

ENNSFELLNER, Karl C. and DORFMAN, Mark S., 1998, The transition to a single insurance market in the European Union. *Risk Management and Insurance Review*. 1998. Vol. 1, no. 2, p. 35–53. DOI 10.1111/j.1540-6296.1998.tb00071.x.

European Commission, 2014. *European Financial Stability and Integration Review (EFSIR)* [online], Available from: https://ec.europa.eu/info/sites/default/files/efsir-2013-28042014_en.pdf

European Commission, 2021. *Reviewing EU insurance rules: Encouraging insurers to invest in Europe's future* [online], Available from: https://ec.europa.eu/commission/presscorner/detail/en/ip_21_4783

European Commission, 2015. *Solvency II overview — frequently asked questions* [online], Available from: https://ec.europa.eu/commission/presscorner/detail/en/MEMO_15_3120

FRANZETTI, Claudio, 2021, Insurance background. *Management for Professionals*. 2021. P. 73–85. DOI 10.1007/978-3-030-70285-4_3.

FOURNIER, J.-M., A. DOMPS, Y. GORIN, et al., 2015. Implicit regulatory barriers in the EU Single Market. *OECD iLibrary* [online]. 2015. Retrieved from: https://www.oecd-ilibrary.org/economics/implicit-regulatory-barriers-in-the-eu-single-market_5js7xj0xckf6-en

FÁVERO, L.P. and P. BELFIORE, 2019. Cluster analysis. *Data Science for Business and Decision Making* [online]. 2019, pp. 311–382. Retrieved from: doi:10.1016/b978-0-12-811216-8.00011-2

GAGANIS, C., I. HASAN, P. PAPADIMITRI, et al., 2019. National culture and risk-taking: Evidence from the insurance industry. *Journal of Business Research* [online]. 2019, vol. 97, pp. 104–116. Retrieved from: doi:10.1016/j.jbusres.2018.12.037

GROUT, Paul A, 2021, Ai, ML, and competition dynamics in Financial Markets. *Oxford Review of Economic Policy*. 2021. Vol. 37, no. 3, p. 618–635. DOI 10.1093/oxrep/grab014.

HAHN, Ute, 2015, A note on simultaneous Monte Carlo tests. *Technical report, Centre for Stochastic Geometry and advanced Bioimaging, Aarhus University* [online]. 2015. Available from: <https://math.au.dk/en/forskning/publikationer/instituttets-serier/arkiv/csgb/publication/publication/1042?cHash=21999b2712c9bf1f0bc928952cf35e1c>

HAN, Liyan, LI, Donghui, MOSHIRIAN, Fariborz and TIAN, Yanhui, 2010, Insurance development and economic growth. *The Geneva Papers on Risk and Insurance - Issues and Practice*. 2010. Vol. 35, no. 2, p. 183–199. DOI 10.1057/gpp.2010.4.

HEIDBREder, E.G., 2011. *The impact of expansion on European Union institutions the eastern touch on Brussels*. New York: Palgrave Macmillan, 2011.

HENNIG, Christian and LIAO, Tim F., 2013, How to find an appropriate clustering for mixed-type variables with application to socio-economic stratification. *Journal of the Royal Statistical Society Series C: Applied Statistics*. 2013. Vol. 62, no. 3, p. 309–369.
DOI 10.1111/j.1467-9876.2012.01066.x.

HILPISCH, Yves, 2020, *Artificial Intelligence in finance: A python-based guide*. Sebastopol (Calif.): O'Reilly Media.

HLÁSKOVÁ, Gabriela and MRKVIČKA, Tomáš, 2022, Functional cluster regression for commodities and the representatives of stock indexes. *DIGITALIZATION. Society and Markets, Business and Public Administration*. 2022. DOI 10.32725/978-80-7394-976-1.33.

IN 'T VELD, J., 2019. The economic benefits of the EU Single Market in goods and services. *Journal of Policy Modeling* [online]. 2019, vol. 41, no. 5, pp. 803–818. Retrieved from: doi:10.1016/j.jpolmod.2019.06.004

JARROW, Robert A., 2021. The economics of insurance: A derivatives-based approach. *Annual Review of Financial Economics* [online]. 2021, vol. 13, no. 1, pp. 79–110. Retrieved from: doi:10.1146/annurev-financial-040721-075128

JAGRIC, T., S. BOJNEC, and V. JAGRIC, 2018. A map of the European Insurance Sector – are there any borders? *ECONOMIC COMPUTATION AND ECONOMIC CYBERNETICS STUDIES AND RESEARCH* [online]. 2018, vol. 52, no. 2/2018, pp. 283–298. Retrieved from: doi:10.24818/18423264/52.2.18.17

JAGRIC, T. and M. ZUNKO, 2013. Neural network world: Optimized spiral spherical SOM (OSS-SOM). *Neural Network World* [online]. 2013, vol. 23, no. 5, pp. 411–426. Retrieved from: doi:10.14311/nnw.2013.23.025

KAUFMAN, Leonard and ROUSSEEUW, Peter J., 1990, Finding groups in Data. *Wiley Series in Probability and Statistics*. 1990. DOI 10.1002/9780470316801.

