

UNIVERSIDAD PONTIFICIA DE COMILLAS
ESCUELA TÉCNICA SUPERIOR DE INGENIERÍA (ICAI)
GRADO EN INGENIERÍA EN TECNOLOGÍAS INDUSTRIALES

TRABAJO DE FIN DE GRADO

**DESARROLLO DE UN KPI SINTÉTICO PARA
MEDIR LA RECUPERACIÓN ESPAÑOLA
FRENTE AL COVID-19**

AUTOR: Gustavo Álvarez Fernández

DIRECTOR: David Roch Dupré

Madrid,
29 de Junio de 2022

Autorizada la entrega del proyecto del alumno:

Gustavo Álvarez Fernández

EL DIRECTOR DEL PROYECTO

David Roch Dupré



Fdo.:

Fecha: 29 / 06 / 2022

V° B° DEL COORDINADOR DE PROYECTOS

Mercedes Fernández García

Fdo.:

Fecha: / /

Resumen

DESARROLLO DE UN KPI SINTÉTICO PARA MEDIR LA RECUPERACIÓN ESPAÑOLA FRENTE AL COVID-19

Autor: Álvarez Fernández, Gustavo.

Director: Roch Dupré, David.

Entidad colaboradora: ICAI - Universidad Pontificia de Comillas

Palabras clave: Indicador sintético, Mercado laboral español, Impacto social, Agregación y ponderación, PCA.

Introducción

El coronavirus, o SARS-CoV-2, ha sido el causante de una de las pandemias más devastadoras de la historia. Tras su irrupción en 2019 en el continente asiático, el virus se extendió rápidamente por todo el planeta. El COVID-19, enfermedad infecciosa propiciada por el virus, enseguida alteró enormemente las vidas de los ciudadanos de naciones tanto desarrolladas como en vías de desarrollo. En este contexto, los países se vieron obligados a tomar medidas políticas extremas tales como la supresión de la movilidad o la interrupción de actividades laborales no esenciales.

La actividad económica y laboral se vio obligada a “hibernar”, de manera que los agentes públicos y privados impulsaron una serie de medidas para mitigar los importantes estragos causados por la pandemia. Pese a ello, las cifras de parados, el coste de los salarios e incluso la productividad se han visto seriamente damnificados. Esta información, en su momento dispersa y desordenada, se encuentra ahora presentada a través del panel interactivo 360 Smart Vision (2021), el cual proporciona una cantidad incommensurable de indicadores que reflejan información precisa y detallada de la recuperación socioeconómica en España. De hecho, la plataforma los ilustra a través de ocho áreas separadas: mercados financieros, movilidad, consumo, sanidad, mercado laboral, ciencia ciudadana, actividad empresarial y sostenibilidad e inclusión.

Pese a ello, si un interesado quiere valorar la situación de algunas de estas áreas, se ve obligado a considerar, en algunos casos, hasta más de 100 indicadores. Por lo tanto, no existe una forma sencilla de evaluar los fenómenos descritos por las áreas. Dado que resulta casi imposible analizar cientos de ellos simultáneamente, los interesados pueden llegar a desechar muchos de ellos por estimarlos carentes de importancia mientras que, en cambio, los considerados relevantes tendrían que ser ponderados sobre la marcha. De esta manera,

tanta información desagregada puede conllevar la toma de decisiones equivocadas a la hora de analizar nuestro entorno. La falta de un indicador sintético que normalice, pondere y agregue dichos indicadores supone un lastre importante cuando se trata de comprender la realidad del mercado laboral español. De esta manera, la elaboración de un indicador sintético a partir de la base de datos de 360 Smart Vision se torna más necesaria que nunca, de forma que los interesados en dicho fenómeno puedan dotarse de un indicador accesible, sencillo y transparente en su elaboración.

Los indicadores sintéticos, en líneas generales, son índices compuestos a partir de indicadores o variables independientes (Freudenberg, 2003), de manera que los agrupan mediante diferentes técnicas en pos de medir un fenómeno multidimensional y complejo. Durante los últimos años, su creación se ha visto potenciada por las bondades que presentan, en tanto que atraen la atención pública sobre problemáticas sociales y dinamizan la toma de decisiones de organismos públicos y privados (Greco *et al.*, 2019). De hecho, poco a poco anteriores detractores de los sintéticos como Amartya Sen, Premio Nobel de Economía en 1998, cambian su postura crítica gracias al gran impacto social que pueden tener los sintéticos, es decir, gracias a su capacidad de transformación social.

Por supuesto, estos detractores siguen existiendo. Como principales argumentos, se destaca la sobreesimplificación de información que acarrearán consigo, sumada a ciertas arbitrariedades presentes en su elaboración (Grupp y Mogee, 2004). Esto puede llevar a los agentes sociales a interpretar erróneamente el sintético y, por tanto, extraer mensajes engañosos del mismo. No obstante, para evitar esos problemas, es fundamental seguir un proceso claro en su creación, tal y como establece la OCDE (2008) y otras entidades públicas (Nardo *et al.*, 2005). Asimismo, la evaluación de resultados debe ser una constante a lo largo del procedimiento, de manera que se identifiquen las fuentes de incertidumbre y sesgo introducidas en el sintético.

En primer lugar, ha de establecerse el marco teórico del indicador, que delimita los conceptos estudiados, determina los subgrupos del sintético y, además, define los criterios de calidad que han de mantenerse a lo largo del proyecto. En segundo lugar, la recolección de datos consiste en la estructuración, depuración y selección de variables de acuerdo con ciertos preceptos como la pertinencia, la oportunidad, la eficiencia o la completitud temporal. Debido a que algunas variables no gozan de la totalidad de los casos, la imputación de información ausente debe ser la tercera etapa del procedimiento. En cuarto lugar, el análisis multivariante arroja conclusiones fundamentales sobre el comportamiento de los datos y las correlaciones existentes entre las variables, de manera que sean consideradas en las posteriores fases de ponderación y agregación.

Una vez comprendido el carácter de los datos, la normalización proporciona una escala común a todas las variables originales. Tras este paso, los datos ya se encuentran preparados para su ponderación y agregación, considerando especialmente la estructura establecida por el marco teórico y la información revelada por el análisis multivariante. De esta manera, la ponderación asigna pesos a cada variable según su importancia relativa, aunque debe ser valorada en conjunción con la técnica de agregación empleada. Al fin y al cabo, algunos métodos permiten la compensación total entre los indicadores componentes, por lo que la ponderación debe tener esto en cuenta, especialmente en caso de que dicho fenómeno no sea adecuado para el sintético. Asimismo, los pesos deben tener en cuenta también la correlación entre variables, ya que se puede incurrir en el problema del doble conteo (Greco *et al.*, 2019). La fase de agregación aúna las diferentes variables en categorías, para después agregar estas

en torno al sintético final. Para ello, destacan las técnicas aritméticas y geométricas. Para concluir, el análisis de robustez cuantifica el impacto de las decisiones tomadas a lo largo de la elaboración, a través de la obtención de 16 resultados distintos generados por decisiones alternativas. Así se pone de manifiesto cuán determinantes son las técnicas empleadas sobre el resultado final, de manera que a mayor variación del índice obtenido menor robustez presenta el sintético. Tras este paso, la divulgación *a posteriori* concluye el proceso de elaboración en cuestión.

Estas etapas se encuentran muy contrastadas por la bibliografía, ya que cada una goza de una aplicación mucho mayor en otros campos de la Estadística. De esta manera, uno de los objetivos del proyecto consistirá en analizar y actualizar los manuales existentes de elaboración a fecha de 2022, de forma que futuros autores dispongan de un complemento renovado y contrastado que permita tomar decisiones adecuadas según el tipo de sintético que se aspire a construir. Por supuesto, el otro objetivo y, si cabe más importante, versa sobre la implementación efectiva del indicador sintético sobre el mercado laboral en España, de manera que se llene el vacío existente al respecto. De esta manera, el proyecto fomentará no solo la transformación social a partir del mencionado indicador, sino también mediante el establecimiento de unos pilares que permitan a futuros autores elaborar sus propios indicadores sintéticos. A tal efecto, la divulgación de los resultados debe también promoverse debidamente.

Metodología

La metodología empleada por el trabajo sigue el proceso establecido por el manual de la OCDE (2008). No obstante, dado que uno de los objetivos consiste en la revisión y actualización del diseño conceptual de sintéticos, el proyecto se divide principalmente en dos capítulos. En primer lugar, se presenta el capítulo 2, donde se diseña conceptualmente un indicador cualquiera, de forma que se muestran críticamente los pormenores de cada técnica aplicable en todas las etapas anteriormente mencionadas. En segundo lugar, el capítulo 3 refleja la implementación efectiva de los hallazgos de su precedente, salvaguardando con sumo cuidado la transparencia y la argumentación de toda decisión tomada en su proceso. Por ello, ambos capítulos han de seguir la estructura de elaboración de un sintético cualquiera planteada anteriormente. Así, tras el análisis de las diferentes técnicas, las decisiones tomadas en cada etapa son las siguientes:

- **Establecimiento del marco teórico:** La *salud del mercado laboral español* representa el fenómeno que describe el sintético construido, comprendida como aquella situación en que los “agentes del mercado ven satisfechas sus necesidades de forma sostenible en el tiempo”. De esta manera, el indicador sintético evalúa no solo los números relativos a la tasa de desempleo, sino también aspectos tales como la productividad, la temporalidad del trabajo, la capacitación de la población, la protección a los desempleados...

Básicamente, cinco categorías componen dicho fenómeno. Se presentan divididas según su carácter: como causas, *Habilidades y capacitación*, *Protección a los desempleados* y *Salarios y costes laborales*; y como consecuencias de la salud del mercado, *Desempleo* y *Empleo*, ya que cada una caracteriza el tipo de paro o de contratos (por ejemplo) presentes en la sociedad. Finalmente, se establecen los criterios de selección de los datos, tales como pertinencia, circunscripción espacial y temporal, accesibilidad y precisión.

El marco teórico del proyecto queda definido de acuerdo con la figura 1:

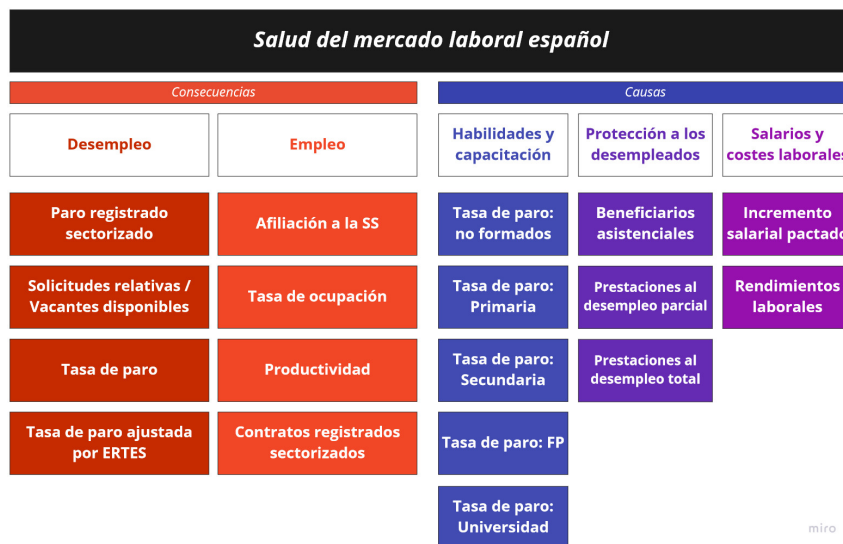


Figura 1. Estructura anidada del sintético construido

- Recolección de datos:** La recolección implica la aplicación de los criterios de calidad establecidos por el marco teórico, así como la estructuración y la depuración de la base de datos proporcionada por 360 Smart Vision. Asimismo, se establece un umbral máximo de un 30 % de datos ausentes para la inclusión de los indicadores en cuestión (Wolf *et al.*, 2022), para salvaguardar así la fiabilidad del sintético elaborado. Considerando este y los demás criterios establecidos, se seleccionan los datos a través de MATLAB.
- Imputación de información ausente:** Llegados a este punto, la literatura enuncia la necesidad de comprender la causa que motiva la falta de ciertos datos, de manera que se caracterice el conjunto según su naturaleza MCAR, MAR o NMAR. En este caso, los análisis realizados refrendan la aleatoriedad de las ausencias, por lo que la aplicación de métodos de imputación es viable. Asimismo, la imputación múltiple es descrita como el método más robusto y fiable en la mayoría de los casos (OCDE, 2008). No obstante, otras técnicas tales como la imputación simple EM, cuya aplicación es además más sencilla, también gozan de buena reputación (Gold y Bentler, 2000). Sin embargo, dado que no existe una regla de oro, se ha implantado un análisis empírico de acuerdo con la comparación entre los valores NRMSE (error cuadrático medio normalizado) que genera cada técnica sobre los datos. Así, con una eliminación de datos conocidos, se estima el error cuadrático medio normalizado generado por cada imputación. Finalmente, este estudio ratifica la literatura, de manera que se aplica el método de imputación múltiple en IBM SPSS para casi la totalidad de las categorías.
- Análisis multivariante:** Tal y como se ha enunciado anteriormente, esta etapa arroja luz sobre el comportamiento multivariado de los indicadores originales. De esta manera, se aplican técnicas conocidas en la Estadística tales como el análisis de componentes principales, el estudio sobre el coeficiente alfa de Cronbach, el análisis *cluster* y otras complementarias. Se revela, por tanto, que los datos de ciertas categorías presentan correlaciones muy elevadas, lo cual puede constituir un problema para la ponderación y la agregación posteriores. De hecho, la matriz de correlación de algunas categorías no es linealmente independiente, es decir, existe cierta sobredimensionalidad en su conjunto de datos. Por ello, una vez se realizan este tipo de hallazgos, el autor debe valorar cómo

abordarlo de forma que se garantice la fiabilidad del indicador y se altere lo mínimo posible la estructura de datos establecida por el marco teórico. A fin de cuentas, todo proceso de elaboración goza de un cierto carácter iterativo (Burgass *et al.*, 2017). Esto se plantea, por tanto, en la evaluación intermedia de resultados.

- **Normalización:** Esta etapa debe valorar las múltiples técnicas de normalización existentes, como por ejemplo la estandarización, que genera una distribución de media nula y varianza unidad; la técnica Mín-Máx, que proporciona un conjunto de datos entre 0 y 1 y es totalmente compatible con todos los métodos de agregación; la puntuación, que es útil para *rankings*... En definitiva, cabe plantearse qué tipo de estructura resulta de cada uno de los métodos para, posteriormente, emplear los indicadores normalizados en la etapa de ponderación y agregación. Así, dadas las ventajas de la normalización Mín-Máx, se aplica esta a todas las categorías.
- **Evaluación intermedia de resultados:** Además de la continua revisión de los resultados obtenidos en cada etapa, este paso permite valorar el punto en que se encuentran los indicadores y las problemáticas reveladas a lo largo del proceso. De esta manera, se mantiene la claridad y la transparencia de las decisiones tomadas, así como se abordan los graves problemas de correlación revelados en el análisis multivariante. Para mitigarlo, se toman decisiones tales como la reducción de dimensiones a partir de la exclusión de indicadores redundantes. En un primer lugar, se decidió mantener un enfoque conservador en cuanto a la pertinencia de cada variable, pero los hallazgos obtenidos a partir del análisis multivariante pueden redefinir hasta cierto punto las decisiones del autor. En este sentido, es importante remarcar que el proceso de elaboración de un sintético debe ser dinámico y flexible. De esta manera, se reduce el número de indicadores de algunas categorías, garantizando así la robustez de los posteriores métodos de ponderación y agregación.
- **Ponderación y agregación:** Estas técnicas, a menudo, se presentan unidas. Después de todo, muchos métodos de ponderación exigen una agregación determinada. Por ello, en este punto, considerando la gran cantidad de datos que se manejan, algunos métodos como BAL (*Budget Allocation*) o AHP (*Analytic Hierarchy Process*) se deben descartar. Asimismo, la intención de elaborar un sintético con poco carácter compensatorio, es decir, que el rendimiento de una de las categorías no pueda simplemente sustituir al mal desempeño de otra, descarta otros como BOD (*Benefit of the Doubt*) o una posible agregación final aritmética. Por ello, se emplea la técnica PCA (*Principal Component Analysis*) para hallar los pesos de los indicadores de 360 Smart Vision que se subagregarán aritméticamente en torno a las categorías (primer nivel de agregación) definidas en el marco teórico, para posteriormente realizar una agregación geométrica en el sintético final (segundo nivel de agregación) que reduzca su carácter compensatorio. Para realizar la ponderación a través de esta técnica, es indispensable que las categorías se sometan a pruebas de adecuación para garantizar la robustez del resultado final.

De esta manera, se obtienen las series temporales que reflejan las categorías del sintético final; las cuales, posteriormente, se ponderan y se agregan geométricamente para dar lugar al índice de la *salud del mercado laboral español*. Dicha ponderación se lleva a cabo de acuerdo con el criterio experto de doctores de la Universidad, además de en consonancia

con el marco teórico elaborado. Así, *Empleo* y *Desempleo* representan cada uno el 25 % del sintético final, mientras que *Salarios y costes laborales* y *Habilidades y capacitación* aportan el 20 % cada uno. La última categoría, *Protección a los desempleados*, pondera un 10 %. Aun así, la existencia de índices subagregados permite establecer unos objetivos más concretos a los agentes sociales, de manera que se pueda monitorizar individualmente cada una de las categorías. Esto debe, por tanto, impulsar el análisis de los diferentes fenómenos que dan pie al mercado laboral.

- **Análisis de robustez:** Finalmente, esta etapa consiste en el estudio del impacto de las decisiones tomadas en las fases previas de la elaboración. Para realizarlo, se pueden llevar a cabo análisis de incertidumbre y de sensibilidad (Saisana *et al.*, 2005). Básicamente, implica considerar la implementación de otras técnicas en la creación del sintético y observar el comportamiento de los resultados obtenidos. En el caso concreto de este proyecto, se generan varios índices diferentes de acuerdo con la aplicación de métodos alternativos en las etapas de normalización, ponderación y agregación; así como el uso del conjunto original de datos sin la reducción de la evaluación intermedia. De esta manera, se generan 16 variaciones posibles, cuyas diferencias con el indicador construido muestran los impactos de cada técnica y la robustez general del sintético. El indicador propuesto por el proyecto es entonces comparado, concluyendo que es robusto y fiable dadas las pequeñas diferencias existentes entre los distintos índices obtenidos.

De esta manera, queda descrita la metodología seguida en el proyecto, la cual es aplicada de forma conceptual y de forma experimental por los capítulos 2 y 3 respectivamente.

Resultados

A través de la metodología explicada, se obtienen los subindicadores para las categorías estudiadas y también el índice final del mercado laboral. La evolución de los subgrupos, considerados fundamentales para caracterizar el mercado de trabajo, se presenta a continuación en la figura 2:



Figura 2. Evolución de los objetivos del indicador sintético

Dado que cada categoría refleja su propio fenómeno, resulta coherente que presenten comportamientos dispares. Además, los indicadores originales fueron orientados de tal manera que, a mayor valor, mejor situación del fenómeno que describían. Así se garantiza la comparabilidad de los datos y los índices contruidos a partir de ellos. De esta manera, las categorías de *Empleo*, *Desempleo* y *Salarios y costes laborales* presentan tendencias muy similares, con un mal desempeño en los peores meses de la pandemia. Por el contrario, *Habilidades y capacitación* y *Protección a los desempleados* son el reflejo del impacto positivo de las medidas de contención llevadas a cabo por los agentes sociales, ya que la irrupción del COVID-19 contribuyó a su mejoría, aunque posteriormente la extenuación de estas iniciativas propicia su reducción. No obstante, estos análisis son ciertamente superfluos, ya que los fenómenos causantes de los comportamientos descritos son harto complejos y deben ser analizados por economistas expertos.

Posteriormente, estas categorías se agregan mediante una técnica geométrica en el indicador sintético final. De esta manera, se sobrepenaliza al sintético si alguna de las categorías rinde de forma especialmente negativa, de manera que los agentes sociales no puedan descuidar ninguna de ellas para garantizar un mercado laboral sano. Este indicador es representado por la figura 3:



Figura 3. Evolución del índice de *salud del mercado laboral español*

En la mencionada figura, se representan algunos hitos de la pandemia de coronavirus. De hecho, al igual que con la figura 2, las causas exactas del comportamiento deben ser analizadas por un experto en la materia, que identifique correctamente los fenómenos propios y ajenos al mercado laboral que influyen sobre el mismo. No obstante, considerando las medidas de choque implementadas por los gobiernos y las empresas privadas, la evolución del sintético parece coherente con los hechos acaecidos. Después de todo, el mercado laboral parece responder a ciertas iniciativas como el impulso de los ERTes (expedientes de regulación temporal de empleo), la creación de los créditos ICO (instituto de crédito oficial) o incluso el IMV (ingreso mínimo vital) que, sin ser algo estrictamente laboral, puede tener cierto impacto sobre este mercado. Finalmente, la salud del mercado laboral parece recuperarse con fuerza tras la vacunación y el fin del segundo estado de alarma impuesto en España, coincidiendo además con hitos en las cifras de turismo nacional (Gutiérrez, 2021).

Finalmente, el proyecto arroja resultados importantes en cuanto a la robustez del sintético, ya que se mide cualitativamente el impacto de las técnicas aplicadas a lo largo de la elaboración.

De esta manera, la figura 4 presenta el sintético elaborado frente a la distribución que toma el conjunto formado por los 16 indicadores construidos a partir de las técnicas alternativas:

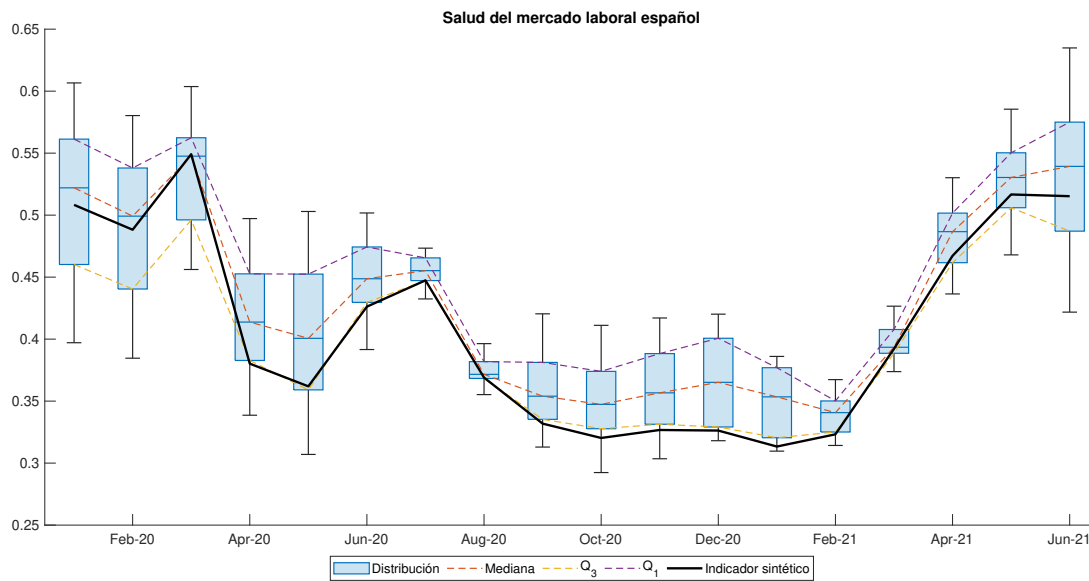


Figura 4. Evolución del sintético elaborado frente a la distribución de sus variaciones

El indicador obtenido, si bien se sitúa en la parte baja de la distribución, es coherente con lo esperado. Este análisis, desarrollado a lo largo de la sección 3.8, destaca al método geométrico y al criterio experto empleados en la ponderación y agregación finales como los elementos más influyentes sobre el resultado final. De hecho, estos factores explican muchos de los comportamientos observados en la figura 4, ya que la agregación geométrica tiende a rebajar el índice cuando alguna categoría se comporta de forma muy negativa (propio de los meses más crudos de la pandemia), así como el criterio experto evita que las categorías sean equiponderadas y, por tanto, que las medidas de protección a los desempleados puedan, por ejemplo, compensar totalmente el gran número de parados que conllevan. Por ello, según el criterio experto de doctores economistas de la Universidad Pontificia de Comillas, la equiponderación no resulta apropiada en este caso. De esta manera, el proyecto da a su fin con la elaboración del presente informe.

Conclusiones

La organización de información representa un valor añadido para todo analista. Por ello, la creación de un indicador sintético inexistente hasta la fecha proporciona a los agentes sociales una herramienta de monitorización de la recuperación socioeconómica en España. Este es, por tanto, un instrumento complementario para todo empresario, ciudadano o legislador interesado en la transformación y la mejora del bienestar social. La transparencia en la elaboración del sintético garantiza el acceso a toda la información subyacente, de forma que todos los interesados puedan comprender y criticar el sintético obtenido. A fin de cuentas, la elaboración de un indicador de esta índole implica una importante cantidad de decisiones a juicio del autor, por lo que se debe fundamentar debidamente cada elección tomada.

Asimismo, el análisis descrito en el capítulo 2 sienta las bases de un complemento útil a la literatura existente que permita optimizar las decisiones de futuros autores en la construcción de

nuevos indicadores, por lo que el impacto del presente trabajo excede el del sintético elaborado.

La creación del índice *salud del mercado laboral* proporciona unos resultados coherentes con la realidad laboral española, aunque el análisis pormenorizado de las causas que motivan su comportamiento queda abierto para futuros trabajos. Como colofón, este proyecto supone la aparición de nuevas líneas de investigación relacionadas con los indicadores sintéticos. En primer lugar, las técnicas utilizadas podrían emplearse para elaborar índices del resto de áreas presentes en la plataforma 360 Smart Vision (sanidad, consumo, actividad económica...), de manera que eventualmente se considere la creación de un sintético general referido a la situación socioeconómica de España. En este sentido, la aplicación de otras técnicas de ponderación y agregación puede resultar enriquecedora en dicho proceso. Adicionalmente, expertos economistas podrían tomar el testigo presentado por este proyecto y analizar los intrincados factores que causan el comportamiento del mercado laboral, tal y como es descrito por el indicador sintético elaborado.

En resumen, el proyecto supone un avance importante para la medición y el análisis del mercado laboral en España. Es, por tanto, una herramienta de transformación social.

Bibliografía

Burgass, M. J., Halpern, B. S., Nicholson, E., y Milner-Gulland, E. (2017). Navigating uncertainty in environmental composite indicators. *Ecological Indicators*, 75: 268-278.

Deloitte (2021). *360 Smart Vision* | Deloitte España. Disponible en: <https://www2.deloitte.com/es/es/pages/strategy/articles/360-smart-vision.html#dashboard> [Consultado 01-10-2021].

Freudenberg, M. (2003). Composite Indicators of Country Performance: A Critical Assessment. *OECD Science, Technology and Industry Working Papers*, 2003/16.

Gan, X., Fernandez, I. C., Guo, J., Wilson, M., Zhao, Y., Zhou, B., y Wu, J. (2017). When to use what: Methods for weighting and aggregating sustainability indicators. *Ecological Indicators*, 81: 491-502.

Gold, M. S. y Bentler, P. M. (2000). Treatments of Missing Data: A Monte Carlo Comparison of RBHDI, Iterative Stochastic Regression Imputation, and Expectation-Maximization. *Structural Equation Modeling: A Multidisciplinary Journal*, 7(3):319-355.

Greco, S., Ishizaka, A., Tasiou, M., y Torrisi, G. (2019). On the Methodological Framework of Composite Indices: A Review of the Issues of Weighting, Aggregation, and Robustness. *Social Indicators Research*, 141(1): 61-94.

Grupp, H. y Moge, M. E. (2004). Indicators for national science and technology policy: how robust are composite indicators? *Research Policy*, 33(9): 1373-1384.

Gutiérrez, H. (2021). El turismo nacional se desboca y supera en julio los récords previos a la pandemia. *El País*. Disponible en: <https://elpais.com/economia/2021-08-24/el-turismo-nacional-se-desboca-y-supera-en-julio-los-niveles-prepandemia.html> [Consultado 10-06-2021].

Nardo, M., Saisana, M., Saltelli, A., y Tarantola, S. (2005). Tools for Composite Indicators Building. *Joint Research Centre*.

OCDE (2008). Handbook on Constructing Composite Indicators: Methodology and User Guide. Disponible en: <https://www.oecd.org/sdd/42495745.pdf> [Consultado 02-10-2021].

Saisana, M., Saltelli, A., y Tarantola, S. (2005). Uncertainty and sensitivity analysis techniques as tools for the quality assessment of composite indicators. *Journal of the Royal Statistical Society: Series A (Statistics in Society)*, 168(2):307-323

Wolf, M. J., Emerson, J. W., Esty, D. C., de Sherbinin, A., y Wendling, Z. A. (2022). *Environmental Performance Index*. Yale Center for Environmental Law & Policy, New Haven.

Abstract

SYNTHETIC KPI TO MEASURE THE SPANISH RECOVERY FROM COVID-19

Author: Álvarez Fernández, Gustavo.

Director: Roch Dupré, David.

Collaborating Entity: ICAI - Comillas Pontifical University

Keywords: Composite indicator, Spanish labour market, Social impact, Weighting and aggregation, PCA.

Introduction

Coronavirus, or SARS-CoV-2, has been the cause of one of the most devastating pandemics in the history of mankind. After its irruption in 2019 in the Asian continent, the virus spread out quickly throughout the world. COVID-19, the disease caused by the virus, strongly disrupted the lives of citizens from both developed and developing countries. During those days, national governments took drastic political measures such as mobility bans or the interruption of all non-essential productive activities.

Labour and economic activities were forced to “hibernate”, so public and private stakeholders decided to foster a range of initiatives to mitigate the havoc wreaked by the pandemics. In spite of them, the figures of unemployment, the cost of salaries and even productivity have been seriously affected. This information, once scattered and disorganized, is now presented through the interactive 360 Smart Vision dashboard (2021), which provides an immeasurable number of indicators that reflect accurate and detailed information on the socio-economic recovery in Spain. In fact, the dashboard illustrates them through eight separate areas: financial markets, mobility, consumption, healthcare, labour market, citizen science, entrepreneurial activity, and sustainability and inclusion. However, if an interested party wants to assess the situation in some of these areas, he or she is forced to consider, in some cases, up to more than 100 indicators, which is probably unfeasible or could lead the person to significant assessing mistakes.

Therefore, there is no simple way to assess the phenomena described by the areas in the dashboard. Since it is almost impossible to analyse hundreds of them simultaneously, stakeholders may discard many of them as unimportant, while those considered relevant would have to be weighted as they go along. In this way, so much disaggregated information can lead to making the wrong decisions when analysing our environment. The lack of a synthetic indicator that standardises, weights and aggregates these indicators is a major handicap when it

comes to understanding the reality of the Spanish labour market. Thus, the development of a composite indicator based on the 360 Smart Vision database is more necessary than ever, so that those interested in this phenomenon can be provided with an accessible, simple and transparent indicator.

Composite or synthetic indicators are indexes composed from independent indicators or variables (Freudenberg, 2003), in such a way that they group them by means of different techniques in order to measure a multidimensional and complex phenomenon. In recent years, their creation has been boosted by the benefits they offer, in that they attract public attention to social problems and stimulate decision-making by public and private organisations (Greco *et al.*, 2019). In fact, previous detractors of synthetics, such as Amartya Sen, Nobel Prize winner in Economics in 1998, are gradually changing their critical stance thanks to the great social impact that synthetics can have, that is, thanks to their capacity for social transformation.

Of course, these detractors continue to exist. As main arguments, they highlight the oversimplification of information they entail, added to certain arbitrariness present in their elaboration (Grupp y Moge, 2004). This can lead social agents to misinterpret the synthetic and, therefore, to extract misleading messages from it. However, to avoid such problems, it is essential to follow a clear process in its creation, as established by the OECD (2008) and other public entities (Nardo *et al.*, 2005). Likewise, the evaluation of results must be a constant throughout the procedure, so that the sources of uncertainty and bias introduced throughout the creation process are identified. This is key to ensure the robustness and reliability of the final composite indicator.

First, the theoretical framework of the indicator must be established, which delimits the concepts studied, determines the subgroups of the synthetic and, in addition, defines the quality criteria to be maintained throughout the project. Secondly, data collection consists of the structuring, filtering and selection of variables according to certain precepts such as relevance, timeliness, efficiency or temporal completeness. Since some variables do not enjoy the totality of cases, the imputation of missing information should be the third stage of the procedure. Fourth, multivariate analysis yields fundamental conclusions about the behavior of the data and the correlations between variables, so that they can be considered in the subsequent stages of weighting and aggregation.

Once the nature of the data is understood, normalisation provides a common scale for all the original variables. After this step, the data are ready for weighting and aggregation, especially considering the synthetic objectives established in the theoretical framework and the information revealed by the multivariate analysis. In this way, the weighting phase assigns weights to each variable according to its relative importance, although this must be assessed in conjunction with the aggregation technique employed afterwards. After all, some methods allow for total trade-off between original variables, so weighting must bear it in mind in case this trade-off is not desired.

Likewise, the weights must also take into account the correlation between variables, since the double counting problem may be incurred (Greco *et al.*, 2019). The aggregation would then unite the different variables around the categories, which would be then aggregated into the final synthetic. The arithmetic and geometric techniques stand out in this regard. To conclude, the robustness analysis quantifies the impact of the decisions made throughout the elaboration by creating different outputs from alternative choices. This shows how decisive the techniques

used are on the final result, so that the greater the variation of the other alternative outputs, the lower the robustness of the synthetic. After this step, the subsequent disclosure concludes the elaboration process.

These stages are very contrasted by the bibliography, since each one enjoys a much greater application in other fields of Statistics. Thus, one of the objectives of the composite will be to analyse and update the existing elaboration manuals as of 2022, so that future authors will have a renewed and contrasted complement that will allow them to make appropriate decisions according to the type of composite they wish to construct. Of course, the other and, if possible, more important objective is the effective implementation of the synthetic indicator on the labour market in Spain, so as to fill the existing gap in this respect. In this way, the project will promote not only social transformation based on the aforementioned indicator, but also through the establishment of pillars that will allow future authors to develop their own synthetic indicators. To this end, the dissemination of the results should also be duly promoted.

Once the initial approach to composite indicators has been made, the methodology of the elaboration process is explained.

Methodology

The methodology applied in the work follows the process set out in the OECD's handbook (2008). However, as one of the objectives is to establish a new conceptual revision on the elaboration process for composite indicators, the project is mainly divided into two chapters. Firstly, the chapter 2, where any given indicator is conceptually designed, critically showing the details of each technique applicable in all the above-mentioned stages. Secondly, the chapter 3 reflects the actual implementation of the findings of its predecessor, carefully safeguarding the transparency and argumentation of any decisions taken in its process. For this reason, both chapters have to follow the structure of elaboration of a synthesis of any of the above.

Thus, after analysing the different techniques and approaches, the decisions taken at each stage are as follows:

- **Theoretical framework:** The *Health of the Spanish labour market* represents the phenomenon described by the constructed synthetic indicator, understood as the situation in which the “market agents see their needs satisfied in a sustainable way over time”. In this way, the synthetic indicator evaluates not only the numbers related to the unemployment rate, but also defines aspects such as productivity, the temporality of work, the training of the population, the protection of the State for the unemployed...

Basically, five categories make up this phenomenon. They are divided according to their nature: as causes, *Skills and training*, *Protection of the unemployed* and *Wages and labour costs*; and as consequences of the health of the market, *Unemployment* and *Employment*, since each one characterises the type of unemployment or contracts (for example) present in society. Finally, the criteria for data selection are established, including relevance, timeliness, spatial and temporal circumscription, accessibility and accuracy.

Thus, the theoretical framework of the project presents a nested structure which is depicted in the figure below:

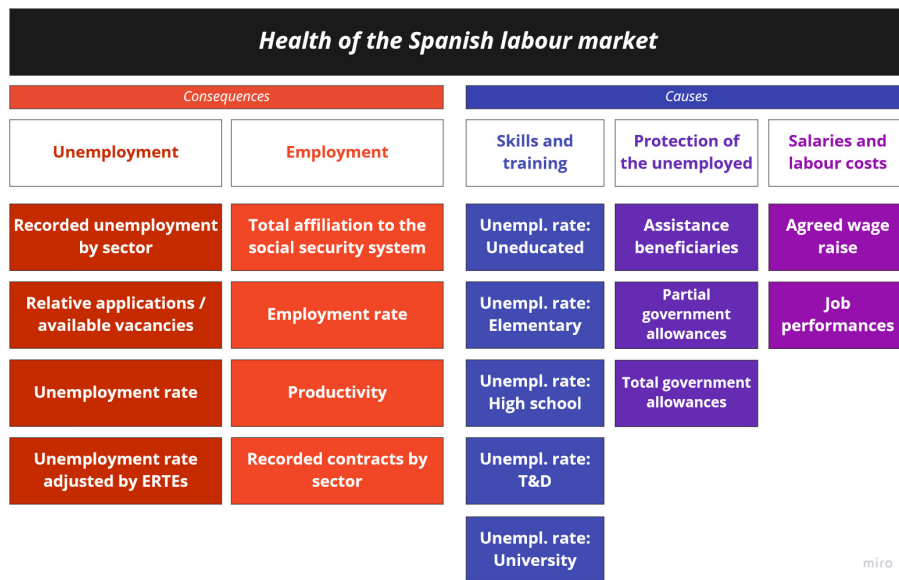


Figure 1. Nested structure of the composite indicator

- Data collection:** The collection involves the application of the quality criteria established by the theoretical framework, as well as the structuring and cleaning of the database provided by 360 Smart Vision. In addition, a maximum threshold of 30% of missing data is set for the inclusion of individual indicators (Wolf *et al.*, 2022), in order to safeguard the reliability of the produced composite. Considering this and the other established criteria, the data are selected using the MATLAB tool.
- Imputation of missing data:** At this point, the literature states the need to understand the cause of the lack of certain data, so as to characterise the set according to its MCAR, MAR or NMAR nature. In this case, the analyses carried out confirm the randomness of the absences, so that the application of imputation methods is feasible. Furthermore, multiple imputation is described as the most robust and reliable method in most cases (OCDE, 2008). However, other techniques such as simple EM imputation, which are also simpler to apply, have good reputation (Gold y Bentler, 2000). However, since there is no golden rule, an empirical analysis has been implemented according to the comparison between the NRMSE (normalised root mean square error) values generated by each technique on the data. Thus, with a known data elimination, the normalised-mean-square-error generated by each imputation is estimated. Finally, this study ratifies the literature, so that the multiple imputation method is applied in IBM SPSS for almost all categories.
- Multivariate analysis:** As stated above, this stage sheds light on the multivariate behaviour of the original indicators. In this way, well-known statistical techniques such as principal component analysis, the determination of Cronbach's alpha coefficient, cluster analysis and other complementary techniques are applied. It is thus revealed that data in certain categories show very high correlations, which can be a problem for subsequent weighting and aggregation. In fact, the correlation matrix of some categories is not linearly independent, i.e. there is some overdimensionality in their data set. Therefore, once such a finding is made, the author must consider how to approach it in a way that ensures the reliability of the indicator and minimises disruption to the data structure

established by the theoretical framework. After all, any elaboration process has a certain iterative character (Burgass *et al.*, 2017). Therefore, this is discussed in the mid-term evaluation of results.

- **Normalisation:** This stage should assess the many existing normalisation techniques, such as standardisation, which generates a distribution of zero mean and unit variance; the Min-Max technique, which provides a dataset between 0 and 1 and is fully compatible with all aggregation methods; scoring, which is useful for rankings.... In short, it is worth considering what type of structure results from each of the methods in order to subsequently use the standardised indicators in the weighting and aggregation stage. Thus, given the advantages of Min-Max normalisation, this technique is applied to all categories.
- **Intermediate evaluation of results:** In addition to the continuous review of the results obtained at each stage, this step allows for the assessment of the indicators' status and the issues revealed throughout the process. In this way, the clarity and transparency of the decisions taken is maintained, as well as addressing the serious correlation problems revealed in the multivariate analysis. To mitigate this, decisions such as the reduction of dimensions through the exclusion of redundant indicators are taken. Initially, it was decided to maintain a conservative approach to the relevance of each variable, but the findings obtained from the multivariate analysis may redefine the author's decisions to some extent. In this sense, it is important to stress that the process of elaborating a composite should be dynamic and flexible. In this way, the number of indicators within some categories is decreased, thus ensuring the robustness of subsequent weighting and aggregation methods.
- **Weighting and aggregation:** These techniques are often presented together. After all, many weighting methods require a certain aggregation. Therefore, at this point, considering the large amount of data involved and the desire to produce a time series index and not a ranking, many methods such as BAL (Budget Allocation) or AHP (Analytic Hierarchy Process) should be discarded. Also, the intention to produce a composite with little compensatory character, i.e. that the performance of one of the categories cannot simply substitute for the poor performance of another, rules out other methods such as BOD (Benefit-Of-The-Doubt) and the arithmetic method for the final aggregation. Thus, the PCA (Principal Component Analysis) technique is used to sub-aggregate the 360 Smart Vision indicators around the categories (first-level aggregation) defined in the theoretical framework, and then perform a geometric aggregation (second-level) into the final synthetic indicator, in order to reduce the chances of trade-offs. In order to perform the weighting through this technique, it is also essential that the categories are subjected to fit tests.

Thus, the time series that reflect the categories of the final synthetic are obtained, which are subsequently weighted and aggregated to give rise to the index of the *health of the Spanish labour market*. This weighting is carried out in accordance with the expert judgement of the University's PhDs, as well as in line with the theoretical framework developed. Thus, *Employment* and *Unemployment* each represent 25% of the final synthetic, while *Wages and labour costs* and *Skills and training* contribute 20% each. The

last category, "Protection of the unemployed", weights 10 per cent. Still, the existence of sub-aggregated indices allows for more concrete targets to be set for the social partners, so that each of the categories can be monitored individually. This should, therefore, boost the analysis of the different phenomena that give rise to the labour market.

- **Robustness analysis:** Finally, this stage consists of the study of the impact of the decisions taken in the previous phases of the elaboration. To do this, uncertainty and sensitivity analyses can be carried out (Saisana *et al.*, 2005). Basically, it involves considering other normalisation, weighting or aggregation techniques and observing the behaviour of the results. In this way, several indicators are generated by the application of alternative methods in the normalisation, weighting and aggregation processes; as well as in the use of the original data set without the reduction of the intermediate evaluation. Hence, 16 possible variations are created, whose differences with the proposed composite indicator show the impacts of each technique and the overall robustness of the synthetic. The proposed composite indicator is compared, thus concluding that it shows due robustness and reliability.

Thus, the methodology followed in the project is described, so the process is established for both the conceptual and the experimental parts (see chapters 2 and 3 respectively).

Results

Through the methodology explained, the sub-indicators for the categories studied and also the final labour market index are obtained. The evolution of the subgroups, considered fundamental to characterise the labour market, is presented in the figure below:

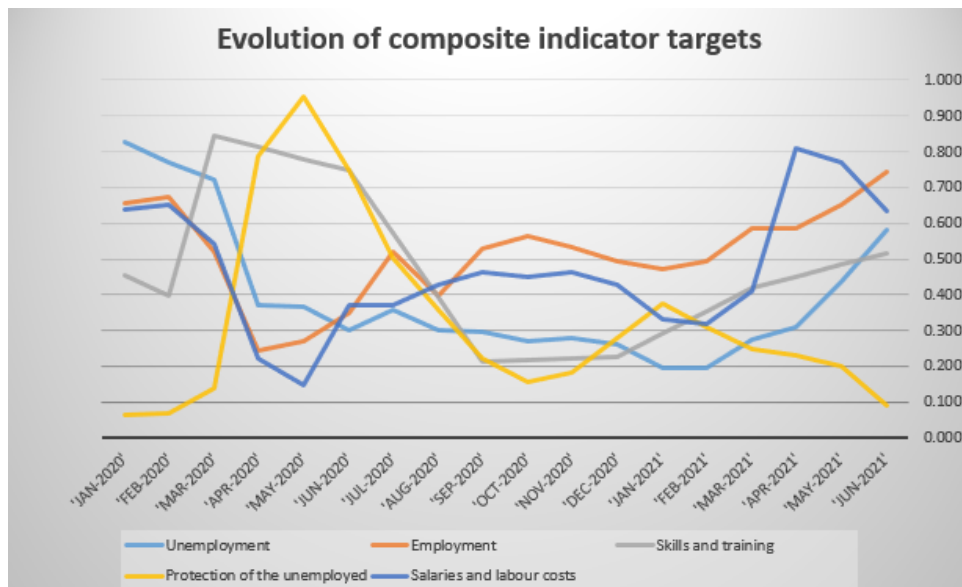


Figure 2. Evolution of synthetic indicator targets

Given that each category reflects its own phenomenon, it is consistent that they have different behaviours. Moreover, the original indicators were oriented in such a way that the higher the value, the better the situation of the phenomenon they describe. This ensures the

comparability of the data and the indices constructed from them. Thus, the categories of *Employment*, *Unemployment* and *Wages and labour costs* present very similar trends, with a poor performance in the worst months of the pandemic.

On the other hand, *Protection of the unemployed* and *Skills and training* reflect the positive impact of the containment measures carried out by the social agents, since the irruption of COVID-19 contributed to their improvement, although the exhaustion of the measures subsequently led to their reduction. However, these analyses are certainly superfluous, as the phenomena causing the behaviour described are highly complex and should be analysed by expert economists.

These categories are then aggregated using a geometric technique in the final synthetic indicator. In this way, the synthetic is over-penalised if any of the categories perform particularly negatively, so that social partners cannot neglect any of them to ensure a healthy labour market. This indicator is represented by figure 3:



Figure 3. Evolution of the index *Health of the Spanish labour market*

In the figure above, some milestones of the coronavirus pandemic are depicted. In fact, as with the figure 2, the exact causes of the behaviour have to be analysed by an expert in the field, who correctly identifies the labour market and non-labour market phenomena influencing it. However, considering the shock measures implemented by governments and private companies, the evolution of the synthetic seems to be consistent with the facts. After all, the labour market seems to respond to certain initiatives such as the promotion of the ERTes, the creation of the ICO credits or even the IMV which, without being strictly labour-related, may have some impact on this market. Finally, the health of the labour market seems to be recovering strongly after the vaccination and the end of the second state of alarm imposed in Spain, coinciding also with milestones in national tourism figures (Gutiérrez, 2021).