KIRMAN, A., 2006. Heterogeneity in Economics. *Journal of Economic Interaction and Coordination* [online]. 2006, vol. 1, no. 1, pp. 89–117. Retrieved from: doi:10.1007/s11403-006-0005-8

KUMAR, Dheeraj and BEZDEK, James C., 2020, Visual approaches for exploratory data analysis: A survey of the visual assessment of Clustering Tendency (VAT) family of algorithms. *IEEE Systems, Man, and Cybernetics Magazine*. 2020. Vol. 6, no. 2, p. 10–48. DOI 10.1109/msmc.2019.2961163.

KWON, W. Jean and WOLFROM, Leigh, 2016, Analytical tools for the insurance market and macro-prudential surveillance. *OECD Journal: Financial Market Trends*. 2016. Vol. 2016, no. 1, p. 1–47. DOI 10.1787/fmt-2016-5jln6hmvwdzn.

KÖHNE, Thomas and BRÖMMELMEYER, Christoph, 2018, The new insurance distribution regulation in the EU—a critical assessment from a legal and Economic Perspective. *The Geneva Papers on Risk and Insurance - Issues and Practice*. 2018. Vol. 43, no. 4, p. 704–739. DOI 10.1057/s41288-018-0089-0.

LIN, Chuyuan, 2019, *Unsupervised Learning on Functional Data with an Application to the Analysis of U.S. Temperature Prediction Accuracy*. thesis. Greater Vancouver, British Columbia, Canada: Simon Fraser University.

LIU, R., C. ZHANG, and T. FENG, 2021. Fusion analysis of economic data of the medical and health industry based on blockchain technology and two-way spectral cluster analysis. *Mobile Information Systems* [online]. 2021, vol. 2021, pp. 1–12. Retrieved from:
doi:10.1155/2021/7731387

LÉGER, Ainhoa-Elena and MAZZUCO, Stefano, 2021, What can we learn from the functional clustering of mortality data? an application to the Human Mortality Database. *European Journal of Population*. 2021. Vol. 37, no. 4–5, p. 769–798. DOI 10.1007/s10680-021-09588-y.

MA, Hongyan, 2021, The role of clustering algorithm-based big data processing in Information Economy Development. *PLOS ONE*. 2021. Vol. 16, no. 3. DOI 10.1371/journal.pone.0246718.

MARDANEH, Karim K, 2015, Functional specialisation and socio-economic factors in population change: A clustering study in non-metropolitan Australia. *Urban Studies*. 2015. Vol. 53, no. 8, p. 1591–1616. DOI 10.1177/0042098015577340.

MARTÍ, Pablo, SERRANO-ESTRADA, Leticia, NOLASCO-CIRUGEDA, Almudena and BAEZA, Jesús López, 2021, Revisiting the spatial definition of neighborhood boundaries: Functional clusters versus administrative neighborhoods. *Journal of Urban Technology*. 2021. Vol. 29, no. 3, p. 73–94. DOI 10.1080/10630732.2021.1930837.

MCGEE, A., 2020. *Single market in insurance: Breaking down the barriers*. S.l.: ROUTLEDGE, 2020.

MÜLLER-REICHART, M., 2005. The EU insurance industry: Are we heading for an ideal single financial services market? *The Geneva Papers on Risk and Insurance - Issues and*

Practice [online]. 2005, vol. 30, no. 2, pp. 285–295. Retrieved from: doi:10.1057/pal-grave.gpp.2510027

NARISSETTY, Naveen N. and NAIR, Vijayan N., 2016, Extremal depth for functional data and applications. *Journal of the American Statistical Association*. 2016. Vol. 111, no. 516, p. 1705–1714. DOI 10.1080/01621459.2015.1110033.

MRKVIČKA, Tomáš, MYLLYMÄKI, Mari, JÍLEK, Milan and HAHN, Ute, 2020, A one-way ANOVA test for functional data with graphical interpretation. *Kybernetika*. 2020. P. 432–458. DOI 10.14736/kyb-2020-3-0432.

MRKVICKA, Tomas, MYLLYMAKI, Mari, KURONEN, Mikko and NARISSETTY, Naveen Naidu, 2019, New methods for multiple testing in permutation inference for the general linear model. *arXiv.org* [online]. 2019. Available from: <https://arxiv.org/abs/1906.09004>

MYLLYMÄKI, Mari and MRKVIČKA, Tomáš, 2020, Comparison of non-parametric global envelopes. *arXiv.org* [online]. 2020. Available from: <https://arxiv.org/abs/2008.09650v1>

MYLLYMÄKI, Mari and MRKVIČKA, Tomáš, 2023, GET: Global envelopes in R. *arXiv.org* [online]. 2023. Available from: <https://arxiv.org/abs/1911.06583>

MYLLYMÄKI, Mari, MRKVIČKA, Tomáš, GRABARNIK, Pavel, SEIJO, Henri and HAHN, Ute, 2017, Global envelope tests for spatial processes. *Journal of the Royal Statistical Society Series B: Statistical Methodology*. 2017. Vol. 79, no. 2, p. 381–404. DOI 10.1111/rssb.12172.