Finally, the project yields important results in terms of the robustness of the synthetic, as it qualitatively measures the impact of the techniques applied throughout the elaboration. In

this way, the figure 4 presents the synthetic developed against the distribution taken by the set formed by the 16 indicators constructed from the alternative techniques:

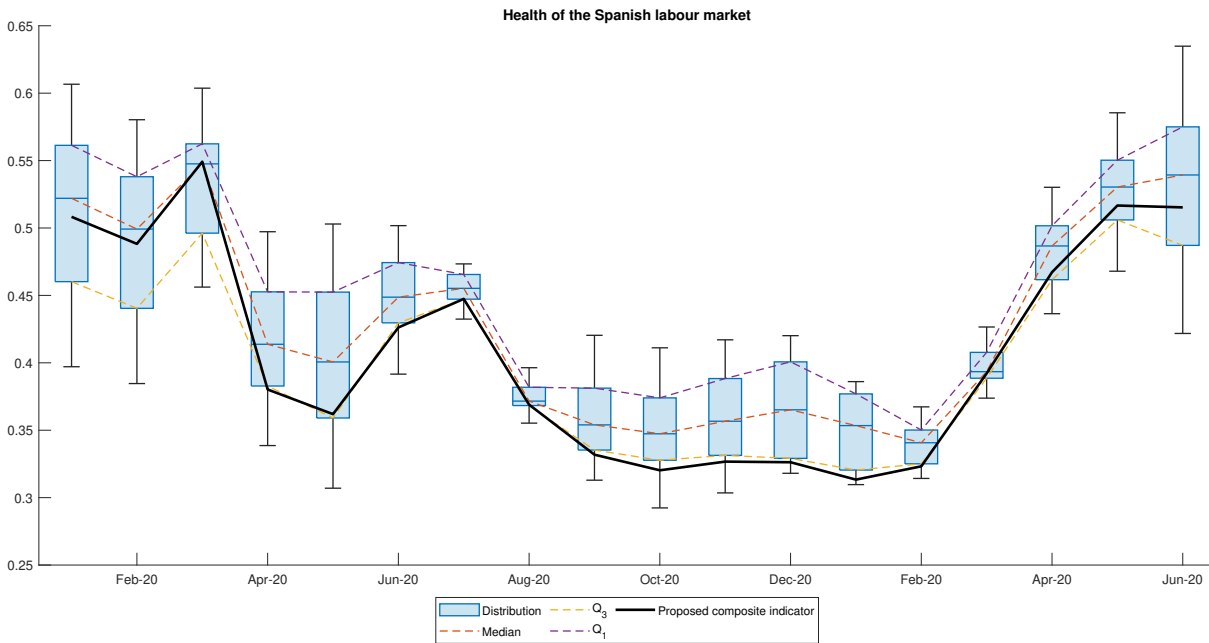


Figure 4. Evolution of the proposed synthetic vs. the distribution of its variations

The indicator obtained, although is located in the lower part of the distribution, is coherent with the expected result. This analysis, developed throughout the section 3.8, highlights the expert criteria and the geometric method used in the final weighting and aggregation phases as the most influential elements on the final result. In fact, these factors explain many of the behaviours observed in the figure, since geometric aggregation tends to lower the index when some categories behave very negatively (typical of the worst months of the pandemic) and the expert criterion prevents the categories from being weighted equally and, therefore, the measures to protect the unemployed can, for example, fully compensate for the large number of unemployed that they entail. According to the expert opinion of economist doctors of the Comillas Pontifical University, equal weighting is not appropriate in this case. The project thus comes to an end with the preparation of this report.

Conclusions

The organisation of information represents an added value for any analyst. The creation of a synthetic indicator that has not existed to date provides social agents with a tool for monitoring socio-economic recovery in Spain. It is therefore a complementary instrument for any businessman, citizen or legislator interested in the transformation and improvement of social welfare. Transparency in the preparation of the composite indicator guarantees access to all the underlying information, so that all stakeholders can understand and criticise the synthetic obtained. At the end of the day, the development of any indicator of this nature involves a significant number of decisions in the author's judgement, and so an attempt has been made to properly inform each choice made.

Furthermore, the analysis described in chapter 2 lays the foundations for a new useful complement to optimise decisions on the elaboration of future indicators, so that the impact

of the present work exceeds that of the developed composite indicator.

The creation of the labour market health index provides results that are coherent with the Spanish labour market reality, although the detailed analysis of the causes that motivate its behaviour remains open for future lines of research. Finally, this work implies the emergence of new lines of research related to synthetic indicators. Firstly, the techniques employed could be used to create indices for the rest of the areas present in the 360 Smart Vision platform (health, consumption, economic activity...), so that the creation of a synthetic index referring to Spain's socio-economic situation could eventually be considered. In this sense, the application of other weighting and aggregation techniques could be enriching for the process. In addition, expert economists could take the baton presented by this project and analyse the reasons behind the behaviour of the labour market, as it is described by the developed synthetic indicator.

In short, the project represents an important step forward in the measurement and analysis of the Spanish labour market. It is, therefore, a helpful tool for social transformation.

References

- Burgass, M. J., Halpern, B. S., Nicholson, E., & Milner-Gulland, E. (2017). Navigating uncertainty in environmental composite indicators. *Ecological Indicators*, 75: 268-278.
- Deloitte (2021). *360 Smart Vision* | *Deloitte España*. Available at: <https://www2.deloitte.com/es/es/pages/strategy/articles/360-smart-vision.html#dashboard> [Accessed 01-10-2021].
- Freudenberg, M. (2003). Composite Indicators of Country Performance: A Critical Assessment. *OECD Science, Technology and Industry Working Papers*, 2003/16.
- Gan, X., Fernandez, I. C., Guo, J., Wilson, M., Zhao, Y., Zhou, B., & Wu, J. (2017). When to use what: Methods for weighting and aggregating sustainability indicators. *Ecological Indicators*, 81: 491-502.
- Gold, M. S. y Bentler, P. M. (2000). Treatments of Missing Data: A Monte Carlo Comparison of RBHDI, Iterative Stochastic Regression Imputation, and Expectation-Maximization. *Structural Equation Modeling: A Multidisciplinary Journal*, 7(3):319-355.
- Greco, S., Ishizaka, A., Tasiou, M., & Torrisi, G. (2019). On the Methodological Framework of Composite Indices: A Review of the Issues of Weighting, Aggregation, and Robustness. *Social Indicators Research*, 141(1): 61-94.
- Grupp, H. & Mogege, M. E. (2004). Indicators for national science and technology policy: how robust are composite indicators? *Research Policy*, 33(9): 1373-1384.
- Gutiérrez, H. (2021). El turismo nacional se desboca y supera en julio los récords previos a la pandemia. *El País*. Available at: <https://elpais.com/economia/2021-08-24/el-turismo-nacional-se-desboca-y-supera-en-julio-los-niveles-prepandemia.html> [Accessed 10-06-2021].
- Nardo, M., Saisana, M., Saltelli, A., y Tarantola, S. (2005). Tools for Composite Indicators Building. *Joint Research Centre*.

OCDE (2008). Handbook on Constructing Composite Indicators: Methodology and User Guide. Available at: <https://www.oecd.org/sdd/42495745.pdf> [Accessed 02-10-2021].

Saisana, M., Saltelli, A., & Tarantola, S. (2005). Uncertainty and sensitivity analysis techniques as tools for the quality assessment of composite indicators. *Journal of the Royal Statistical Society: Series A (Statistics in Society)*, 168(2):307-323

Wolf, M. J., Emerson, J. W., Esty, D. C., de Sherbinin, A., & Wendling, Z. A. (2022). *Environmental Performance Index*. Yale Center for Environmental Law & Policy, New Haven.

*A vosotros,
a mi familia.*

Temo al hombre de un solo libro
SANTO TOMÁS DE AQUINO

Agradecimientos

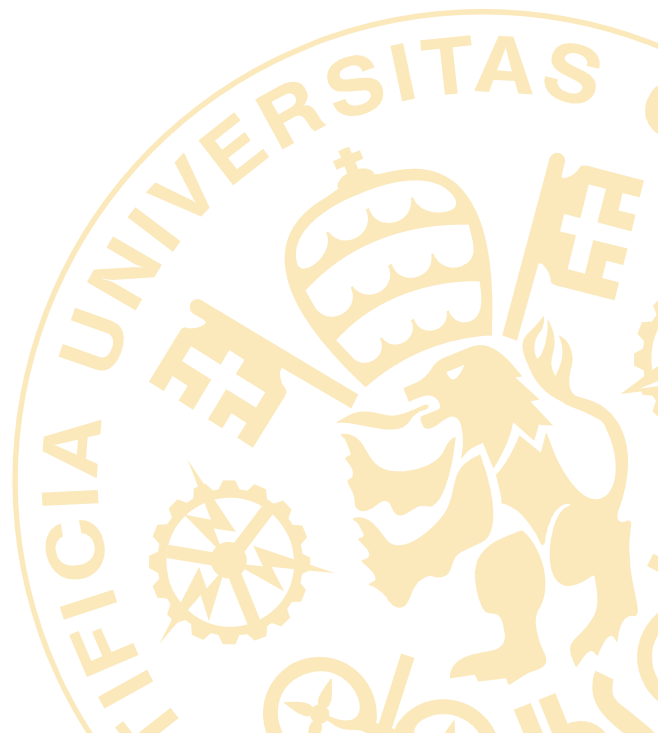
No puedo sino comenzar por ti, David, que me has acompañado desde el 1 de septiembre. Creo, sinceramente, que eres el tutor que todo alumno querría tener. Tus ganas, tu cercanía, tu disposición... todo suma para haber hecho de este largo proyecto un trabajo del que estar orgullosos. Efectivamente, este trabajo es nuestro: de los dos. Muchas gracias por tus consejos, tu constancia y tu amistad. Sin ninguna duda, volvería a elegir este proyecto contigo.

Os doy también las gracias a vosotros, a mi familia, por ser mi norte durante la larga noche. Os lo debo todo, por lo que este TFG es para vosotros.

Por último, gracias a los que, por un motivo u otro, habitáis una pequeña parte de mi corazón. Gracias por permitirme crecer a vuestro lado.

DOCUMENTO I

MEMORIA



Índice

I. Memoria	13
1. Introducción	15
1.1. La pandemia y el mercado laboral español	15
1.2. Estado de la cuestión	16
1.3. Motivación	21
1.4. Objetivos	22
1.5. Recursos	23
1.6. Metodología	24
2. Diseño conceptual de un indicador sintético	27
2.1. Marco teórico	27
2.2. Recolección de datos	31
2.3. Imputación de información ausente	33
2.4. Análisis multivariante	38
2.5. Normalización	43
2.6. Ponderación y agregación	47
2.6.1. Ponderación	47
2.6.2. Agregación	53
2.7. Análisis de robustez	56
2.8. Divulgación	58
3. Indicador sintético 360 Smart Vision sobre el mercado laboral español	61
3.1. Marco teórico aplicado	61
3.2. Datos de 360 SMART VISION	65
3.3. Imputación de información ausente	69
3.4. Análisis multivariante	73
3.5. Normalización	79
3.6. Evaluación intermedia de resultados	81
3.7. Ponderación y agregación	85
3.8. Análisis de robustez	95
4. Conclusiones	103
Bibliografía	106
II. Anexos	115
A. Alineación con los ODS	117
B. Tablas de datos	119
C. Figuras	131
D. Código fuente	145

Índice de figuras

1. Representación de la estructura de los indicadores sintéticos	19
2. Diagrama de flujo sobre el proceso de construcción del indicador sintético y sus objetivos . .	24
3. Dendograma a partir del enlace promedio y del cuadrado de la distancia euclídea	42
4. Matriz de histogramas para diferentes técnicas de normalización	46
5. Frontera óptima de desempeño creada por los países A, B y C mediante DEA	50
6. Estructura anidada del sintético junto a los subgrupos e indicadores componentes principales	65
7. Muestra de los datos originales proporcionados por la plataforma	66
8. Instantánea de la tabla de estructuración y selección de datos en el subgrupo <i>Empleo</i>	68
9. Resumen de los criterios de calidad aplicados y su implementación	68
10. Instantánea del subgrupo <i>Empleo</i> tras la eliminación de casos con información ausente	71
11. Instantánea de la categoría <i>Empleo</i> tras la imputación múltiple en IBM SPSS	72
12. Representación del subgrupo <i>Empleo</i> tras la aplicación de la imputación múltiple	73
13. Matriz de correlaciones de los indicadores de la categoría <i>Empleo</i>	74
14. Instantánea del subgrupo <i>Empleo</i> tras la normalización y reorientación de los datos	81
15. Evolución de los subgrupos que conforman la <i>Salud del mercado de trabajo</i>	91
16. Gráfico apilado de la evolución de los subgrupos del indicador sintético	91
17. Evolución individualizada de la categoría <i>Habilidades y capacitación</i>	92
18. Evolución de la <i>salud del mercado laboral español</i> tras la agregación final	94
19. Árbol de decisiones potenciales para cada etapa del proceso	96
20. Evolución de los objetivos del sintético según las diferentes técnicas de la elaboración	97
21. Evolución del indicador sintético según las diferentes técnicas de la elaboración	98
22. Evolución del sintético elaborado frente a la distribución de sus variaciones	101
23. Instantánea de la tabla de estructuración y selección de datos en el subgrupo <i>Desempleo</i> . . .	131
24. Estructuración y selección de datos en el subgrupo <i>Habilidades y capacitación</i>	131
25. Estructuración y selección de datos en el subgrupo <i>Protección a los desempleados</i>	132
26. Estructuración y selección de datos en el subgrupo <i>Salarios y costes laborales</i>	132
27. Instantánea de la categoría <i>Empleo</i> tras la imputación EM en IBM SPSS	133
28. Instantánea de la categoría <i>Empleo</i> tras la imputación por regresión	133
29. Subgrupo <i>Desempleo</i> tras la aplicación de la imputación regresiva	133
30. Subgrupo <i>Habilidades y capacitación</i> tras la aplicación de la imputación múltiple	133
31. Matriz de correlaciones de los indicadores de la categoría <i>Desempleo</i>	134
32. Matriz de correlaciones de los indicadores de <i>Habilidades y capacitación</i>	134
33. Matriz de correlaciones de los indicadores de <i>Protección a los desempleados</i>	135
34. Matriz de correlaciones de los indicadores de <i>Salarios y costes laborales</i>	135
35. Dendograma del análisis <i>cluster</i> de la categoría <i>Desempleo</i>	136
36. Dendograma del análisis <i>cluster</i> de la categoría <i>Empleo</i>	136
37. Dendograma del análisis <i>cluster</i> de la categoría <i>Habilidades y capacitación</i>	136
38. Dendograma del análisis <i>cluster</i> de la categoría <i>Protección a los desempleados</i>	137
39. Dendograma del análisis <i>cluster</i> de la categoría <i>Salarios y costes laborales</i>	137

40. Instantánea del subgrupo <i>Desempleo</i> tras la normalización y reorientación de los datos . . .	137
41. <i>Habilidades y capacitación</i> tras la normalización y reorientación de los datos	138
42. <i>Protección a los desempleados</i> tras la normalización y reorientación de los datos	138
43. Instantánea de <i>Salarios y costes laborales</i> tras la normalización y reorientación de los datos .	138
44. Evolución individualizada de la categoría <i>Desempleo</i>	139
45. Evolución individualizada de la categoría <i>Empleo</i>	139
46. Evolución individualizada de la categoría <i>Protección a los desempleados</i>	140
47. Evolución individualizada de la categoría <i>Salarios y costes laborales</i>	140
48. Evolución de los objetivos del sintético con saturación en la normalización	141
49. Evolución de los objetivos del sintético a partir de los indicadores originales	141
50. Evolución de los objetivos del sintético original con saturación en la normalización	142
51. Indicador sintético con todos los indicadores originales tras la agregación final aritmética . .	142
52. Evolución del indicador sintético tras la agregación geométrica	143
53. Evolución del sintético con todos los indicadores originales tras la agregación aritmética . .	143
54. Evolución del sintético con los indicadores originales tras la agregación geométrica	144

Índice de tablas

1. Oportunidades e inconvenientes de un indicador sintético cualquiera	18
2. Resumen de las características del ordenador empleado en el proyecto	24
3. Síntesis de las principales técnicas de imputación de información ausente	37
4. Ejemplo de una matriz de componentes obtenida a partir del método PCA	40
5. Matriz de componentes rotada con el método <i>varimax</i>	40
6. Comparación de las matrices rotada y rotada cuadrática normalizada	49
7. Prueba MCAR de Little de las categorías <i>Desempleo</i> , <i>Empleo</i> y <i>Habilidades y Capacitación</i> .	70
8. Análisis preliminar de las variables con datos ausentes	70
9. Comparativa de los NRMSE para los métodos de imputación disponibles en IBM SPSS . . .	72
10. Matriz de varianza total explicada por los componentes obtenidos por PCA para <i>Empleo</i> . .	75
11. Matriz de componentes rotados obtenidos por PCA para la categoría <i>Empleo</i>	76
12. Adecuación de <i>Protección a los desempleados</i> y <i>Salarios y costes laborales</i> para PCA . . .	77
13. Coeficientes alfa de Cronbach para las diferentes categorías	78
14. Número de indicadores linealmente dependientes en cada categoría	83
15. Nueva adecuación de <i>Desempleo</i> , <i>Empleo</i> y <i>Habilidades y capacitación</i> para PCA	85
16. Matriz de componentes rotada con rotación <i>varimax</i> de la categoría <i>Empleo</i>	86
17. Comparación de las matrices rotada y rotada cuadrática normalizada	86
18. Indicadores intermedios obtenidos para la categoría <i>Empleo</i>	87
19. Pesos de los indicadores intermedios de <i>Empleo</i> obtenidos de la varianza explicada	89
20. Pesos de los diferentes indicadores para la subagregación en la categoría <i>Empleo</i>	89
21. Subindicadores de las diferentes categorías que forman la <i>Salud del mercado de trabajo</i> . . .	90
22. Indicador sintético final: <i>La salud del mercado laboral español</i>	93
23. Técnicas aplicadas en cada etapa del proyecto y sus alternativas	95
24. Lista completa de los indicadores componentes del trabajo junto a su granularidad	122
25. Matriz de componentes rotados obtenidos por PCA para la categoría <i>Desempleo</i>	122
26. Matriz de varianza total explicada por los componentes obtenidos por PCA para <i>Desempleo</i> .	123
27. Matriz de componentes rotados de PCA para la categoría <i>Habilidades y capacitación</i>	123
28. Matriz de varianza explicada por los componentes de PCA para <i>Habilidades y capacitación</i> .	124
29. Matriz de componentes rotados de PCA para la categoría <i>Salarios y costes laborales</i>	124
30. Matriz de varianza explicada por los componentes de PCA para <i>Salarios y costes laborales</i> .	124
31. Correlación y alfa de Cronbach tras excluir variables de <i>Desempleo</i>	125
32. Correlación y alfa de Cronbach tras excluir variables de <i>Empleo</i>	125
33. Correlación y alfa de Cronbach tras excluir variables de <i>Habilidades y capacitación</i>	126
34. Correlación y alfa de Cronbach tras excluir variables de <i>Protección a los desempleados</i> . . .	126
35. Correlación y alfa de Cronbach tras excluir variables de <i>Salarios y costes laborales</i>	127
36. Número de indicadores que cumplen los criterios para la transformación logarítmica	127
37. Matriz de componentes rotada con rotación <i>varimax</i> de la categoría <i>Desempleo</i>	128
38. Matriz de componentes rotada por <i>varimax</i> de la categoría <i>Habilidades y capacitación</i> . . .	128
39. Matriz de componentes rotada por <i>varimax</i> de la categoría <i>Salarios y costes laborales</i>	129

40. Pesos de los diferentes indicadores para la subagregación aritmética en <i>Desempleo</i>	129
41. Pesos de los indicadores para la subagregación aritmética en <i>Habilidades y Capacitación</i> . .	129
42. Pesos de los indicadores para la subagregación aritmética en <i>Salarios y costes laborales</i> . . .	130

Índice de extractos de código

1. Selección de datos y creación de archivos para la imputación	151
2. Comparación de métodos de imputación según su error NRMSE	152
3. Proceso de normalización Mín-Máx y la reorientación de escala posterior	154
4. Identificación de sobredimensionalidad en las matrices de correlación de las categorías	155
5. Reducción de las categorías según la evaluación intermedia de resultados	156
6. Ponderación y agregación del sintético final junto al análisis de robustez	168

Acrónimos

<i>AHP</i>	Analytical Hierarchy Process
<i>BAL</i>	Budget Allocation
<i>BOD</i>	Benefit of the Doubt
<i>CA</i>	Conjoint Analysis
<i>DEA</i>	Data Envelopment Analysis
<i>EC</i>	European Commission
<i>EM</i>	Expectation-Maximization
<i>EPI</i>	Environmental Performance Index
<i>ERTE</i>	Expediente de Regulación Temporal de Empleo
<i>ESI</i>	Environmental Sustainable Index
<i>FA</i>	Factor Analysis
<i>ICAI</i>	Instituto Católico de Artes e Industrias
<i>ICO</i>	Instituto de Crédito Oficial
<i>IDH</i>	Índice de Desarrollo Humano
<i>IM</i>	Imputación Múltiple
<i>IMV</i>	Ingreso Mínimo Vital
<i>INE</i>	Instituto Nacional de Estadística
<i>KMO</i>	Kaiser-Meyer-Olkin
<i>KPI</i>	Key Performance Indicator
<i>MAE</i>	Mean Absolute Error
<i>MAR</i>	Missing at Random
<i>MCA</i>	Multi-Criteria Approach
<i>MCAR</i>	Missing Completely at Random
<i>NMAR</i>	Not Missing at Random
<i>NRMSE</i>	Normalized Root Mean Square Error
<i>OCDE</i>	Organización para la Cooperación y el Desarrollo Económicos
<i>OHI</i>	Ocean Health Index
<i>PCA</i>	Principal Component Analysis
<i>RMSE</i>	Root Mean Square Error
<i>SEPE</i>	Servicio Público de Empleo Estatal
<i>SSI</i>	Sustainable Society Index
<i>TAI</i>	Technology Achievement Index
<i>TFG</i>	Trabajo de Fin de Grado

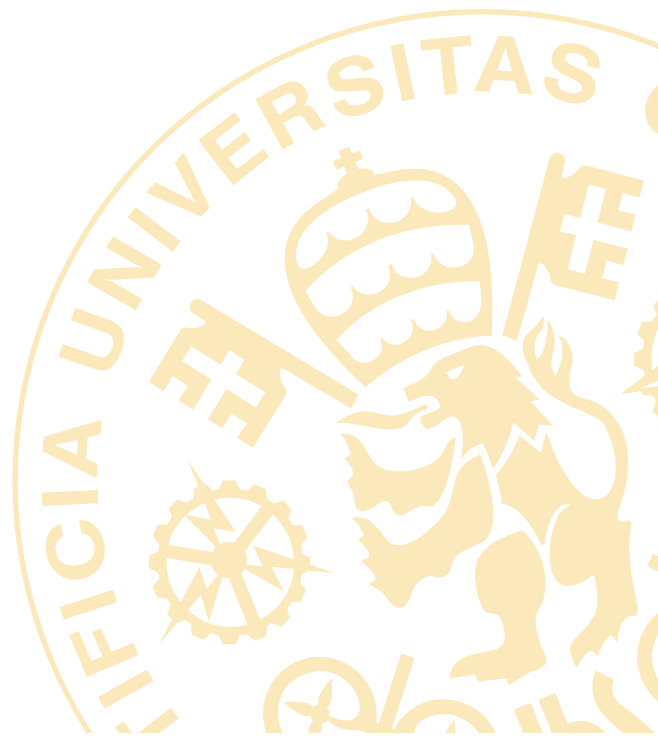
Símbolos

Como regla general, el subíndice i representa a cada indicador componente y t el momento del tiempo que describe (en caso de una serie temporal):

α	Coeficiente alfa de Cronbach
α_i	Peso resultante de un indicador i sobre el total en la técnica PCA
CI_t	Valor del indicador sintético para un tiempo t
F_{it}	Valor del factor i para un tiempo t obtenido por FA
I_{it}	Valor normalizado del indicador i para un tiempo t
λ_{ji}	Carga factorial rotada del indicador i para el componente j
$NRMSE_i$	Valor del error cuadrático medio normalizado del indicador i
N_i	Número de datos del conjunto del indicador i
Q_n	Cuartil n de una distribución
$RMSE_i$	Valor del error cuadrático medio del indicador i
σ_{ji}	Carga factorial rotada, cuadrática y normalizada del indicador i para el componente j
x_{it}	Valor sin normalizar del indicador i para un tiempo t
$\overline{x_i}$	Media de los valores sin normalizar del indicador i
w_{ir}	Peso asignado al indicador i para un método de agregación r
y_{it}	Valor del indicador i para un tiempo t tras la transformación de escala
Z_{jt}	Valor del componente j para un tiempo t obtenido por PCA

PARTE I

MEMORIA



Capítulo 1

Introducción

Antes de entrar propiamente en materia, resulta esencial situar el contexto del presente trabajo para posteriormente introducir la figura de los indicadores, de forma que el lector se dote de unas herramientas fundamentales para comprender la parte más técnica del proyecto.

1.1. La pandemia y el mercado laboral español

La pandemia de coronavirus supuso un antes y un después para todo el planeta en su conjunto. Generaciones que nunca habían vivido una situación de emergencia nacional afrontaron sus peores días en un contexto marcado por la incertidumbre y el confinamiento total. Los estragos causados por la enfermedad detuvieron el mundo y la práctica totalidad de la actividad económica en España. Muchos aspectos de la vida humana se vieron interrumpidos mientras la sociedad abordaba la mayor de las problemáticas: su salud. Esta interrupción, si bien repentina y casi absoluta, fue impulsada por los agentes socioeconómicos. De hecho, desde el propio Gobierno se incentivó la “hibernación” de la economía y toda actividad no esencial se vio obligada a parar su producción (Carlos E. Cué, 2022).

La interrupción de toda actividad económica implicó importantes consecuencias para el mercado laboral, que también se vio comprometido por la situación sanitaria. De hecho, la Cámara de Comercio de España (2021) ya anunciaba en su momento la significativa debilidad del mercado laboral durante los últimos meses de 2020, cuyos datos presentaban preocupantes tendencias de desempleo. Aun así, el análisis de los mismos resultaba y sigue resultando harto complejo, ya que los agentes públicos y privados trataron de mitigar el impacto sobre el mercado laboral a través de varias iniciativas, tales como el fomento de los ERTes o el amparo económico de los créditos ICO. De hecho, todavía hoy en día siguen presentes en nuestra sociedad algunas de sus consecuencias (Alejandro Martínez Vélez, 2022). De esta manera, se realizaron importantes esfuerzos por salvaguardar el tejido productivo español, para que así el mercado laboral no se viera damnificado como resultado.

Por ello, la monitorización del mercado laboral, tanto en años de pandemia como en un momento de normalidad, es compleja. Multitud de factores y agentes influyen sobre este fenómeno, que no debe ser caracterizado únicamente en torno a una o dos variables sobre el desempleo. De hecho, durante la pandemia no se disponía de ninguna herramienta accesible ni transparente para evaluarlo. Incluso aunque algunas fuentes proporcionaran algunos datos de empleo o de costes laborales, estos se encontraban dispersos y, a menudo, incompletos. Por supuesto, también se carecía de un indicador agregado que permitiera monitorizar el fenómeno

en su conjunto.

Para paliar esta problemática, 360 Smart Vision (2021) surgió como una plataforma tecnológica independiente impulsada por la Universidad Pontificia de Comillas y Deloitte, donde se proporciona acceso libre y transparente a indicadores de la realidad socioeconómica en España. En forma de panel interactivo, se presentan multitud de indicadores sobre los ámbitos de sanidad, movilidad, mercado laboral, consumo, mercado financiero, actividad empresarial, ciencia ciudadana y sostenibilidad e inclusión. De esta manera, se aborda uno de los principales vacíos existentes para la evaluación de la recuperación socioeconómica española. Este panel interactivo tiene por objetivo promover el análisis del impacto de las iniciativas públicas y privadas, de manera que todos los agentes sociales dispongan de las herramientas adecuadas para evaluar sus acciones y cómo estas repercuten en el beneficio social.

Sin embargo, si bien la plataforma proporciona una inmejorable cantidad de información sobre el mercado laboral, esta se encuentra totalmente desagregada en la actualidad. Por ello, el análisis del fenómeno descrito requiere del estudio simultáneo de demasiados indicadores, lo cual puede desincentivar a ciertos agentes de tratar de comprender el verdadero alcance de ciertas iniciativas. Al fin y al cabo, tanta información puede hacer que el público, el empresario o incluso el político se vean obligados a desechar ciertos indicadores –a sus ojos irrelevantes– para evaluar lo que cada uno considera como el mercado de trabajo. Esto, por tanto, puede llevar a cualquier interesado a tomar decisiones erradas a la hora de analizar el entorno laboral.

Por todo ello, es necesaria la existencia de un indicador sintético que agrupe todos los indicadores individuales presentados por la plataforma 360 Smart Vision en lo concerniente al mercado laboral español. De esta manera, se dotaría a todos los agentes sociales de una verdadera herramienta que, en conjunción con los KPIs originales, permita monitorizar de forma efectiva y transparente el estado del mercado laboral español. Por ello, se propone la elaboración de un indicador sintético sobre este fenómeno.

1.2. Estado de la cuestión

Los indicadores sintéticos son un tipo especial de indicadores que están formados típicamente por indicadores componentes, los cuales a su vez son variables directamente mensurables tales como la temperatura meteorológica, el paro, la inversión en sanidad e incluso la cantidad de goles marcados por Borja Bastón en el Real Oviedo. Por supuesto, incluso los indicadores más simples pueden medir algo con mayor o menor eficacia (Burgass *et al.*, 2017), dependiendo de su método de medición, su fiabilidad, su transparencia... Sin embargo, cuando la realidad que se busca medir no es medible por un solo valor (como, por ejemplo, la situación económica), varios indicadores pueden aportar su granito de arena para conformar lo denominado como *indicador sintético*, que vendría a reflejar dicha realidad compleja haciendo uso de la información aportada por los indicadores “componentes”. En ocasiones, algunos sintéticos se forman a partir de otros subindicadores sintéticos que sí están compuestos por indicadores componentes individuales.

La construcción de indicadores sintéticos es algo muy estudiado y contrastado. Al fin y al cabo, un indicador sintético no solo agrupa información, sino que otorga una visión general de un fenómeno no mensurable. Esta estructuración y presentación de la información siempre ha generado valor *per se*, lo cual se acentúa especialmente en el campo de los indicadores con aplicaciones sociales, donde hay una gran cantidad de datos desagregados cuyo análisis resulta

inviabile por separado. Además, dado su impacto multidisciplinar, su presencia se extiende por multitud de áreas, tanto en la vida pública como en la privada. En este sentido, resulta reseñable cómo “la tentación de las partes por sintetizar un proceso complejo y, en ocasiones, ambiguo (*i.e.* la sostenibilidad) en una simple cifra para comparar el desempeño nacional respecto de la aplicación de políticas parece irresistible” (Saisana *et al.*, 2005, p. 308).

No obstante, la complejidad de los mismos aumenta enormemente cuando se trata de describir el entorno en que se encuentran. Por ello, tanto empresas privadas como organismos públicos de todas las esferas han contribuido con el paso de los años a estandarizar los mecanismos que permiten la construcción de indicadores sintéticos con fiabilidad y trazabilidad. En este sentido, la OCDE (2008) ha preparado un manual que aspira a establecer un marco común en el proceso de creación de indicadores sintéticos. Así, no solo aquellos deseos de elaborar uno tendrán a su alcance los pasos bien detallados, sino que los interesados en comprender la lógica subyacente podrán hacerse con un manual que les ahorre mucho tiempo y quebraderos de cabeza.

Sin embargo, si bien dicho documento se considera un buen punto de comienzo para comprender ciertos aspectos fundamentales, la elaboración de un indicador completo requiere de una mayor profundización en cada una de las etapas descritas en el propio manual. Por ello, también se habrá de analizar el estado de la técnica aludiendo a los diferentes momentos de la construcción del medidor, tales como la recolección de datos, el análisis multivariante, la ponderación... No obstante, en primer lugar habrá de comprenderse bien su definición y el porqué de su importancia.

Los indicadores sintéticos, definidos por autores como Freudenberg (2003) como aquellos contruidos a partir de múltiples indicadores individuales (de ahora en adelante, indicadores componentes), son útiles para medir variables multidimensionales difíciles de acotar con un solo medidor. La mayoría de autores coinciden en la agregación de indicadores simples como elemento vertebrador de cualquier sintético, aunque para algunos resulta indispensable la existencia de un sistema complejo que, a partir de un modelo, permita comprender una realidad multidimensional (Nardo *et al.*, 2005). Por ejemplo, ¿cómo se podría medir el bienestar? A partir de indicadores componentes sobre la capacidad económica, los servicios a los que se tiene acceso, el nivel de educación, la salud se podría atinar a dar un número que, por sí solo, no significaría nada concreto; pero que serviría para establecer comparativas con otros países y consigo mismo.

De esta manera, los indicadores sintéticos, siguiendo un método de agregación aritmética, tendrían la siguiente forma matemática:

$$CI_t = \sum_{i=1}^n w_i I_{it} \quad (1)$$

CI_t : Valor del indicador sintético en un momento del momento t .

I_{it} : Valor normalizado del indicador i para un tiempo t .

w_i : Peso asociado al indicador i en un método de agregación aritmética

A lo largo del texto se presentarán otras fórmulas de acuerdo con la variedad de técnicas existentes referentes a la normalización, la ponderación y la agregación, así como el uso de otra simbología cuyo significado aparece en el prólogo Símbolos de la página 11 del documento.

Por todo ello, y recordando la cita de Saisana et al. (2005), los indicadores sintéticos gozan de una creciente importancia en el mundo académico y profesional, ya que aportan cierta monitorización a aspectos típicamente cualitativos y generalmente complejos. Su incipiente uso está documentado por numerosos autores, especialmente en la esfera pública y mediática (Saltelli, 2007). Sin embargo, pese a que el uso de estos medidores está muy extendido – siendo el paradigma el Índice de Desarrollo Humano (PNUD, 1990)–, también ha habido mucha discusión académica al respecto. Los detractores los critican por carecer de significado estadístico y por las arbitrariedades presentes en su elaboración (Grupp y Moge, 2004). En cambio, los partidarios destacan cómo los indicadores permiten poner el foco del debate en ciertas problemáticas sociales, a la par que facilitan la comprensión de realidades complejas. De hecho, estos argumentos lograron que Amartya San, Premio Nobel de Economía en 1998, cambiara su postura crítica. El economista comenzó a apoyarlos debido al impacto social propiciado por el debate en torno al IDH, tanto por atraer la atención pública como por dinamizar la toma de decisiones de entidades públicas y privadas (Greco *et al.*, 2019).

Así pues, los indicadores sintéticos también son especialmente útiles al reducir un gran número de medidores iniciales, lo cual además los convierte en una herramienta especialmente didáctica para aquellos poco versados en el entorno observado por el indicador. Al fin y al cabo, resulta más sencillo comparar los mercados productivos de dos países a partir de un indicador sintético llamado *Producción* que con los medidores IPRIX, IPRIM, IPRI... los cuales resultan complejos fuera del entorno industrial. Sin embargo, también suponen un riesgo notable para las instituciones allí donde son construidos negligentemente (Cherchye *et al.*, 2007). Para evitarlo, el proceso de elaboración ha de ser lo más transparente y justificado posible, de forma que el análisis del entorno proporcionado por el indicador no induzca a errores o manipulaciones. A tal efecto, la trazabilidad de sus medidores componentes es esencial.

Por todo lo expresado, queda patente que los indicadores sintéticos tienen un enorme potencial de transformación social, pero también pueden ser engañosos si su elaboración no es minuciosa. A continuación, se presenta la tabla 1 sobre las oportunidades y los inconvenientes que presentan los indicadores sintéticos, basada en los hallazgos de Saisana y Tarantola (2002):

Indicadores sintéticos	
Oportunidades	Inconvenientes
Sintetiza y esclarece problemas multidimensionales	Puede arrojar mensajes engañosos
Proporciona una imagen holística de la realidad compleja	Depende de la interpretación correcta del usuario
Genera concienciación social en el público y sus legisladores	Se elabora en torno a decisiones a veces arbitrarias o subjetivas
Estimula la implantación de políticas y contribuye a cuantificar su impacto	Exige una gran cantidad de datos y su fiabilidad

Tabla 1. Oportunidades e inconvenientes de un indicador sintético cualquiera

El mundo académico ha sido muy prolijo en la elaboración de indicadores medioambientales, los cuales son además especialmente complejos por reflejar realidades multidimensionales.

Para realizar un diseño útil para todos los actores involucrados, este habrá de gozar de ciertas características (Burgass *et al.*, 2017): ser rentable en términos de tiempo y dinero, proporcionar información fiable, ser relevante y responder acordemente a lo que refleja. Su capacidad de influir sobre el público y los políticos los hace una herramienta de transformación social, propósito perseguido por los siguientes indicadores:

- ***Ocean Health Index (OHI, Índice de Salud de los Océanos)***: Los autores Halpern *et al.* (2012) elaboraron un índice que representa el estado del ecosistema acuático en base a diez subindicadores (también sintéticos) que reflejan distintos objetivos que los legisladores deben cuidar. Así, un set inicial de indicadores se agrupó en torno a metas legislativas como la limpieza de agua, la biodiversidad, la economía costera, la disponibilidad de productos naturales... Estos subindicadores, descriptores de realidades económicas, sociales, químicas... son posteriormente agregados en el índice sintético final. Resulta de especial interés la agrupación según objetivos, lo que en cierta manera recuerda a las áreas descritas por el panel interactivo 360 Smart Vision (Deloitte, 2021), en tanto que cada una puede ser abordada por los *stakeholders* de forma independiente. Esto repercutiría acordemente sobre el indicador sintético final, según la importancia de dicho categoría concreta.
- ***Environmental Performance Index (EPI, Índice de Desempeño Medioambiental)***: Evalúa el progreso general de los países respecto de su desarrollo sostenible (Hsu y Zomer, 2016). En su haber se encuentran multitud de indicadores institucionales, socio-económicos, medioambientales... puesto que todos gozan de una influencia significativa sobre el desarrollo sostenible. Su razón de ser persigue el fortalecimiento del análisis técnico de la sostenibilidad para dar lugar a decisiones más acertadas y concienciar a la población. Al igual que el anterior, este indicador vuelve a tomar una estructura anidada de dos categorías (vitalidad medioambiental y salud medioambiental), siendo cada una un subindicador de varios indicadores individuales.

De acuerdo con los ejemplos y la práctica asentada en el sector, los indicadores sintéticos suelen tomar la estructura anidada representada por la figura 1, la cual ha de ser definida en el marco teórico del indicador en cuestión:

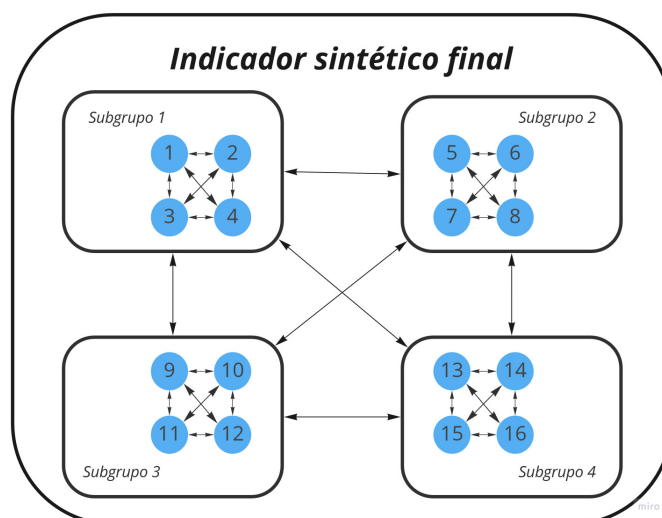


Figura 1. Representación de la estructura de los indicadores sintéticos

Estos ejemplos permiten comprender la estructura típica que siguen los indicadores sintéticos, donde básicamente se realizan las técnicas de normalización, ponderación y agregación preliminares en categorías relacionadas o aspiracionales, las cuales vendrían a ser los subgrupos y conformarían subindicadores sintéticos. Finalmente, estas técnicas se reevalúan y se aplican de nuevo para agregar dichos subgrupos en el indicador sintético final. Por supuesto, cada proceso de construcción es único, puesto que ha de evaluarse detenidamente qué realidad se busca medir, junto a la calidad y cantidad de los datos de que se disponen. Por ello, la toma de decisiones es el eje central de toda elaboración de indicadores sintéticos y, de hecho, donde ocurren las arbitrariedades achacadas por sus detractores (Grupp y Moge, 2004).

Para minimizar esto, numerosos organismos han tratado de proporcionar fiabilidad y establecer un marco repleto de propuestas e hitos para avanzar en la elaboración de los mismos. Sin embargo, la toma de decisiones no puede ser desterrada, ya que la elección de ciertas técnicas no siempre responde o puede responder a información técnica. Esto ya se menciona en el propio manual de la OCDE (2008) que engloba artículos de organismos tales como la Comisión Europea, donde se resalta la importancia de las decisiones, sin proporcionar tampoco ninguna regla de oro que solucione la cuestión. En cambio, encomiendan a la comunidad académica la crítica de las técnicas empleadas y el marco general creado para el indicador.

Este manual, no obstante, sí establece una serie de pasos que contribuyen a esclarecer el proceso de elaboración, mas no la toma de decisiones:

- **Marco teórico general:** Cada decisión tomada a lo largo del proyecto ha de ser debidamente justificada con base en las experiencias previas de otros autores. A pesar de que no existe un indicador idéntico a otro, ciertas características comunes permitirán el establecimiento de un marco general que defina los conceptos estudiados y los posible subgrupos que conformarán el sintético final. Además, cada paso posterior, en tanto que involucra una toma de decisiones concreta, debe estar relacionado con el marco teórico establecido al comienzo del proceso, de forma que el sintético se alinee finalmente con los objetivos con que fue propuesto.
- **Recolección de datos:** Los indicadores componentes han de seleccionarse según unos criterios de calidad, tales como la relevancia, la utilidad y la consistencia en el ámbito de medición (en este caso, nacional). Los indicadores componentes son facilitados por el panel interactivo 360 Smart Vision (Deloitte, 2021). Asimismo, los datos están presentados de forma que la recolección permite desechar los indeseados y estructurar los restantes fácilmente.
- **Imputación de información ausente:** El panel interactivo 360 Smart Vision proporciona indicadores con escalas de tiempo muy dispares entre sí. Por lo tanto, cuando se trata de homogeneizar la estructura de datos para analizarlos y realizar los siguientes pasos de la construcción, resultará apropiado estudiar el empleo de ciertas técnicas de imputación. En primer lugar, habrá de analizarse qué tipo de ausencias se dan y a qué se deben estas, ya que de ello dependerá el abanico de métodos que se podrán usar (Donders *et al.*, 2006). Según la literatura, existen los patrones MCAR (missing completely at random o ausentes totalmente al azar), MAR (missing at random o ausentes al azar) y NMAR (not missing at random o no ausentes al azar). Una vez determinado el patrón, existen muchos métodos

de imputación tales como técnicas regresivas, la esperanza-maximización, la imputación múltiple... Si bien el estado de la cuestión está muy avanzado, cada método cuenta con sus virtudes y defectos.

- **Análisis multivariante:** El análisis multivariante proporciona al investigador información relevante acerca de la estructura de datos que maneja. Este análisis resulta indispensable para comprender las correlaciones entre indicadores y, por tanto, ser acertados a la hora de ponderarlos y agregarlos más adelante (OCDE, 2008). El estado de la técnica está también muy desarrollado, ya que existen multitud de métodos para profundizar tanto como se desee, tales como el análisis de componentes principales (PCA), la técnica ANOVA, el análisis *cluster* o el método de regresiones auxiliares.
- **Normalización:** La normalización se basa en reducir todos los indicadores a una escala común (Gan *et al.*, 2017). Esta parte resulta indispensable para los posteriores pasos de ponderación y agregación, ya que en un comienzo las escalas de los diferentes medidores son totalmente dispares. Después de todo, los indicadores componentes son en ocasiones tasas de variación interanual y en otros valores absolutos formados por varias cifras. En estas circunstancias, ponderar escalas tan diferentes haría del proceso de ponderación algo muy complejo y tedioso. Para solucionarlo, diversos mecanismos ya han sido analizados extensamente en diferentes ámbitos profesionales, siendo algunos ejemplos el método de puntuaciones (*ranking*), la técnica de estandarización, el método Mín-Máx...
- **Ponderación y agregación:** Este paso resulta clave a la hora de elaborar el indicador, ya que se toman las principales decisiones en cuanto a la priorización de unos medidores componentes sobre otros. Esto debe hacerse teniendo en cuenta la literatura existente, así como comprobando posteriormente la robustez y utilidad del indicador creado según los diferentes métodos empleados. De esta manera, el autor podrá identificar mediante datos empíricos qué decisiones han sido erróneas en la elaboración y, si fuera posible, corregirlas a tiempo. Según Salvatore Greco (2019), hay una gran variedad de métodos de ponderación: EW (equal weighting o ponderación equilibrada), UCM (unobserved components models o modelos de componentes no observados) o AHP (analytic hierarchy processes o procesos de jerarquía analítica), entre otros. Considerando la agregación, hay técnicas tanto compensatorias como no compensatorias, además de mixtas (Gan *et al.*, 2017), muy bien documentadas y analizadas por multitud de autores.

Posteriormente, cabría realizar la etapa final que consiste en estudiar la robustez del resultado obtenido, ya sea a través de estudios de incertidumbre o de sensibilidad. De esta manera, el indicador sintético se pondrá a prueba, tratando de evaluar el impacto de cada uno de los pasos realizados y observar la trazabilidad de los indicadores componentes. Los análisis de robustez son muy frecuentes en la estadística, por lo que hay abundantes referencias al respecto.

1.3. Motivación

Por todo lo expuesto, queda claro que el estado de la técnica está muy avanzado en lo referente a la construcción del indicador. Cada uno de los pasos involucrados en la misma está

profundamente detallado por la literatura existente, puesto que además todas las técnicas se pueden aplicar a otras labores estadísticas. Sin embargo, dada la particularidad del indicador sintético en cuestión, muchas decisiones habrán de individualizarse basándose en situaciones análogas de acuerdo con el juicio del investigador. No hay indicadores idénticos entre sí, por lo que la construcción de cada uno de ellos reviste de especial complejidad.