OSTROWSKA-DANKIEWICZ, Anna and SIMIONESCU, Mihaela, 2020, Relationship between the insurance market and macroeconomic indicators in the EU member states. *Transformations in Business & Economics*. 2020. Vol. 19, no. 3, p. 175–187.

PELECKIENĖ, V., K. PELECKIS, G. DUDZEVIČIŪTĖ, et al., 2019. The relationship between insurance and economic growth: Evidence from the European Union countries. *Economic Research-Ekonomska Istraživanja* [online]. 2019, vol. 32, no. 1, pp. 1138–1151. Retrieved from: doi:10.1080/1331677x.2019.1588765

POLAŃSKI, P.P., 2018. Revisiting country of origin principle: Challenges related to regulating e-commerce in the European Union. *Computer Law & Security Review* [online]. 2018, vol. 34, no. 3, pp. 562–581. Retrieved from: doi:10.1016/j.clsr.2017.11.001

PUŁAWSKA, K., 2021. Financial Stability of European insurance companies during the COVID-19 pandemic. *Journal of Risk and Financial Management* [online]. 2021, vol. 14, no. 6, p. 266. Retrieved from: doi:10.3390/jrfm14060266

SEDELMEIER, U., 2014. Europe after the eastern enlargement of the European Union: 2004-2014: Heinrich Böll stiftung: Brussels Office - European Union. *Heinrich-Böll-Stiftung* [online]. 2014. Retrieved from: <https://eu.boell.org/en/2014/06/10/europe-after-eastern-enlargement-european-union-2004-2014>

SHEVCHUK, O., I. KONDRAT, and J. STANIENDA., 2020. Pandemic as an accelerator of digital transformation in the insurance industry: Evidence from Ukraine. *Insurance Markets and Companies* [online]. 2020, vol. 11, no. 1, pp. 30–41. Retrieved from: doi:10.21511/ins.11(1).2020.04

SINGH, D., 2021. The ECB Guide to internal liquidity adequacy: A principles-based approach. *SSRN Electronic Journal* [online]. 2021. Retrieved from: doi:10.2139/ssrn.3846317

SUGIMOTO, Nobuyasu and WINDSOR, Peter, 2020, Regulatory and Supervisory Response to Deal with Coronavirus Impact: The Insurance Sector. *Special series on covid-19*

[online]. 2020. Available from: [https://www.imf.org/en/Publications/SPROLLs/covid19-special-](https://www.imf.org/en/Publications/SPROLLs/covid19-special-notes)
notes

Swiss Re Sigma database, 2021. *Sigma Explorer - catastrophe and insurance market data*
[online], Available from: <https://www.sigma-explorer.com/>

TOSHKOV, D.D., 2017. The impact of the eastern enlargement on the decision-making
capacity of the European Union. *Journal of European Public Policy* [online]. 2017, vol. 24, no. 2,
pp. 177–196. Retrieved from: doi:10.1080/13501763.2016.1264081

TSEPENDA, Igor and HURAK, İhor, 2021, Perspectives of the EU membership for
Ukraine: The main challenges and threats. *İnsan ve Toplum Bilimleri Araştırmaları Dergisi*
[online]. 2021. Available from: <https://dergipark.org.tr/en/pub/itobiad/issue/60435/880128>

VAGNOZZI, Anna Marie, 2020, Detecting Partisan Gerrymandering through Mathemati-
cal Analysis: A Case Study of South Carolina. *All Theses*. 3347 [online]. 2020. Available from:
https://tigerprints.clemson.edu/all_theses/3347

WARSAME, Mohamed A., 2021, Clustering algorithms for economic policy guidance.
Medium [online]. 2021. [Accessed 2023]. Available from: [https://towardsdatascience.com/clus-](https://towardsdatascience.com/clustering-algorithms-for-economic-policy-guidance-45f469704815)
tering-algorithms-for-economic-policy-guidance-45f469704815

WEINGARTH, Janina, HAGENSCHULTE, Julian, SCHMIDT, Nikolaus and BALSER,
Markus, 2019, Building a digitally enabled future: An insurance industry case study on Digitali-
zation. *Management for Professionals*. 2019. P. 249–269. DOI 10.1007/978-3-319-95273-4_13.

Wilson Center, 2004. *European Union enlargement: Pushing the Frontiers of Coopera-*
tion [online], Available from: [https://www.wilsoncenter.org/article/european-union-enlargement-](https://www.wilsoncenter.org/article/european-union-enlargement-pushing-the-frontiers-cooperation)
pushing-the-frontiers-cooperation

ZHU, Zhangyao and LIU, Na, 2021, Early warning of financial risk based on K-means clustering algorithm. *Complexity*. 2021. Vol. 2021, p. 1–12. DOI 10.1155/2021/5571683.

ZLATEVA, Ioanna, 2022, *O solvency* [online]. 2022. Available from:
<https://www.scribd.com/document/490515795/o-solvency>