Dicha dificultad representa un reto atractivo, puesto que hoy en día España carece de un indicador de la recuperación socioeconómica. En un mundo tan damnificado por la pandemia de COVID-19, la inexistencia de un medidor que aporte una visión global, actualizada y también histórica de la recuperación parece un vacío de información impropio de un mundo digital como el nuestro. Por lo tanto, este trabajo tiene como motivación principal aportar valor a la sociedad española mediante dicho indicador sintético, de manera que la construcción aporte también transparencia y pedagogía a cualquiera que busque información para elaborar futuros medidores.

Asimismo, dado que el proceso para elaborarlo requiere de decisiones activas por parte del autor, la abundante documentación servirá como una base sólida sobre la que construir el indicador. La tarea resultará ardua y compleja, ya que la constante elección de métodos a cada uno de los pasos de la construcción ha de ser sólida y justificada. De esta manera, se garantiza la total comprensión del proceso de elaboración por parte del autor. Además, será fundamental un correcto manejo de las soluciones tecnológicas que permiten trabajar con los datos, así como el aprendizaje de muchos procesos y técnicas pertenecientes a diferentes ramas de la estadística.

Por último, la elaboración del indicador sintético supone un atractivo muy significativo para el entorno del panel interactivo 360 Smart Vision. La aparición de un indicador que compare históricamente la evolución de la recuperación, especialmente centrada en el mercado laboral, aumentaría la capacidad del panel de ser una herramienta divulgadora eficaz para la comunidad científica y empresarial.

Por todo ello, este proyecto se revela como un reto y, al mismo tiempo, como una solución tangible para un problema muy actual. En el proceso, se pretende interiorizar muchas técnicas de análisis y profundizar en aquellas ya conocidas gracias a los años de estudio de Ingeniería Industrial y ADE. Así pues, esta investigación es una oportunidad que aportará valor a la sociedad y que permitirá, tanto al autor como a futuros investigadores, dotarse de valiosas herramientas a la hora de analizar otros indicadores sintéticos y el entorno que estos describen.

1.4. Objetivos

Los objetivos del proyecto son los siguientes:

- **Diseño conceptual del indicador sintético (objetivo 1):** Cada paso requerirá de la elaboración de un marco teórico que justifique la implementación posterior. De esta manera, este objetivo alude a la continua revisión y crítica de la bibliografía que permita tomar las decisiones sobre qué técnicas usar a la hora de construir el indicador. Para completar esta meta con eficacia, habrán de llevarse a cabo acciones como el entendimiento del marco genérico de construcción propuesto por la OCDE (2008), la búsqueda de artículos específicos de cada método empleado y la revisión de cada una de las decisiones tomadas una vez sus consecuencias han sido medidas en la práctica. Así, se

creará un marco teórico sólido particularizado para la elaboración del indicador sintético estudiado.

- **Implementación efectiva del indicador (objetivo 2):** A nivel nacional no existe precedente de un indicador que establezca una comparativa histórica del proceso de recuperación del mercado laboral en un contexto pospandémico. A partir del marco referencial establecido gracias a la literatura, este objetivo se centra en la puesta en práctica de las decisiones llevadas a cabo en dicha revisión bibliográfica. La programación de las técnicas de normalización, ponderación o agregación, además del análisis inicial de la estructura de datos proporcionada por 360 Smart Vision, permitirá la construcción del agregador sintético. Asimismo, también se propone una subagregación de los indicadores por áreas que, además de aunarse posteriormente para dar lugar al medidor sintético final, también arrojen luz sobre cada área por separado. Esto proporcionará valor adicional al proyecto.

Los objetivos enunciados configuran las metas principales del proyecto. No obstante, tras la obtención de resultados, también se promoverá su divulgación y la transformación social que pueda derivar de ellos. Actores tales como la comunidad científica, los organismos públicos o ciertas entidades privadas encontrarán gran utilidad en el indicador construido en tanto que les permite evaluar su realidad de forma más efectiva. Asimismo, dado que algunos de los indicadores componentes se encuentran alineados con los Objetivos de Desarrollo Sostenible (ODS), se buscará potenciar la concienciación social sobre los mismos. A tal efecto, se presenta el anexo A. Por tanto, para cumplir estos propósitos ha de garantizarse el carácter didáctico del informe, así como llevar a cabo una presentación ilustrativa en el panel interactivo 360 Smart Vision en pos de su integración con el resto de la plataforma. Finalmente, dados los estándares de calidad establecidos en la construcción del sintético, se explorarán *a posteriori* otras vías como la publicación de un artículo académico que asiente el proyecto como una referencia para futuros trabajos.

Estas metas persiguen la creación de valor social y personal. Como objetivo general, el proyecto aspira a proporcionar una herramienta útil y transparente que permita analizar el entorno laboral en España, a la vez que dinamiza la toma de decisiones de los actores interesados y arroja luz sobre la construcción de indicadores sintéticos.

1.5. Recursos

Los recursos necesarios están disponibles, ya que la Universidad Pontificia de Comillas proporciona licencias para el programa MATLAB y la herramienta IBM SPSS. Además, también se empleará el programa Gretl de forma puntual, el cual es de código abierto. De esta manera, contando además con el acceso a las bibliotecas virtuales de muchas instituciones, no parece que se den limitaciones en cuanto al uso de recursos. La plataforma 360 Smart Vision proporcionará una base de datos que cuenta con información tanto pública (datos referentes al mercado de trabajo provenientes del INE, por ejemplo) como reservada, la cual es aportada en muchos casos por entidades privadas. Cabe destacar que este panel forma parte del Observatorio Critería, una iniciativa surgida de prestigiosas instituciones y que aspira a ayudar a la sociedad española a conocer y comprender los efectos de los principales retos demográficos

y socioeconómicos (CRITERIA, 2022). Los datos del panel constituyen, por tanto, información de primera calidad.

Estos datos son a veces tan cuantiosos que requieren de cierta capacidad computacional para su procesamiento. Para ello, el PC de que se dispone se presenta en la tabla 2:

Recursos			
Modelo de PC	Vivobook S-14	Sistema operativo	Windows 10 Pro
Procesador	AMD Ryzen 7 4700U	Tarjeta Gráfica	AMD Radeon Graphics
Disco Duro	512 GB SSD	Memoria RAM	16 GB

Tabla 2. Resumen de las características del ordenador empleado en el proyecto

1.6. Metodología

La metodología de trabajo seguirá el marco explicitado en el apartado del estado de la cuestión, a saber: marco teórico general, recolección de datos, imputación de información ausente, análisis multivariante, normalización, ponderación y agregación y posterior revisión del método según la robustez y la sensibilidad del indicador construido. Tras concluir estas etapas, solo restaría la divulgación del proyecto. Para poder llevar a buen término el proyecto descrito, se ha establecido una serie de pasos que habrán de cumplirse. Esta secuencia se presenta en forma de diagrama de flujo, el cual es reflejado por la figura 2 presentada a continuación:

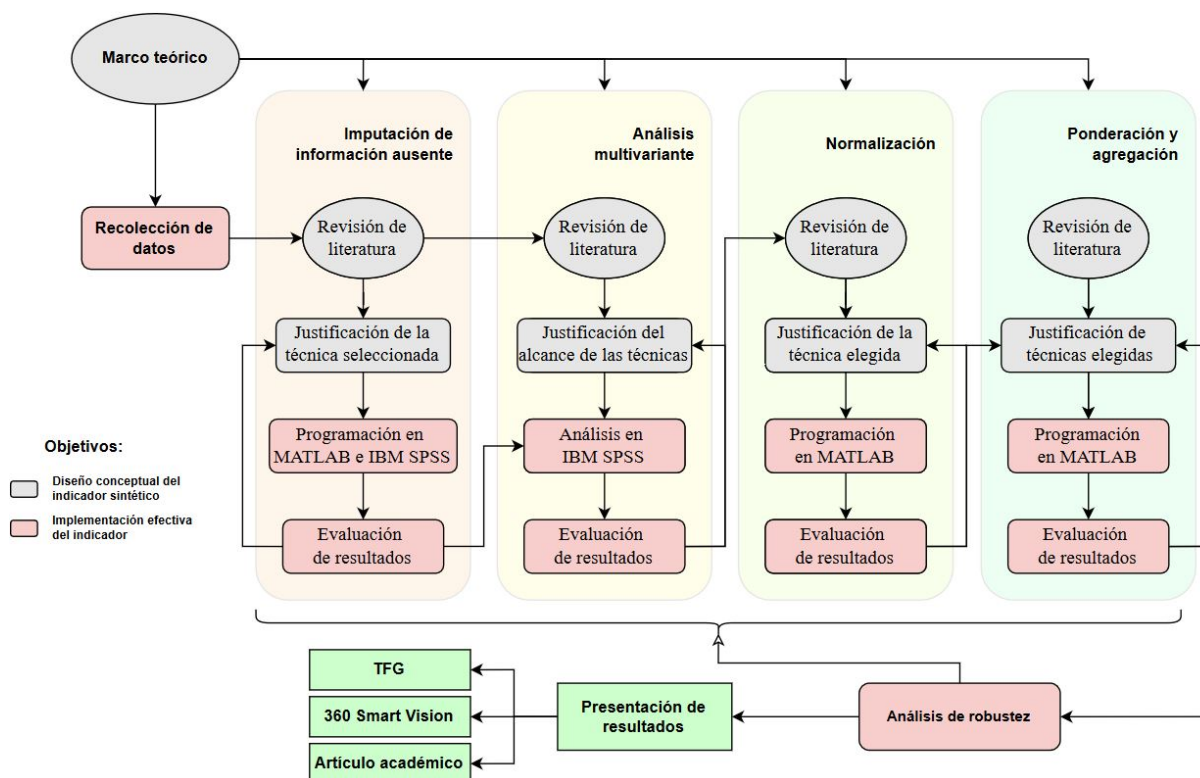


Figura 2. Diagrama de flujo sobre el proceso de construcción del indicador sintético y sus objetivos

La construcción del marco teórico y la recolección de datos son procesos estrechamente relacionados, ya que uno se basa en el estudio del estado de la cuestión y el establecimiento de una estructura para el indicador y el otro, en cambio, busca aplicar los criterios asentados para elegir las variables que dotarán de información al indicador sintético final. Por supuesto, también será importante disponer de una fuente de datos suficientemente extensa como para aplicar dichos criterios sin perjudicar la cantidad de información disponible.

En cuanto a otras tareas de carácter simultáneo, las revisiones de bibliografía de la imputación de información ausente y del análisis multivariante son consideradas como tal, ya que ninguna involucra una toma de decisiones que suponga cambios para la otra. Pese a ello, la implementación del análisis sí requiere la finalización de la imputación; ya que, en caso contrario, faltarían los datos ausentes y los resultados estarían sesgados. Tras finalizar estos pasos, las revisiones de literatura de la normalización, la ponderación y la agregación darán comienzo. La técnica empleada para ponderar no debería depender de los resultados de la normalización solo la importancia relativa de cada indicador componente, pero resulta prudente espaciar temporalmente las fases críticas de la elaboración.

Finalmente, cabe resaltar que las evaluaciones de resultados siempre pueden conducir a un replanteamiento de las técnicas empleadas en cada caso, ya que las conclusiones podrían revelar alguna decisión errónea. Estas evaluaciones son esenciales para garantizar la transparencia y la fiabilidad del indicador. Asimismo, una vez se finalice la construcción del indicador sintético, el estudio de robustez y sensibilidad podría retocar cualquiera de los mecanismos presentes en la elaboración del medidor, en tanto que este identifica las principales fuentes de incertidumbre a lo largo de la construcción del sintético (Cherchye *et al.*, 2008). Tras la obtención y el análisis de los resultados, la presentación de los mismos a través de distintos canales es importante para impulsar el impacto del proyecto, por lo que algunos de ellos son posteriores a la elaboración del presente texto. El diagrama de flujo de la 2 aporta leyenda de color para facilitar la identificación de los objetivos explicados en el presente capítulo, donde la presentación de resultados se muestra diferenciada por su índole posterior, en cierta manera excediendo los límites del presente proyecto.

Las tareas descritas en dicha figura se agrupan en torno a los dos objetivos principales, aunque algunas de ellas, tales como el establecimiento del marco teórico, participan de ambos. Los dos objetivos planteados por el capítulo y reflejados en la figura 2 definirán los dos próximos capítulos. En primer lugar, el capítulo 2 profundizará y recopilará las técnicas elaboradas en el campo de la estadística, de modo que se presente un análisis contrastado y actualizado de las ventajas y desventajas que presenta cada técnica. Se tratará, por tanto, de un proceso de elaboración conceptual y teórico, que sirva de guía a futuros autores. Posteriormente, el capítulo 3 versará sobre la construcción efectiva del indicador sintético del proyecto. En él se aplicarán las enseñanzas halladas en el capítulo 2 y se obtendrán, por tanto, los resultados buscados. Su obtención dará por finalizado el proyecto.

De esta manera, la metodología del proyecto ha sido debidamente explicada, reflejando fielmente el proceso a través del cual se elaborará el indicador y su marco conceptual. De esta manera, tal y como se ha descrito a lo largo del capítulo, el primer paso en la creación del sintético será establecer un marco teórico general que, a modo de acercamiento conceptual, permita asentar la bibliografía existente y comprender las principales técnicas aplicables a lo largo del proceso.

Capítulo 2

Diseño conceptual de un indicador sintético

EL marco teórico general sobre el que se asienta el proyecto consta de un análisis pormenorizado de lo que comúnmente se conoce como *estado de la cuestión*, con la salvedad de que en este caso el estado de la cuestión no se referirá a la existencia genérica de indicadores sintéticos, sino principalmente a cada una de las etapas presentes en la elaboración de los mismos. De esta manera, se presentarán las principales técnicas y ventajas de aquellos métodos que resulten plausibles a la hora de configurar un indicador sintético. Además de sentar los pilares del proyecto, este análisis comparado otorgará a futuros investigadores una sólida base sobre la que argumentar sus decisiones, por lo que este estado de la cuestión goza de una importancia sustancial.

Por estos este motivo, este capítulo abordará el primero de los objetivos del proyecto, que versaba sobre el diseño conceptual del indicador sintético, de manera que se proporcione a próximos autores una crítica holística de la bibliografía existente y un manual de aplicación de la misma que les permita adaptarla a sus propios trabajos. Por tanto, se reunirán también mecanismos, procedimientos y técnicas cuya utilización excede la propia de los indicadores sintéticos, pero que eventualmente pueden verse involucrados en su elaboración.

Dada la abundante y variada bibliografía al respecto, esta colección pormenorizada gozará de especial significación, en tanto que supondrá una revisión actualizada de las técnicas y las conclusiones que diferentes autores han ido aportando a lo largo de los años. Su realización constituye, por ello, uno de los objetivos del proyecto.

2.1. Marco teórico

“Lo que está mal definido está, probablemente, mal medido.”
(OCDE, 2008, p. 22)

El marco teórico en que se asienta un indicador alude, *grosso modo*, a la estructura y objetivo que le otorga el autor. Aunque posteriormente se presentará el marco teórico concreto que definirá la estructura del indicador del proyecto, ha de establecerse en primer lugar el estado de la cuestión sobre el mismo. La bibliografía no es precisamente prolija (Burgass *et al.*, 2017), debido principalmente a la dificultad de establecer criterios comunes a la creación de una estructura totalmente dependiente del indicador que se desea construir. En este sentido,

no solo influirá la propia estructura de los datos, sino también las intenciones del autor en su elaboración; puesto que estas son, al final, las que orientan el marco que seguirá el indicador sintético. No obstante, sí es cierto que existen unas líneas generales que todo marco teórico ha de seguir, a la par que ha de abordar ciertos preceptos que salvaguarden la calidad del indicador.

Como ya se ha mencionado, autoridades europeas del campo de la Estadística (Eurostat, 2017) establecen una serie de dimensiones de calidad que todo resultado estadístico ha de preservar:

- *Pertinencia:* Este parámetro representa la capacidad del indicador sintético de tener impacto sobre los usuarios, legisladores o cualquier otra parte involucrada. Por lo general, la mayoría de sintéticos nacen debido a una necesidad social de mostrar una realidad harto compleja, de forma que esta sea medible y cercana para el público (Freudenberg, 2003). El Código de Buenas Prácticas de las Estadísticas Europeas (2017) también sugiere que han de establecerse mecanismos de consulta y seguimiento para garantizar la pertinencia de la producción estadística en el futuro, anticipándose a las necesidades de los usuarios.
- *Precisión y fiabilidad:* Este precepto alude a la exactitud con que los datos reflejan la realidad que describen. A tal efecto, Fabio Santeramo (2015) ya indica cuán importante es la selección de datos y su calidad a la hora de lograr un sintético de calidad. Este parámetro goza de especial importancia a la hora de elaborar un indicador sintético; puesto que, además de las imperfecciones debidas a la calidad de los datos originales, cada técnica de imputación, normalización o ponderación añade nuevas fuentes de incertidumbre. En aras de medir esto, la etapa final de análisis de sensibilidad debe tratar de cuantificar cuán trazable y preciso es el indicador construido, reflejando por tanto cualquier error detectado.
- *Oportunidad y puntualidad:* La cualidad de ser oportuno se refiere al periodo de tiempo que transcurre entre la medición de los datos y la realidad que describen. La puntualidad, en cambio, representa cuán cercana es la fecha prevista de presentación de resultados y la de entrega efectiva. En cuanto a la elaboración de indicadores, el primer precepto resulta más relevante, ya que la utilidad de los mismos depende en buena parte de la actualidad de sus datos.
- *Accesibilidad y claridad:* Esta dimensión aborda la facilidad con que un usuario corriente puede acceder y comprender fácilmente los datos. Para lograrlo, distintos metadatos deben acompañar la presentación de resultados a través de también diferentes canales de comunicación. Así, se fomenta el alcance del proyecto y su impacto en el público general, lo cual representa generalmente uno de los objetivos de los indicadores sintéticos.
- *Coherencia y comparabilidad:* Estos preceptos aluden a la consistencia de la producción estadística a lo largo del tiempo. La coherencia representa cómo una diferente combinación de los datos originales seguiría reflejando la realidad, de forma que ha de contrastarse especialmente aquellos que provengan de fuentes de información diferentes, garantizando la fiabilidad de las mismas. La comparabilidad representa el análisis de impacto que cada variable o procedimiento tendría sobre el resultado estadístico, lo cual es medido en el

análisis de robustez y sensibilidad en el caso de los indicadores. Además, también aborda la comparación entre países (en el caso de indicadores nacionales comparativos, como el Índice de Desarrollo Humano o el *Technology Achievement Index* (Guz *et al.*, 2017)) o en una determinada serie temporal.

Si bien la satisfacción de estos parámetros atañe a todas las fases de la elaboración, el marco teórico es, tanto por su prelación en la toma de decisiones como por su importancia, la fase del proyecto que más influirá en la consecución de estas dimensiones de calidad. Al fin y al cabo, el marco teórico debe definir los conceptos involucrados en el indicador y proponer qué subgrupos de variables, de existir, lo compondrán (Burgass *et al.*, 2017). Esto suele ser necesario para la mayoría de sintéticos, en los cuales incluso se pueden establecer objetivos intermedios cuya situación puede ser cuantificada (Mazziotta y Pareto, 2017). En este sentido, cabe recordar que el ejemplo del *Ocean Health Index* (Halpern *et al.*, 2012) presentado en el epígrafe 1.2, donde los 10 subgrupos que formaban el sintético final eran indicadores que reflejaban objetivos por sí mismos: Limpieza del Agua, Biodiversidad, Turismo y Ocio, Recursos Naturales...

De esta manera, mediante el marco teórico quedará establecida la estructura del indicador, lo que afectará a las técnicas elegidas posteriormente y, por tanto, a todos los parámetros de calidad anteriormente mencionados. A este asunto se refieren por ejemplo Paruolo *et al.* (2013), quienes establecen la existencia de un coste de sobresimplificación debido a la propuesta de un marco teórico excesivamente simple, que conlleva métodos de ponderación y agregación también sencillos. Además, la falta de un marco teórico bien definido implica que, en muchas ocasiones, los autores no evalúen adecuadamente el alcance de las técnicas empleadas. Por todo ello, esta etapa de la elaboración del indicador sintético se torna fundamental, puesto que reconoce la naturaleza multidimensional del fenómeno a describir (Nardo *et al.*, 2005). De hecho, aunque en muchas ocasiones pueda parecer obvio, demostrar dicha multidimensionalidad habría de ser lo primero de todo proceso de construcción de un indicador sintético.

Por todo lo mencionado, y de acuerdo con el manual de la OCDE (2008), un buen marco teórico ha de determinar lo siguiente:

- **Definición de los conceptos:** Aquellos presentes en el indicador han de definirse apropiadamente, concretándolos tanto como sea posible y refiriéndolos a autoridades expertas en el tema, de forma que dichas definiciones estén debidamente contrastadas. También es vital para evitar cualquier tipo de ambigüedad o interpretación equivocada (Saltelli, 2007). Además, de este punto se podrían percibir fácilmente las relaciones entre los subgrupos que se establecerán después, ya que estos podrían formar parte de la definición del fenómeno medido por el sintético. Esto queda reflejado en la elaboración del Índice de Desarrollo Humano (PNUD, 1990), donde Amartya Sen y el resto de autores se ven obligados a comenzar definiendo el término *desarrollo humano*, ciertamente complejo y subjetivo.

Para los académicos involucrados, el desarrollo humano goza de dos vertientes: el cultivo de capacidades (conocimientos, salud, destrezas...) y la aplicación de dichas capacidades (descanso, producción, ocio...). Por lo tanto, la economía participaría del desarrollo humano, pero no lo conformaría en exclusiva (Haq, 1995). El crecimiento económico, como tal, se convierte sólo en un subconjunto del paradigma del desarrollo humano, de manera que ciertos parámetros sociales tomaran una importancia otrora desconocida. Así,

definieron el desarrollo en torno a tres elementos esenciales de la vida humana y vitales para estas dos vertientes: la longevidad (salud para cultivar y aplicar capacidades), el alfabetismo (formación de conocimientos) y el acceso a recursos (producción). Como ningún elemento mencionado es independiente en la vida humana, la elaboración de un indicador sintético es más que pertinente (Abell, 2010) para la comprensión del fenómeno del desarrollo humano. De hecho, tal es la complejidad del susodicho que el mismo Amartya Sen llegó a dedicar libros enteros para explorar filosóficamente el concepto (Sen, 1998).

- **Determinación de los subgrupos:** La existencia de subconjuntos es típica de los indicadores sintéticos, que tienden a presentar estructuras anidadas que facilitan la comprensión del usuario (Santeramo, 2015). Además, los propios subgrupos tienden a ser sintéticos intermedios (Mazziotta y Pareto, 2017) que son muy útiles para legisladores y otros actores involucrados. Así, en ocasiones ocurre que ciertos países tienen un desempeño muy bueno en ciertos subgrupos del sintético final (como en el IDH o en el TAI), pero son lastrados en otras categorías y, por ello, no presentan un valor muy positivo en el indicador final. Por lo tanto, la existencia de estos subgrupos y su divulgación permite al entorno público-privado identificar correctamente los problemas y monitorizar más fácilmente el impacto de las políticas implementadas. Básicamente, los subgrupos son ciertamente útiles, tanto por facilitar la comprensión del sintético final, como representar en sí mismos objetivos dignos de estudio y análisis.

Asimismo, habrán de describirse la posible interdependencia entre los subgrupos, lo cual será refrendado posteriormente por el análisis multivariante. Esto arrojará luz sobre posibles técnicas de ponderación y agregación, de manera que no se incurra en el mencionado coste de sobresimplificación mencionado por Greco *et al.* (2019).

- **Establecimiento de un criterio de selección de variables:** Típicamente, el marco teórico ha de establecer un criterio que permita posteriormente dilucidar qué variables han de tomarse y cuáles excluirse. A tal efecto, el objetivo del indicador y la definición del mismo resultan muy importantes. Las variables han de estar circunscritas al lugar concreto en el momento determinado en que el sintético pretende reflejar su realidad, ya que mezclar variables de ámbitos local y nacional puede complicar mucho la interpretabilidad del indicador. Por lo tanto, queda patente que todo criterio aplicado sea lo más concreto y esclarecedor posible (OCDE, 2008). En casos como el EPI (Hsu y Zomer, 2016), el OHI (Halpern *et al.*, 2012) o el SSI (Van de Kerk y Manuel, 2008) se aplican algunos criterios como la relevancia, la mensurabilidad, la accesibilidad, el método científico, la consistencia... de manera que indirectamente las variables alcancen cotas de fiabilidad más que satisfactorias.

De hecho, autores como Burgass *et al.* (2017) promueven la aplicación de los criterios de la calidad del Código de Buenas Prácticas de las Estadísticas Europeas (2017) a la hora de seleccionar las variables. Además, en aquellas ocasiones en las que la calidad de los datos sea cuestionable, las matrices de pedigrí presentadas por Van der Sluijs *et al.* (2005) podrían arrojar luz sobre la fiabilidad de los datos. Básicamente, un panel de expertos otorgaría calificaciones a la calidad de las variables (Van Der Sluijs *et al.*, 2002) en base

a parámetros tales como su método de medición, su representatividad, su validación...

Este método resulta muy útil cuando la calidad de las variables iniciales es falible; pero, si la base de datos con que se cuenta es extensa, esto puede llegar a consumir mucho tiempo en lo referente tanto a la planificación como a la ejecución a través del panel de expertos. Por lo tanto, los autores deberán evaluar si su realización merece la pena de acuerdo con la fiabilidad de las variables disponibles.

En resumen, el marco teórico establecido es indispensable a la hora de comprender el funcionamiento del indicador. La definición de los conceptos involucrados, tanto a nivel de sintético final como de los subgrupos o categorías, es indispensable para evitar toda ambigüedad técnica o lingüística. Además, la determinación de dichos subgrupos tiene impacto notable sobre las técnicas posteriores de ponderación y agregación; a la vez que otorga al usuario una comprensión de la influencia que cada variable tiene sobre el sintético final. Finalmente, el establecimiento de un criterio para seleccionar variables y desecharlas permite salvaguardar los preceptos propuestos por las autoridades europeas (Eurostat, 2017), de forma que se minimice tanto como sea posible la incertidumbre proveniente de la selección de datos.

En este sentido, la incertidumbre como fuente de error está presente a lo largo de la construcción de un indicador. Cada decisión o etapa del proceso es imperfecta y, por tanto, no refleja la realidad con exactitud (Burgass *et al.*, 2017). Cada variable empleada como fuente de información presenta cierta desviación respecto de lo que aspira a cuantificar, ya sea por errores en la medición, en la muestra estudiada, en el método empleado... Además, las técnicas posteriores de imputación, normalización, ponderación y agregación añaden otro grado más de incertidumbre que, aunque se pueda reducir tomando decisiones acertadas, nunca llegará a cero.

A este respecto, el trabajo de Burgass *et al.* (2017) es muy interesante, ya que diferencia la incertidumbre producida por cuatro fuentes distintas: el marco teórico, los datos originales, el proceso de construcción y, también, la divulgación final de resultados. Por tanto, todo autor debe comprender que cada decisión tomada implica una serie de desajustes respecto de la realidad medida y que, tal y como afirma Van Der Sluijs en sus reflexiones sobre la incertidumbre (2005), tratar de desterrar a este 'monstruo' resulta poco viable. En cambio, su asimilación a través del análisis y el estudio de la incertidumbre está a la orden del día en el mundo académico. Por ello, un pormenorizado análisis de robustez y sensibilidad resulta muy importante para cualquier indicador de reciente creación.

Aunque, como ya se ha mencionado, la bibliografía no es muy extensa en cuanto al marco teórico de los indicadores, el presente epígrafe ha descrito las líneas generales que todo marco debe salvaguardar, así como una crítica referenciada de algunas estructuras empleadas en el mundo académico. Dado que cada indicador sintético tiene un objetivo y una base de datos particular, la estructura concreta que seguirá cada uno depende, exclusivamente, del autor.

2.2. Recolección de datos

La recolección de datos representa, por tanto, el segundo paso en toda elaboración de indicadores sintéticos. Dejando de lado la viabilidad o acierto con que se elijan ciertas técnicas de normalización, imputación o ponderación, la calidad de los datos originales es vital para

dotar al indicador sintético de fiabilidad, puesto que este vendría a ser una suma de sus partes (Freudenberg, 2003). Construir un sintético consistente será imposible si sus indicadores componentes no son calidad. Como ya se ha explicado, las fuentes de incertidumbre introducen cierto error en el indicador, por lo que la calidad del sintético final no puede mejorar a la de sus indicadores originales. A tal efecto, la aplicación de los preceptos y el criterio de selección elegido en el marco teórico deben minimizar este problema y garantizar la calidad de los indicadores sintéticos tanto como sea posible.

Una apropiada selección de datos explicita claramente qué criterio se ha seguido y cómo se ha aplicado, tal y como ocurre en la versión más reciente del EPI (Wolf *et al.*, 2022), donde se enuncian y explican los criterios de inclusión: relevancia, abordabilidad legislativa, metodología de medición uniforme, verificación de resultados, completitud espacial y temporal, novedad y acceso libre. Estos criterios protegen los preceptos establecidos por la autoridad europea (Eurostat, 2017), a la par que añade alguno más de interés para el indicador en cuestión. Así, este ejemplo resulta edificante para otros autores, en tanto que toda selección de datos ha de contar con algún elemento particularizado para el indicador sintético que se desea construir.

Además, algunos autores defienden la inclusión de variables según la relevancia que estos posean a la hora de describir el fenómeno explorado (Riedler *et al.*, 2015), pero la realidad es más compleja. En ocasiones, no toda la información tiene por qué estar disponible o ser del todo fiable. A veces, incluso, la aplicación de los preceptos de calidad afecta demasiado a la cantidad de información apta. De hecho, Freudenberg (2003) ya describe esta problemática, asegurando que la escasez de conjuntos completos de información cuantitativa y comparativa suele conllevar la introducción de datos cualitativos menos fiables que los remplacen.

En este sentido, también resulta complejo obtener la totalidad de la información para ciertas variables. Si la ausencia de ciertos datos se mantiene en niveles aceptables, esto abrirá la puerta a la imputación de datos ausentes, algo típico y frecuente en la elaboración de indicadores. Sin embargo, si una variable presentara solo el 5 % de los casos, ¿merecería la pena imputar los demás? Sainani (2015) asegura que ha de establecerse un umbral a partir del cual, en lugar de emplear técnicas de imputación o eliminación de casos, se desechen las propias variables, aunque no existe consenso académico al respecto. Esta exclusión podría formar parte del esfuerzo del autor por garantizar el empleo de indicadores fiables y completos, pero sí se debe ser cuidadoso de no eliminar variables significativas para el sintético (Horton y Kleinman, 2007).

Asimismo, el mismo autor asegura que, si un sintético tiene por objetivo reflejar un cierto *output* (el desempeño en innovación en una universidad, por ejemplo), sus indicadores componentes solo deberían albergar información del mismo cariz (producción de patentes, número de artículos académicos publicados...) y no introducir *inputs* (inversión en I+D de la universidad, por ejemplo). La diferencia principal entre ambos conceptos radica en que uno mide efectivamente cierta parte del resultado (*output*), mientras que el otro persigue alterarlo a través de cierta conexión o impacto (*input*). Así, mayor inversión en I+D no implica *per se* un mejor desempeño sino que busca generarlo, aunque su alta correlación con el número de patentes producidas pueda llevar a engaño.

Esta regla de asociación se cumpliría en el caso del IDH (PNUD, 1990), cuyos indicadores componentes representan, al igual que el desarrollo humano, variables de tipo *output* como son la longevidad, la tasa de alfabetismo o el PIB per cápita. Esto se cumple también en el

caso del OHI (Halpern *et al.*, 2012), pero no siempre ocurre. En algunos indicadores sintéticos ampliamente contrastados en el mundo académico, como el EPI de 2008 (Emerson, Jay *et al.*, 2008), se produce una cierta mezcla de ambos tipos de indicadores. En este caso concreto, los subsidios a la agricultura formaban parte del área “Vitalidad del Ecosistema”. Está claro que ambos conceptos no están intrínsecamente relacionados, pero sí se da una cierta relación de causa-efecto que puede justificar esta agrupación. Aun así, esto fue alterado en versiones posteriores del indicador, a medida que la calidad y la fiabilidad del mismo iba en aumento (Hsu y Zomer, 2016).

Además, la ausencia de cierta información puede animar al autor a introducir variables de tipo *input* que, sirviendo de *proxy*, reflejan los datos *output* que se busca representar. El uso de estas variables está extendido, mas su utilización debe analizarse con detenimiento y de forma particularizada, ya que son un elemento más de creación de incertidumbre. Aun así, su uso está ciertamente extendido debido a su utilidad y la tradicional escasez de información (Alam *et al.*, 2016). En este sentido, cabe recordar que los propios indicadores sintéticos podrían definirse como variables *proxy*, en tanto que reflejan una realidad inmensurable por ningún indicador individual (José Antonio *et al.*, 2022).

Por todo ello, no existe un consenso evidente respecto de la metodología que un indicador sintético ha de seguir a la hora de seleccionar y recopilar datos, puesto que dependerá enormemente de su objetivo y del juicio del autor. No obstante, sí resulta importante seguir los criterios de calidad y fiabilidad debidamente establecidos en el marco teórico, de forma que el usuario comprenda fácilmente los motivos que llevaron al autor a incluir los indicadores componentes seleccionados.

2.3. Imputación de información ausente

La información ausente (*missing data* en inglés) representa información no disponible para alguno de los casos medidos por un indicador componente (Gómez *et al.*, 2017). Su aparición resulta un obstáculo en la elaboración del sintético, en tanto que las técnicas posteriores de ponderación y agregación exigen un conjunto de datos para operar correctamente. De hecho, es un problema recurrente en muchos ámbitos académicos, no solo en el mundo de los indicadores sintéticos (Donders *et al.*, 2006). Ciertas herramientas estadísticas no funcionan en un contexto de datos ausentes o, de forma análoga, sus resultados no son realmente fiables. Esto ha contribuido a que la práctica más extendida para solventar el problema sea la eliminación de dichas variables (Buuren, 2018). De hecho, los autores Orchard y Woodbury (1972) aseguran que “obviamente la mejor manera de lidiar con los problemas de información ausente es no tenerlos” (p. 697). Sin embargo, también reconocen que su existencia es a veces inevitable y que, por tanto, concierne a toda la Estadística evaluar las diferentes técnicas que puedan aplicarse a situaciones específicas. Gracias a su cuasiomnipresencia, existe una abundante literatura referente a métodos estadísticos que evitan comprometer la robustez del indicador sintético a causa de una negligente imputación.

El primer paso para dilucidar cómo afrontar la imputación de información ausente radica en comprender las causas que condujeron a dicha pérdida de datos (Li, 2013), según si las causas que la ocasionaron están relacionadas con el propio valor que toma el indicador. Por ejemplo, en una encuesta de predilección política, si los votantes de cierta ideología tienen mayor propensión a no dar información al respecto, su información faltará de forma sistemática

y no totalmente aleatoria. En cambio, si la gente que no responde a la pregunta sobre su ideología tuviera, de media, la misma ideología que los que sí lo hacen, la información ausente sería totalmente aleatoria. De acuerdo con esta explicación, similar a la presente en el manual de la OCDE (2008), se establecen tres principales categorías:

- **Missing Completely At Random (MCAR, ausentes totalmente al azar):** La probabilidad de que un dato concreto falte es puramente aleatoria. Como consecuencia de que la distribución de ausencias sea así, la muestra efectivamente medida tendrá las mismas características que la muestra ausente (Buuren, 2018). Básicamente, los datos ausentes serían indistinguibles de los datos presentes. Esta distribución totalmente al azar no resulta realista, puesto que por lo general hay motivos –conocidos a veces e ignotos otras– que provocan algún tipo de relación en las ausencias que afecta a su aleatoriedad (Bennett, 2001).
- **Missing At Random (MAR, ausentes al azar):** Esta distribución es mucho más realista, en tanto que la muestra de datos ausentes no es idéntica a la medida, pero sí rastreable en base al resto de respuestas o variables medidas. De esta manera, dichos datos pueden estimarse a partir del resto de la muestra medida. Los valores ausentes no dependerían del propio indicador medido (la ideología, en el caso anterior), sino de otras variables conocidas (como, por ejemplo, la renta familiar). Esta distribución opera de forma similar a la MCAR, sin ser tan restrictiva. De acuerdo con Donders *et al.* (2006), las técnicas simples de manejo de datos ausentes proporcionan resultados sesgados, por lo que sería imprescindible emplear técnicas avanzadas de imputación.
- **Not Missing At Random (NMAR, no ausentes al azar):** En el ejemplo sobre la ideología explicado anteriormente, la propia ausencia de un caso depende de su valor: la respuesta a la pregunta sobre predilección política solo dependería de la ideología de la persona. Esto haría que los datos siguieran una distribución NMAR, de forma que no habría manera de predecir su valor a partir de otras variables. Asimismo, otro origen para estas distribuciones podría ser el desconocimiento respecto de las causas que originan dichos datos ausentes o, también, un error sistemático en el instrumento de medida. En el caso NMAR, ninguna técnica de imputación producirá resultados insesgados y debidamente robustos, siendo el caso más complejo de todos (Buuren, 2018).

Básicamente, una distribución MCAR implica que las ausencias son aleatorias, independientes tanto de la variable en cuestión como del resto de variables observadas. Una distribución MAR, en cambio, establece que las ausencias pueden predecirse en base al resto de variables observadas, siendo independientes de la propia variable que presenta la ausencia. Estas dos distribuciones son en ocasiones indistinguibles (Gómez *et al.*, 2017), pero demostrar su existencia es vital para la imputación de datos ausentes, ya que la NMAR presenta difícil solución. De hecho, en estos casos habría que modelar un patrón implícito que justifique dichas ausencias, lo cual acarrearía grandes perjuicios a la robustez del indicador sintético final (OCDE, 2008).

¿Cómo averiguar entonces ante qué distribución se encuentra el autor? Roderick J. A. Little (1988) desarrolló un método a través del cual se evalúa si los datos estudiados siguen una distribución MCAR. Sin entrar en sus complejidades matemáticas, esta prueba permite

descartar, a partir del típico p-valor estadístico (en relación al valor del nivel de significación estadístico), la naturaleza MCAR de un indicador. Por tanto, cuanto mayor sea el p-valor, menos dudas habrá de la aleatoriedad de su información ausente. Además, por si fuera poco, otros autores (Li, 2013) han realizado contribuciones a la técnica de forma que el método distinga también entre MAR y NMAR. De esta manera, a través de cierto análisis empírico (test de Little) pero también cualitativo de las causas que motivan la aparición de información ausente (juicio del autor), los métodos de imputación producirán resultados fiables, siempre y cuando se haya logrado demostrar la ausencia de naturaleza NMAR en los datos.

Una vez explicado esto, cabe enunciar los diferentes métodos de imputación desde un punto de vista cualitativo, ya que los fundamentos matemáticos de los mismos son harto complejos y, además, se encuentran debidamente explicados en las referencias. No obstante, sí resulta edificante comprender las ventajas y las desventajas que presenta cada técnica, así como el momento y la estructura de datos en que cada una debe ser aplicada. En líneas generales, hay tres formas de enfocar esta problemática: la eliminación del caso (*case deletion*), que suprime los casos que no dispongan de toda la información; la imputación simple, que presenta varias técnicas de único proceso; y la imputación múltiple, que itera varios procesos de imputación para luego combinarlos en un conjunto de datos final (Calafati, 2017). Por supuesto, cada una de estas tres categorías encierra técnicas variadas con impactos bien diferenciados sobre los datos. De esta manera, los principales métodos para lidiar con información ausente se presentan a continuación:

■ **Eliminación del caso** (*case deletion*):

- **Análisis de casos completos o eliminación de la lista** (*listwise deletion*): Si un cierto caso (país o una fecha para series temporales) no presenta información para una de los indicadores componentes, el caso al completo sería desechado del análisis y, por tanto del indicador. Esto solo proporcionaría resultados aceptables en el caso de distribución MCAR, lo cual como ya se ha mencionado es poco realista en la mayoría de ocasiones (Bennett, 2001). Su principal problema radica en la reducción del tamaño de la muestra y, por tanto, de la potencia de cualquier análisis posterior (Newman, 2014). De hecho, desechar información por no estar completa implicaría generar un vacío de información aún mayor, lo cual es especialmente problemático si la aleatoriedad de las ausencias no es totalmente perfecta, lo cual sería una situación utópica.
- **Análisis de casos disponibles o eliminación del par** (*pairwise deletion*): Este método surge como una evolución del anterior. En esta técnica, los casos no son eliminados por completo, sino que se emplean para el análisis estadístico en aquellas variables para las que tengan datos (Bennett, 2001). En un estudio estadístico con MCAR y una pequeña ratio de información ausente, este método arroja resultados prometedores (Newman, 2014). Sin embargo, su aplicación a los indicadores es discutible, ya que estos requieren de conjuntos de datos completos para realizar debidamente las técnicas de ponderación y agregación. En este sentido, futuros trabajos podrán profundizar en la posibilidad de codificar cierto *machine learning* que permita dinamizar estos procesos según la disponibilidad de los datos.

■ **Imputación simple:** Tiende a infravalorar la varianza de los datos ausentes:

- *Hot-deck imputation*: Este método, basado en un algoritmo de similitud, completa la información ausente con los datos disponibles en casos parecidos (OCDE, 2008). Asume que los datos son, por tanto, MAR.
- *Cold-deck imputation*: Esta técnica funciona de forma similar, pero la información es obtenida de fuentes externas a la base de datos o en base a información ya recolectada (Bennett, 2001).
- Imputación de la media: El promedio de la muestra disponible se introduce en los datos faltantes, de manera que se deprecia la variabilidad de la muestra. Por supuesto, este método no funciona bien para datos MAR o NMAR y tiende a producir resultados sesgados (Donders *et al.*, 2006).
- Métodos de regresión: Cada indicador componente con algún dato ausente se convierte en la variable dependiente de una función, cuyos regresores o variables independientes serían el resto de indicadores, especialmente aquellos altamente correlados (Bennett, 2001). Por ello, asume una distribución MAR para los datos. Para mejorar su robustez, la inclusión de residuos con distribución normal en la regresión es típica (Musil *et al.*, 2002).
- *Expectation-Maximization (EM) imputation* (Imputación de esperanza-maximización): Esta técnica está a caballo entre la imputación simple y la múltiple, ya que consta de un algoritmo iterativo en dos pasos: la etapa de esperanza, donde se obtienen valores para los datos ausentes en base a la información disponible; y la de maximización, donde se recalculan los parámetros de la distribución (medias y covarianzas, por ejemplo) con la imputación del paso anterior incluida. De esta forma, el proceso vuelve a comenzar con los nuevos parámetros, de manera que se calculan nuevos valores para los datos ausentes y así hasta que converge en estimadores de máxima verosimilitud (OCDE, 2008).

Suele generar resultados muy cercanos a la realidad para estructuras de datos MCAR y MAR, especialmente si es comparado con otros métodos de imputación simple (Gold y Bentler, 2000). Asimismo, su uso está extendido para muestras de varios tamaños y su desempeño ha sido cuantificado también en relación a un posterior análisis de componentes principales (Malan *et al.*, 2020), arrojando muy buenos resultados. Su único inconveniente es que, en ocasiones, la convergencia es lenta.

- **Imputación múltiple (IM)**: Consta de un proceso aleatorio que crea N conjuntos de datos, de forma que la incertidumbre queda representada. Los parámetros de cada conjunto se calculan a través de imputaciones subyacentes para después combinarlos en un conjunto final (Bennett, 2001). Así, cada valor ausente gozará de N valores imputados que son posteriormente fusionados. Las técnicas de imputación subyacentes pueden ser de todo tipo, como las detalladas en la imputación simple (métodos regresivos, por ejemplo). No obstante, una de las técnicas más robustas consiste en el empleo de métodos MCMC (*Markov-Chain Monte Carlo*), basados en procesos estocásticos (Medel Esquivel *et al.*, 2019), donde cada estado, variable o conjunto de datos depende exclusivamente

del valor que lo precede. Su justificación técnica es compleja, pero presenta muy buenos resultados en términos de robustez y representatividad de los datos imputados (OCDE, 2008).

La elección del método de imputación es compleja; pero, al igual que en el caso de la determinación de la naturaleza MCAR, MAR o NMAR, el análisis técnico debe ser acompañado del juicio crítico y referenciado del autor. En este sentido, la literatura establece ciertas reglas mínimas, tales como que la eliminación de casos ha de evitarse o que, a partir de un 10 % de información ausente, métodos robustos como IM o EM han de ser empleados (Newman, 2014). Además, para poder comparar las técnicas apropiadamente, Derrick A. Bennett (2001) ya presenta una síntesis cualitativa de los principales métodos de imputación, la cual ha sido actualizada y contrastada con la literatura posterior (Buuren, 2018), (OCDE, 2008), (Gómez *et al.*, 2017). La tabla 3 resume lo mencionado:

Método	Sesgo	Variabilidad
<i>Listwise deletion</i>	Insegada para MCAR; si no, muy sesgada	Subestimada
<i>Pairwise deletion</i>	Insegada para MCAR; si no, muy sesgada	Subestimada
<i>Imputación de la media</i>	Insegada para MCAR; si no, muy sesgada	Subestimada pero supone una mejoría
<i>Métodos de regresión</i>	Insegado para MCAR y MAR	Varianza subestimada pero mejorable por residuos
<i>EM imputation</i>	Insegado para MCAR y MAR	Estimadores fiables
<i>Imputación múltiple</i>	Insegado para MCAR y MAR	Estimadores fiables

Tabla 3. Síntesis de las principales técnicas de imputación de información ausente

Para concluir este capítulo, solo resta comprender qué técnica podría resultar más fiable a la hora de aplicarse para un conjunto de datos concretos. Por supuesto, el propio manual de la OCDE (2008) asevera que no existe una “regla de oro”. Al fin y al cabo, si bien el método de imputación múltiple y la imputación EM presentan las mejores cualidades, la puesta en práctica puede acarrear sorpresas. Para evitarlo, al igual que hicieron Gold y Bentler (2000), se propone llevar a cabo un análisis empírico que cuantifique qué técnica es más apropiada. El autor habría de tomar parte disponible del conjunto de datos, eliminar algunos de ellos de forma aleatoria y, posteriormente, aplicar varios métodos de imputación. El valor imputado presentará una desviación respecto del valor observado, lo cual habrá de ser reflejado por algún índice o medida como, por ejemplo, el error cuadrático medio normalizado (NRMSE, por sus siglas en inglés).

Otros parámetros útiles podrían ser el coeficiente de correlación o el índice de acuerdo de Willmot (1985), ambos propuestos junto al error cuadrático medio (RMSE) o el error absoluto medio (MAE) como instrumentos de medición de la idoneidad de una técnica de imputación (OCDE, 2008). Sin embargo, estas herramientas no parecen considerar la naturaleza

multidimensional de los indicadores sintéticos, especialmente en una etapa del proceso de elaboración donde las escalas de los indicadores componentes aún no han sido normalizadas. Esto podría hacer que, aunque un método tuviera un desempeño general mejor, una imputación pobre en la variable cuya escala es mayor implique obtener el error cuadrático más alto. Por ello, parece más razonable aplicar un índice normalizado por la media de la distribución (*i.e.*, NRMSE) que, comparado para las diferentes técnicas, arroje una luz fiable y empírica sobre las mismas (Shcherbakov *et al.*, 2013). En caso de anhelar una comparación por subgrupos o categorías, la normalización permite sumar los NRMSE de los distintos indicadores componentes, cuya abundancia provoca a menudo que un análisis individualizado sea inviable en términos de tiempo y esfuerzo. De esta manera, la ecuación del NRMSE para n indicadores componentes es la siguiente:

$$NRMSE = \sum_i^n NRMSE_i = \sum_i^n \frac{1}{\bar{x}_i} \cdot RMSE_i = \sum_i^n \frac{1}{\bar{x}_i} \cdot \sqrt{\frac{\sum_{\forall t} (x_{it}^{imp} - x_{it}^{obs})^2}{N}} \quad (2)$$

Básicamente, este método empírico analizaría el comportamiento de los datos para cada técnica de imputación, de manera que el procedimiento que presentara un menor NRMSE, para un mismo subgrupo o indicador componente (dependiendo de su cantidad y la capacidad computacional disponible), sería la técnica menos lesiva para el conjunto de datos. Pese a ello, es importante recordar que esta comprobación no es la panacea ya que, en tanto que las técnicas se comparan exclusivamente sobre la base de datos disponible, ciertos datos ya habrían sido descartados a la hora de analizar empíricamente qué técnica emplear. Esto supone una fuente de incertidumbre y refuerza la idea de que, en la elaboración de los indicadores sintéticos, toda decisión debe ser respaldada por tanto análisis cuantitativos como cualitativos en la medida de lo posible. Al fin y al cabo, en este proceso de elaboración toda decisión goza de ventajas y desventajas.

De esta forma, el estado de la cuestión sobre la imputación de datos ha sido explicada, por lo que a continuación se explorará la situación del análisis multivariante, imprescindible para comprender las relaciones existentes entre los propios indicadores componentes. De sus conclusiones, los autores podrán juzgar con mayor tino qué técnicas de ponderación o agregación son más apropiadas.

2.4. Análisis multivariante

Tal y como se ha mencionado, el objetivo del análisis multivariante consiste en esclarecer las tendencias y correlaciones establecidas entre los diferentes indicadores componentes (Lebart *et al.*, 1984). En esta etapa se revelan las interacciones presupuestas en la elaboración del marco teórico, puesto que en ese momento el autor no dispone de la información estadística completa. Por ello, que estructura y análisis muestren consistencia es importante en la construcción del indicador sintético en términos de fiabilidad. La información obtenida en el análisis multivariante debería arrojar una luz imprescindible para aplicar técnicas de ponderación y agregación idóneas, así como evaluar de forma general la estructura de los datos, de ahí su relevancia (Floridi *et al.*, 2011).

Para averiguar dichas correlaciones (multicolinealidad si es lineal), existen infinidad de técnicas exploradas por el campo de la Estadística. De forma ilustrativa, se recomienda la

lectura del trabajo de Carles M. Cuadras (2007), donde se describen una cantidad ingente de las susodichas. Como el análisis multivariante no forma parte del objetivo de los indicadores sintéticos, se recomienda a los autores el empleo de técnicas que faciliten la comprensión de los datos, en tanto que desvelen cuestiones útiles para la elaboración del indicador. Al fin y al cabo, ciertos procedimientos aportan la misma información, así que queda a juicio del autor del sintético identificar cuántos y qué tipos de análisis multivariantes son necesarios para caracterizar correctamente los datos. A continuación, se presentan algunas de las técnicas más empleadas (Santos, 1996), con atención a las aplicables en la elaboración de indicadores:

- **Regresiones auxiliares:** Dependiendo de la cantidad de datos, la posibilidad de analizar el coeficiente de correlación de Pearson entre dos variables es viable (Floridi *et al.*, 2011), lo cual puede ser estudiado a partir de una matriz de correlaciones o incluso una matriz de gráficos de regresiones auxiliares (Cuadras, 2007). Básicamente, representa cada indicador componente en función de cada uno de los demás componentes. Por ello, este método solo goza de utilidad para un conjunto de datos de tamaño pequeño o mediano, ya que evaluar matrices de 150x150 se torna una empresa muy compleja. Así, un valor del coeficiente muy alto entre dos variables implicará una gran colinealidad (Chumacero, 2015).
- **Principal Component Analysis (PCA):** Joya de la corona del análisis multivariante. Esta técnica es capaz de reducir un número importante de variables o indicadores correlacionados en unos componentes linealmente independientes, los cuales reflejan casi toda la variabilidad estadística del conjunto original (OCDE, 2008). Por ello, el método requiere de cierta correlación en las variables iniciales, por lo que su aplicación puede apoyarse en otras técnicas multivariantes que la demuestren (Illahi y Mir, 2021). Además, un análisis preliminar de adecuación de la muestra para PCA es necesario. Este test se conoce como *Kaiser-Meyer-Olkin* (KMO) y refleja cuán idóneo es el conjunto de datos para ser sometido a un análisis PCA. Típicamente, se suele establecer la regla $KMO > 0.6$ como suficiente para emplear la técnica (Reisi *et al.*, 2014).

Asimismo, hay una cantidad importante de suposiciones que deben hacerse para obtener resultados fiables del análisis PCA (OCDE, 2008), como son disponer de un número suficiente de casos, típicamente un mínimo de diez (aunque hay ratios de casos y variables más restrictivas); seleccionar los indicadores originales sin un sesgo particular, lo que debe establecerse en el marco teórico y en la selección de datos; evitar los valores atípicos y las variables ordinales; gozar de linealidad y altas correlaciones, lo cual se demuestra mediante el mencionado KMO, la prueba de esfericidad de Bartlett y métodos similares; e idealmente seguir una distribución multivariante de normalidad, lo cual es analizable a través de pruebas como las de Kolmogorov-Smirnov o Saphiro-Wilk (Razali y Wah, 2011). En ocasiones, estos supuestos no se cumplirán a la perfección, lo cual disminuirá la potencia del análisis de componentes principales.

La técnica muestra la importancia de cada indicador original (carga factorial) sobre los componentes, así como cuánta variabilidad estadística explica cada uno de ellos. Por lo general, los indicadores más correlacionados se agruparán en un componente común, lo que mantiene la interpretabilidad del conjunto en niveles aceptables (OCDE, 2008). Así, la siguiente tabla 4 presenta las cargas factoriales de un conjunto de datos cualquiera:

Matriz de componentes			
	Componentes		
Indicadores	1	2	3
Tasa de Paro Ajustada por ERTES y Salidas	-0.264	0.783	0.297
Paro Registrado	0.819	0.435	-0.356
Paro Registrado Mujeres	0.862	0.350	-0.361
Solicitudes Relativas Vacantes por Mes	-0.205	0.808	0.326
Tasa de Paro Ajustada Eurostat	0.562	-0.139	0.754
Tasa de Paro Trimestral	0.785	-0.265	0.413

Tabla 4. Ejemplo de una matriz de componentes obtenida a partir del método PCA

Se han resaltado las cargas factoriales mayores de 0.5, que serán las que mayor impacto tendrán en la composición de los componentes. La fórmula que muestra cómo se extraen valores para unos indicadores intermedios a partir de las cargas factoriales es la siguiente:

$$Z_{jt} = a_{1j} \cdot x_{1t} + a_{2j} \cdot x_{2t} + \dots + a_{ij} \cdot x_{it} \quad (3)$$

Sin embargo, en ocasiones un indicador aparece con fuerza en varios componentes, tal y como queda patente en la tabla 4. Esto empobrece la interpretabilidad y comprensión del componente, lo cual enturbia futuras etapas de ponderación y agregación. Por ello, la rotación *varimax* es una herramienta útil para solventar este problema. Así, esta rotación de componentes provoca que cada variable aparezca mayormente en un único componente. Ejemplo de ello se muestra en la siguiente tabla 5:

Matriz de componentes rotada			
	Componentes		
Indicadores	1	2	3
Tasa de Paro Ajustada por ERTES y Salidas	-0.016	-0.086	0.874
Paro Registrado	0.986	0.114	0.051
Paro Registrado Mujeres	0.986	0.153	-0.036
Solicitudes Relativas Vacantes por Mes	0.026	-0.034	0.894
Tasa de Paro Ajustada Eurostat	0.021	0.949	0.056
Tasa de Paro Trimestral	0.302	0.840	-0.245

Tabla 5. Matriz de componentes rotada con el método *varimax*

En este caso, se puede comprobar claramente cómo la asociación de componentes con indicadores es más sencilla, ya que cada variable solo tiene impacto sobre uno de los componentes. De hecho, el primer componente podría entenderse como una variable sobre el *número de parados*, mientras que los dos últimos versan ambos sobre la *tasa de paro*, con algún matiz. Pese a que no solucione a la perfección el problema, sí es notoria la mejora en la interpretabilidad del significado de los componentes. Por tanto,

cabe reconocerle a la rotación su utilidad a la hora de facilitar la comprensión del PCA.

Aun así, todavía no se ha planteado una importante pregunta: ¿cuántos componentes han de tomarse? En realidad, aparecerán tantos como variables originales haya en el conjunto (Cuadras, 2007). No obstante, sus autovalores (obtenidos a partir de las cargas factoriales) serán ínfimos en la mayoría de los casos, lo que implica que apenas explicarán nada del conjunto de datos original. Esta es la razón por la que las tablas 4 y 5 solo disponen de tres componentes, puesto que se ha aplicado en ellas el criterio de Kaiser que se explicará a continuación. Al fin y al cabo, depende del autor establecer el criterio a partir del cual se descarten los componentes (Reisi *et al.*, 2014). De hecho, existen multitud de métodos para hacerlo (Peres-Neto *et al.*, 2005), algunos de los cuales se muestran a continuación:

- *Criterio de Kaiser*: Defendido también por Guttman (1954), exige obtener los componentes cuyo autovalor sea superior a uno. Después de todo, no tendría sentido introducir un componente que explique menos que una variable cualquiera (Illahi y Mir, 2021). Es uno de los métodos más empleados por su simplicidad y sus buenos resultados.
- *Criterio del bastón roto*: Implica la división aleatoria de la varianza explicada según el número de componentes (Peres-Neto *et al.*, 2005). Si la varianza explicada por el componente k es mayor que el trozo asignado de varianza (para cuatro componentes: $E_1=52,08\%$; $E_2=27,08\%$; $E_3=14,58\%$...), este debe tenerse en cuenta. La fórmula empleada para la “romper el bastón” se encuentra en las referencias (Cuadras, 2007). Presenta unos resultados sorprendentemente positivos (Jackson, 1993).
- *Criterio del porcentaje de varianza explicada*: Si el autor se propone un umbral (80 % por ejemplo), tomará tantos componentes como sea posible para alcanzar dicha varianza explicada (OCDE, 2008).

La lista de criterios es prácticamente inabarcable. Aun así, otros métodos interesantes serían el test de esfericidad (Peres-Neto *et al.*, 2005), el criterio de Joliffe (similar al de Kaiser, pero con un umbral de autovalor en 0.7) o a través del gráfico de sedimentación de Cattell (OCDE, 2008). Asimismo, el propio manual también expresa una tercera vía, que vendría a denominarse “comprensibilidad”, donde básicamente todo depende de lo que el autor encuentre coherente e interpretable. Aun así, por lo general, suele darse una combinación de dos o más criterios sosteniendo la decisión del autor.

- **Análisis de factores** (FA, por sus siglas en inglés): Muy similar al PCA, salvo porque no se emplea un modelo de combinación lineal en su procedimiento. Esta técnica proporciona factores “latentes” (Cuadras, 2007) que explicarían cada uno de los indicadores originales. Así, se “invierte” el proceso, de forma que cada indicador es explicado por varios factores por determinar:

$$x_{it} = a_{i1} \cdot F_{1t} + a_{i2} \cdot F_{2t} + \dots + a_{ii} \cdot F_{it} \quad (4)$$

Resulta interesante compararla con la ecuación 3, donde cada componente era explicado por los indicadores originales. Los criterios de selección de factores funcionan de manera similar, así como las posibles rotaciones a las que se pudieran someter las matrices.

- **Coefficiente alfa de Cronbach:** Mide la consistencia interna de un conjunto de datos (Heo *et al.*, 2015), es decir, cuánto mayor sea el coeficiente ($0 < \alpha < 1$) más consistentes son las variables entre sí. Esta cualidad aludiría a la capacidad del conjunto de datos de representar una realidad subyacente común a todos ellos (OCDE, 2008), por lo que indirectamente es un indicador clave para cuantificar la multicolinealidad. El valor umbral a partir del cual resulta aceptable considerar la existencia de correlación en la muestra depende del autor, aunque se considera una buena referencia el valor $\alpha > 0,7$ (Cortina, 1993). Asimismo, también cabe analizar los coeficientes alfa de Cronbach en caso de eliminarse alguna variable en particular, puesto que su variación implicaría que el indicador excluido presenta poca correlación con el resto de la muestra.

Otras técnicas de uso estadístico, como la MANOVA (ANOVA multivariante), no tienen mucha utilidad en el caso de los indicadores, ya que no existen grupos predefinidos con los que comparar medias. En cambio, métodos como el análisis de grupos (*cluster analysis*) sí son mencionados en el manual de la OCDE (2008), cuya principal aplicación consistiría en la agrupación de casos (países, fechas...) mediante la ejecución de un algoritmo. Suelen funcionar de acuerdo con la “distancia” existente entre las variables, siendo las técnicas principales la distancia euclídea, la distancia euclídea al cuadrado, la distancia de Manhattan, la de Mahalanobis o la de Chebychev (Matthiesen, 2010). Un ejemplo del análisis de grupos podría ser el siguiente dendrograma, que resulta visualmente muy didáctico a la hora de analizar el parecido entre casos y, si así desea, entre variables:

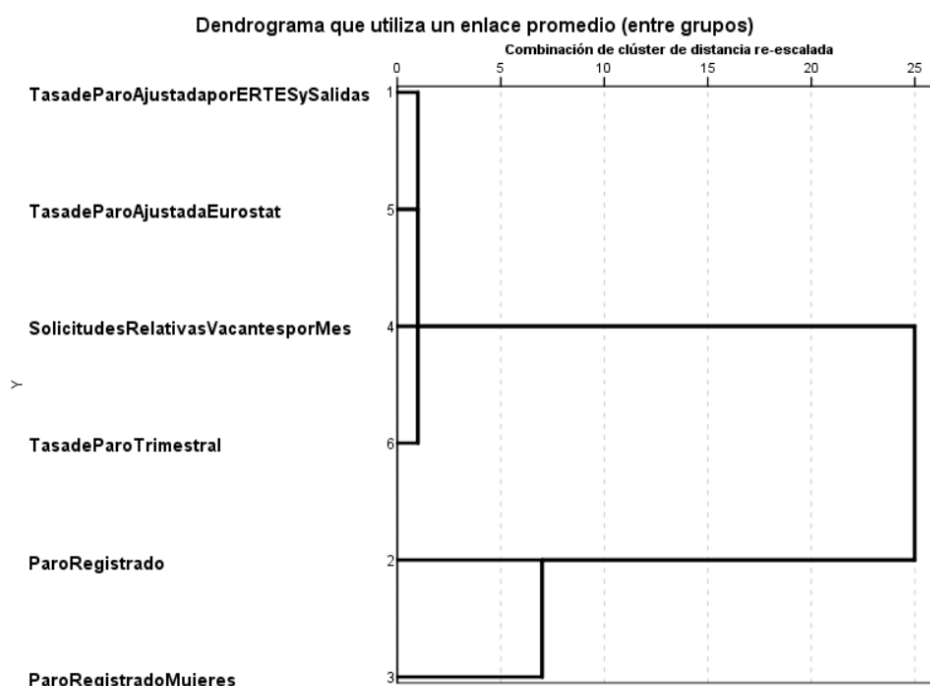


Figura 3. Dendrograma a partir del enlace promedio y del cuadrado de la distancia euclídea

Al igual que ocurría con la matriz rotada de componentes de PCA, se percibe cómo las variables referidas a la *tasa de paro* (que antes conformaban los componentes dos y tres), mientras que el primer componente *número de parados* se encuentra a mayor distancia.

Para más información sobre otras formas de agrupar (vecino más lejano, método de Ward...) o incluso una técnica ligeramente diferente (*k-means clustering*) se recomienda acudir a la

literatura (Hartigan, 1975), (Matthiesen, 2010), (OCDE, 2008). Si bien estos métodos son muy interesantes, el análisis multivariante realizado debe describir la estructura de datos y esclarecer correlaciones, por lo que la utilización de técnicas poco usuales no es primordial. Por ello, si los análisis presentados ya son capaces de cumplir con creces los objetivos, el empleo de otras técnicas como, por ejemplo, el análisis de correlación canónica (muy limitado a muestras pequeñas de datos) es innecesario.

De esta manera, se da por finalizado el estado de la cuestión referido al análisis multivariante, por lo que la normalización centrará ahora el foco del estudio.

2.5. Normalización

La normalización representa el paso en que los autores homogeneizan su estructura de datos, dotando a sus indicadores componentes de una escala idéntica. Poniendo un ejemplo, si alguien quisiera estimar cuál ha sido el desempeño del Real Oviedo en la pasada temporada, podría medir variables tan distintas como el número de puntos logrados y el presupuesto (medido en millones de euros) para todos los equipos de fútbol de la categoría. Si la intención es obtener un indicador sobre el desempeño, ambas variables (y otras varias) han de ser agregadas. Entonces, ¿cómo habría de hacerse? ¿Deben los pesos encargarse de ello? Proponer que los pesos de la ponderación equilibren las diferentes escalas de medida daría al traste con el objetivo de la ponderación. Después de todo, los pesos deben reflejar en cierta medida la importancia de cada variable, así como su posible sustituibilidad (Greco *et al.*, 2019). Por lo tanto, un paso imprescindible es la normalización de las dos variables, de manera que ambas empleen las mismas unidades.

Por supuesto, al igual que en otras etapas de la elaboración, existen muchos métodos para lograr la deseada normalización. Este abanico varía desde la típica estandarización (conocida y aplicada a menudo por todo estudiante de estadística que se precie) hasta métodos menos ortodoxos, como pueda ser asignar una puntuación a modo de *ranking* (Saisana y Saltelli, 2011). No obstante, en ocasiones cierta transformación previa a la normalización ha de tenerse en cuenta (Freudenberg, 2003). Algunos cambios de escala, como pasar de kilómetros a millas (ratio) o de Fahrenheit a Celsius (transformación lineal), no afectan en gran medida a la posterior normalización. Sin embargo, autores como Groenveld y Meeden (1984) sí describen como problemática la existencia de una distribución con una asimetría absoluta mayor que 1 y una curtosis superior a 3.5, algo que las mencionadas transformaciones de escala no solucionan. En ese caso, el autor debe plantearse qué otras medidas puede tomar para evitar grandes desajustes en el sintético final. Por ejemplo, cabría realizar una transformación logarítmica:

$$y_{it} = \log(x_{it}) \quad (5)$$

Sin embargo, también cabe resaltar que esta transformación también acarrea problemas. Después de todo, exige valores positivos y penaliza demasiado a los valores cercanos al cero, ya que el logaritmo les da un valor negativo muy grande. Su aplicación es, por ello, arbitraria y debe realizarse con sumo cuidado (OCDE, 2008), evaluando si el remedio es o no peor que la propia enfermedad, puesto que se pueden alterar la estructura de los datos de forma irreversible.

Por ejemplo, si se midiera la temperatura diaria máxima en Pamplona a lo largo del año, la variable gozará de un rango desde 0°C hasta 40°C, aproximadamente. Dado el aumento en

la frecuencia de las olas de calor (Núñez, 2019), las efemérides cálidas son más frecuentes y, por tanto, se dará una serie de valores atípicos que provocarán una distribución asimétrica positiva. En este caso, uno podría valorar la aplicación de logaritmos. Sin embargo, se enfrentará a dos problemas: En primer lugar, los datos cercanos al cero se convertirán, casi con total seguridad, en valores atípicos extremos en la nueva distribución, lo cual ya de por sí desvirtúa la realidad reflejada por los datos originales. En segundo lugar, incluso aunque el anterior problema no ocurriera, el autor ha de plantearse qué pretende con el indicador: ¿son los *outliers* algo necesariamente indeseado? Tal vez un estudio de efemérides, basado en fenómenos poco frecuentes, no deba tratar de deshacerse de dichos valores atípicos.

Pese a ello, tal y como aseveran Saltelli y Saisana (2011), los valores atípicos a veces tienen un impacto indeseado sobre los indicadores, ya que en ocasiones se deben a errores de medida o ajenos a la realidad descrita y alteran toda la estructura de los datos. Para detectarlos, es común en la estadística emplear el método del rango intercuartílico (Hudrliková, 2013):

$$Outlier \in (-\infty, Q_1 - 1,5 \cdot (Q_3 - Q_1)) \cup (Q_1 + 1,5 \cdot (Q_3 - Q_1), \infty) \quad (6)$$

Ante su existencia, Freudenberg (2003) propone truncar dichos casos para que no afecten al indicador final. Otro proceder sería el empleado en la elaboración del ESI (precursor del EPI), donde los valores por encima o por debajo de los percentiles 97.5 % y 2.5 % se “saturan” a dichos percentiles (Nardo *et al.*, 2005). Sea cual fuere la decisión del autor, está claro que cierta decisión ha de tomarse respecto de los valores atípicos (Burgass *et al.*, 2017). Esta decisión ha de estar basada en múltiples factores, como son el objetivo del indicador sintético, la fiabilidad de la toma de datos y el resultado que se obtiene de la supresión de dichos valores. Dado que este paso puede ser muy sensible, el análisis de robustez final debe estudiar con detenimiento el impacto de esta decisión.

Una vez esto se ha decidido, cabe estudiar las principales técnicas de normalización:

- **Puntuación:** Se trataría de asignar valores enteros en base al desempeño (Saisana y Saltelli, 2011) o incluso una “categoría” (excelente, bueno, regular...) según su posición relativa a umbrales o percentiles (Freudenberg, 2003). Estos métodos son insensibles a los valores atípicos, pero desperdician mucha información debido a que obvian la “distancia” entre los valores originales. Presentan una gran arbitrariedad en el establecimiento de los umbrales, por lo que se desaconseja su uso (EC, 2022). De forma similar, establecer valores según la cercanía a la media (-1 muy inferior, 0 en una región determinada y 1 muy superior) está también presente en la literatura (OCDE, 2008). Esta técnica sí es sensible a valores atípicos, por lo que se aconseja emplear otro estadístico referente, como pueda ser la mediana.
- **Estandarización:** Transforma la distribución –suponiendo normalidad– en una de media nula y varianza unidad:

$$I_{it} = \frac{x_{it} - \bar{x}_i}{\sigma_{x_i}} \sim N(0, 1) \quad (7)$$

Los *outliers* tienen por tanto un gran impacto, en tanto que distorsionan mucho la media (Burgass *et al.*, 2017). Los valores resultantes se encontrarán en torno al valor 0, pudiendo tomar valores negativos. Esto puede resultar contraproducente a la hora de ponderar y agregar. Además, en el caso de series temporales, cabe la posibilidad de tomar valores

referenciales para la media y la desviación típica, de forma que la normalización no tenga que recalcularse (OCDE, 2008).

- **Método Mín-Máx:** Está técnica, ampliamente refrendada por la literatura (Zhou *et al.*, 2007), normaliza los valores de acuerdo con el rango que puede tomar el indicador en cuestión. Así, se obtiene la siguiente ecuación:

$$I_{it} = \frac{x_{it} - \min(x_{it})}{\max(x_{it}) - \min(x_{it})} \quad (8)$$

donde $\max(x_i)$ y $\min(x_i)$ son los valores máximo y mínimo que toma la variable para todos los casos (Saisana *et al.*, 2005). Los valores normalizados tendrán, por tanto, un rango de entre 0 y 1, lo cual facilita enormemente su ponderación y agregación. No obstante, sigue siendo muy sensible a la presencia de valores atípicos. Al contrario que con la estandarización, la inclusión de nuevos casos o fechas (en series temporales) genera problemas, puesto que una valor antes máximo, puede ser superado por el nuevo conjunto de datos, lo cual generaría valores normalizados mayores que 1 y obligaría a recalcular todo el proceso (EC, 2022).

- **Distancia a referencia:** A modo de ratio, divide el valor del indicador por una referencia u objetivo definido por el autor, que bien podría ser el valor máximo (Floridi *et al.*, 2011) o incluso la media (Freudenberg, 2003):

$$I_{it} = \frac{x_{it}}{x_i^{ref}} \quad (9)$$

Este método es también muy sensible a valores atípicos, aunque es de esperar que, si es elegido por el autor, su influencia se haya considerado anodina.

Como en otras etapas de la elaboración, existe una cantidad casi inabarcable de técnicas de normalización disponibles (Jain *et al.*, 2005). Sin embargo, parece haber cierto consenso en cuanto a la utilización de métodos que conserven las “distancias” entre los valores, es decir, que generen escalas proporcionales (Talukder *et al.*, 2017). Pese a ello, se recomienda un análisis de robustez y sensibilidad que permita dilucidar el efecto concreto de cada técnica sobre el sintético final, de manera que toda decisión del autor esté debidamente justificada.

Asimismo, no parece razonable finalizar la sección sobre normalización sin aludir a otro fenómeno relevante que facilitará la labor de ponderación y agregación de forma notable: el sentido de la escala, puesto que algunos indicadores medirán aspectos negativos que irán contra el objetivo del sintético. Por ejemplo, si medimos la situación del mercado laboral de un país, mezclar variables positivas (tasa de empleo) con negativas (cierre de empresas) puede ser problemático para asignar pesos posteriormente. De hecho, muchas técnicas de ponderación emplean solo pesos positivos, asumiendo esta condición (OCDE, 2008). Por lo tanto, se recomienda orientar todos los indicadores componentes de manera que tengan el mismo “sentido”: un valor mayor del indicador implicaría un mejor desempeño.

Esto es muy fácil de realizar en el caso de la estandarización, donde cambiar el signo de los valores para las variables que sean “negativas” es suficiente, de forma que se mantengan las

condiciones de normalidad (media nula y varianza unidad). En el caso de otras normalizaciones, bastaría con invertir los valores de la siguiente manera:

$$I'_{it} = I_{it}^{max} - I_{it} \quad (10)$$

Nótese cómo en el caso de una normalización Mín-Máx, la ecuación 10 tomaría la forma:

$$I'_{it} = 1 - I_{it} \quad (11)$$

De esta manera, las principales técnicas de normalización han sido explicadas. A modo de resumen, se muestra a continuación la figura 4 que muestra el impacto de cada una sobre la variable ejemplo *Paro Registrado* (ya empleado en las tablas 4 y 5) para unas fechas concretas:

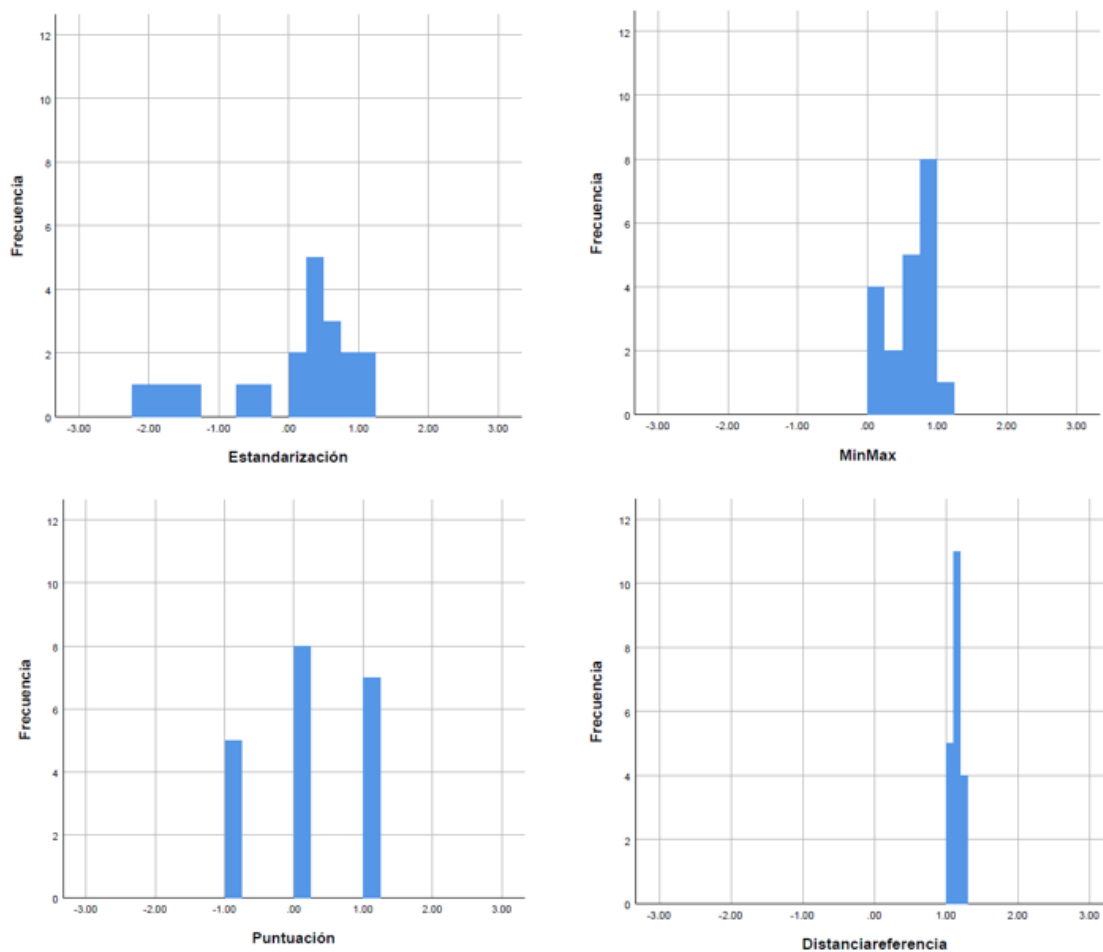


Figura 4. Matriz de histogramas para diferentes técnicas de normalización

La técnica de puntuación se basa en la cercanía a la media (región intermedia: 97 % – 103 % del promedio de la serie), mientras que la distancia a la referencia proviene de la ratio de cada valor respecto del líder de la serie (menor paro registrado). Los métodos de estandarización y Mín-Máx reflejan mejor la verdadera estructura de los datos, ya que la técnica de puntuación desvirtúa las distancias entre las medidas, mientras que la distancia al valor líder comprime excesivamente los datos. Por ello, la aplicación de uno u otro método dependerá, principalmente, del objetivo del indicador, puesto que algunos *rankings* no requieren de distancia entre valores para ser realizados. En el caso de un indicador de serie temporal, mantener la distancia para

garantizar la comparabilidad de los datos es primordial.

En conclusión, la toma de decisiones vuelve a ser vital. Ningún método está exento de criticismo y, aunque algunos parecen ser más ventajosos, comprender el impacto de cada uno sobre el indicador final es necesario para tomar dichas decisiones de forma debidamente justificada, en tanto que se alineen con los objetivos establecidos en el marco teórico. Asimismo, diferentes técnicas de normalización y su influencia sobre el sintético final pueden ser estudiadas en el análisis de robustez.

2.6. Ponderación y agregación

Una vez se han escalado y orientado uniformemente todos los indicadores componentes, la ponderación es la siguiente etapa del proceso de elaboración de los indicadores sintéticos.

2.6.1. Ponderación

La ponderación de los indicadores componentes se refiere al proceso mediante el cual se asignan pesos según cierto criterio para ser después agregadas y formar el indicador sintético final. Típicamente, se considera los pesos una medida de la importancia de cada componente sobre el sintético final (Abdella *et al.*, 2021). Tal vez para alguien poco versado en la materia esta sea la forma más intuitiva de comprender el proceso de ponderación; pero, por desgracia, es un error (Becker *et al.*, 2017), al menos desde la definición clásica de “importancia”.

A la hora de elegir un sistema de ponderación, la importancia relativa efectivamente ha de considerarse, pero también ha de tenerse en cuenta las posibles compensaciones (*trade-offs*) que ocurren debido al método de agregación elegido (Gan *et al.*, 2017). De hecho, de acuerdo con Greco *et al.* (2019), esto distinguiría la importancia explícita, la más intuitiva y atribuible a la relevancia común, de la implícita, más compleja e ignorada sin un análisis *ad hoc*. Además, el hecho de que cierta ponderación interprete cómo serán agregados después los indicadores, provoca que a menudo un método imponga otro (Dialga y Thi Hang Giang, 2017).

Para facilitar la comprensión del fenómeno de la ponderación, se propone el siguiente ejemplo: un sintético *CI* está formado por los indicadores componentes ya normalizados I_1 , I_2 e I_3 , de forma que I_1 e I_2 son, para las partes interesadas, los más “importantes”. Así, el autor del sintético decide emplear una técnica basada en su opinión y asignar unos pesos que reflejen dicha relevancia, asignando los pesos $w_1 = 0,4$ y $w_2 = 0,4$. La última variable, considerada menos influyente o relevante, obtendría un $w_3 = 0,2$. La ecuación 12 representa la agregación aritmética de dicho sintético:

$$CI = w_1 \cdot I_1 + w_2 \cdot I_2 + w_3 \cdot I_3 \quad (12)$$

Sin embargo ¿qué pasa si la variable 1 y la variable 2 están muy correladas? Cuando una aumente, con casi total probabilidad lo hará la otra, en tanto que describirían una realidad común, causa de dicha correlación (Becker *et al.*, 2017). Por ello, este fenómeno común gozaría de un peso asignado de –aproximadamente– 0.8, mientras que la tercera variable dispondría solo de 0.2, lo cual ya iría en contra del objetivo original del autor. Por ello, si bien la importancia explícita, aunque subjetiva, puede ser evaluada por un panel de expertos, la implícita exige una mayor atención (Nardo *et al.*, 2005) y se podría medir, de acuerdo con Paruolo *et al.* (2013),

según el coeficiente de correlación de Pearson, que mide el impacto en la variabilidad del sintético si una de sus variables se fija.

Después de todo, la definición de los pesos es más compleja de lo que *a priori* podía parecer. De hecho, la propia definición de importancia aplicada a la ponderación de indicadores ocupa hoy en día artículos enteros (Becker *et al.*, 2017). Por ello, ante tamaña complejidad, diferentes técnicas de ponderación han surgido para abordar la cuestión. Están clasificadas según el origen de su criterio:

- **Técnicas estadísticas:** Este conjunto de métodos evita la utilización de criterios arbitrarios o subjetivos en la asignación de pesos para la ponderación:
 - *PCA / FA:* Los componentes o factores obtenidos a partir de PCA tienen la ventaja de estar incorrelados los unos de los otros, es decir, que no permite asignar pesos de acuerdo con el coeficiente de Pearson explicado anteriormente. Sin embargo, esto evita el problema de doble conteo explicado anteriormente (Greco *et al.*, 2019). Además, la aparición de componentes o factores complica a menudo la interpretación de los pesos asignados.

Sin embargo, si uno es capaz de identificar a qué dimensión se refiere cada componente (según qué indicadores gozan de más carga factorial), podría aplicarse criterio experto en su interpretación (Gan *et al.*, 2017). Esto es importante, ya que teóricamente, el peso de cada componente se obtendría a partir del porcentaje de varianza explicada por el mismo. Esto refleja cuánta variabilidad explica el factor en cuestión, lo que tampoco es en sí mismo un reflejo de su importancia relativa (Greco *et al.*, 2019). Sencillamente, si varios indicadores correlados entre sí pero independientes del resto forman uno de los factores, la varianza explicada por ellos no tiene por qué tener ninguna relación con su importancia explícita para el sintético final. Este razonamiento debe llevar a los autores a analizar con detenimiento la asignación de pesos en este método, ya que en ocasiones puede no ser fiable.

Rememorando lo explicado en el apartado 2.4, para aplicar estas técnicas ha de darse una gran correlación (Reisi *et al.*, 2014) y un número importante de indicadores que justifique su reducción. Se recomienda acudir a dicho apartado para comprender los criterios de selección de componentes y la rotación que facilita la interpretabilidad antes mentada. Aun así, todavía se puede mejorar la simplicidad de interpretación escalando la matriz de cargas factoriales a la unidad. El proceso para hallar los coeficientes que facilitan la comprensión de la matriz rotada se rige por la siguiente fórmula (OCDE, 2008):

$$\sigma_{ij} = \frac{\lambda_{ij}^2}{\sum_{\forall i} (\lambda_{ij})^2} = \frac{\lambda_{ij}^2}{\phi_j} = \frac{\text{Varianza explicada por el indicador } i \text{ en el factor } j}{\text{Varianza total explicada por el factor } j} \quad (13)$$

Así pues, la tabla 5 de la sección 2.4 se transformaría en:

Matriz	Rotada			Rotada cuadrática normalizada		
	<i>Componentes</i>			<i>Componentes</i>		
<i>Indicadores</i>	<i>1</i>	<i>2</i>	<i>3</i>	<i>1</i>	<i>2</i>	<i>3</i>
1	-0.016	-0.086	0.874	0.000	0.004	0.468
2	0.986	0.114	0.051	0.477	0.014	0.001
3	0.986	0.153	-0.036	0.478	0.008	0.002
4	0.026	-0.034	0.894	0.000	0.000	0.490
5	0.021	0.949	0.056	0.000	0.545	0.002
6	0.302	0.840	-0.245	0.045	0.427	0.037

Tabla 6. Comparación de las matrices rotada y rotada cuadrática normalizada

Además, si se sumaran por columnas los valores al cuadrado de los componentes en la matriz original rotada, se obtendría la varianza explicada de cada uno de los factores, tal y como se asevera en la ecuación 13. Además, pese a que el conjunto de datos usado como ejemplo no es el adalid de la complejidad, la nueva matriz cuadrática normalizada es, si cabe, más sencilla de interpretar que la matriz previa. De hecho, de ella se podrían generar una suerte de indicadores intermedios (OCDE, 2008, p. 90) cuya agregación posterior según la varianza explicada por cada uno proporcione finalmente los subindicadores de las categorías del proyecto.

A este respecto, la creación de dichos indicadores intermedio puede emplear solo los indicadores relevantes (Nicoletti *et al.*, 2000) resaltados en la matriz o, por el contrario, emplear todos. Sin embargo, dado que el método PCA ya es incapaz de representar la totalidad de la varianza del conjunto original (cierto porcentaje se pierde), parece poco riguroso descartar aún más de dicha variabilidad mediante la aplicación de un umbral arbitrario en la matriz rotada cuadrática y normalizada. Sin embargo, en tanto que los valores que tomen dichos indicadores intermedios no aportan información pertinente ni interpretable, su obtención es inane.

A pesar de ello, algo que sí es relevante es el peso proporcionado a cada factor, el cual proviene de su porcentaje de varianza explicada sobre el total. Recordando que ϕ_j es la varianza que cada factor j (una suerte de indicador intermedio) explica del conjunto original de datos, el peso asignado a cada uno sería:

$$w_{jr} = \frac{\phi_j}{\sum \phi_j} \quad (14)$$

Así, para los datos mostrados en la tabla 6, los pesos otorgados a cada factor según los porcentajes de varianza explicada serían:

$$w_{1r} = 0.383 \quad w_{2r} = 0.310 \quad w_{3r} = 0.307$$

En este caso, los pesos son muy similares, como consecuencia de que cada factor explica una parte de la varianza de forma bastante equitativa. Como colofón al ejemplo, resultaría interesante que un panel de expertos ratificara estos pesos en base a la interpretación que se pueda hacer de la matriz rotada cuadrática y normalizada. Dado que no hay correlación entre ellos, no habría riesgo de doble conteo (Greco

et al., 2019). Por ello, la importancia explícita goza de relevancia siempre y cuando sea identificable a partir del análisis, lo cual no siempre ocurre. Así pues, una vez más vuelve a depender del autor la interpretación y la toma de decisiones al respecto. Asimismo, la técnica PCA puede operar tanto con agregación aritmética como geométrica, tal y como demuestran los exhaustivos análisis de robustez de Dialga y Thi Hang Giang (2017). Por si fuera poco, este método se empleó para la elaboración del ESI y constituye el segundo más usado en la muestra de 96 sintéticos analizado por Gan *et al.* (2017).

Análisis envolvente de datos (DEA, por sus siglas en inglés): Este método de ponderación asigna pesos específicos a cada caso, dependiendo de su desempeño (Zhou *et al.*, 2007). Emplea la programación lineal (Nardo *et al.*, 2005) para obtener los pesos a partir de la distancia que separe el caso de una frontera óptima de desempeño. Al igual que en otras técnicas, la suma de los pesos se restringe a la unidad. En el caso explicado por Mahlberg y Obersteiner (2001), donde cuatro países son clasificados según dos indicadores, la frontera vendría dada por la imagen 5:

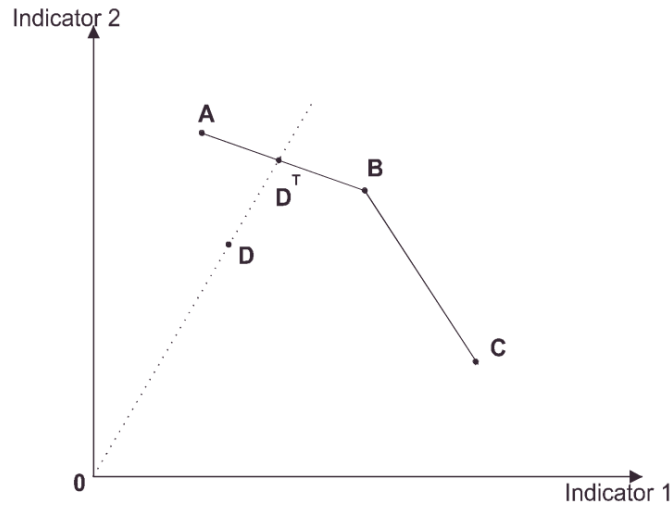


Figura 5. Frontera óptima de desempeño creada por los países A, B y C mediante DEA

La ratio $\overline{OD'} / \overline{OD}$ es utilizada por el programa lineal para asignar los pesos estadísticos para cada país. La gran desventaja del método radica en que tanto A, B y C recibirían la puntuación máxima, incluso aunque A y C tuvieran un rendimiento pobre en uno de los indicadores. Esto se puede solventar introduciendo restricciones adicionales a los pesos (Mahlberg y Obersteiner, 2001) o estableciendo una frontera artificial de acuerdo con los criterios de un panel de expertos (Korhonen *et al.*, 2001). Esto permitiría reconciliar un método muy matemático como es el DEA con los métodos participatorios donde se toman decisiones de índole más subjetiva.

Beneficio de la duda (BOD, por sus siglas en inglés): Vendría a ser la particularización del método DEA para indicadores sintéticos (Freudenberg, 2003). El indicador sintético se definiría como la ratio entre el desempeño del caso y un valor objetivo:

$$CI_i = \frac{\sum I_{ij} \cdot w_{ij}}{\sum I'_{ij} \cdot w_{ij}} \quad (15)$$

La aportación clave de Cherchye *et al.* (2007) en la búsqueda de I'_{ij} como resultado de programación lineal de maximización, provoca que los pesos asignados a cada país sean aquellos que más le beneficien según sus propias características, lo cual comporta de por sí un serio problema. Después de todo, la técnica BOD permite una importante sustitutabilidad entre los indicadores componentes, ya que los pesos finales se ajustarían para potenciar aquellos en los que rindan mejor (Dialga y Thi Hang Giang, 2017), lo que también puede perjudicar la labor de los sintéticos de revelar los problemas sociales.

En este sentido, su uso solo se debe emplear cuando la sustitutabilidad es consistente con el objetivo del indicador (Murias *et al.*, 2008). Asimismo, otra desventaja importante del método es su no comparabilidad con otros métodos de ponderación, ya que la suma de los pesos en este caso no da necesariamente uno (Nardo *et al.*, 2005). Comparar pesos entre países tampoco es factible, dada la total personalización de la técnica. Además, no funciona correctamente con agregación geométrica (Dialga y Thi Hang Giang, 2017).

A su favor, al mismo tiempo, se encuentra dicha personalización. Cada país o caso tiene sus propias preferencias y esfuerzos, por lo que el método permite premiarlos en este sentido. Asimismo, en su versión más estricta tampoco se toman decisiones arbitrarias ni se establecen criterios subjetivos, pese a que la propia elección del método de ponderación ya lo es. No obstante, al igual que con el método DEA, la aplicación de ciertas restricciones puede asistir a los autores para alinear la técnica con sus objetivos concretos, tal y como establece Murias *et al.* (2008).

Finalmente, cabe resaltar la posibilidad de aplicar un BOD invertido, es decir, un BOD “pesimista” (Greco *et al.*, 2019, p. 74). Esta técnica minimizaría la función objetivo, proporcionando un indicador sintético donde se penaliza a cada caso lo máximo posible. Esto tiene una utilidad intrínseca, que soluciona el problema de la sustitutabilidad: para rendir de forma decente, todo caso o país debe cuidar mínimamente todas las variables. Esto le da una utilidad muy aplicable a sintéticos de sostenibilidad donde ningún indicador debería verse compensando por otro.

La literatura es muy abundante en lo referente a los métodos estadísticos de ponderación (Saisana y Tarantola, 2002). En este sentido, cabe mencionar métodos como el modelo de componentes no observables (UCM, por sus siglas en inglés) o los métodos regresivos, que tratan de describir el comportamiento del sintético final mediante una función lineal de variables dependientes desconocidas (UCM) o constituidas por los indicadores componentes. Estos métodos, sin embargo, integran tanto ponderación como agregación, ya que generan un modelo lineal completo (OCDE, 2008).

Pese a ello, los valores atípicos y una gran cantidad de indicadores complican en demasía su ejecución, a la par que el problema del doble conteo sigue presente debido a la posible multicolinealidad (Nardo *et al.*, 2005). Si bien es cierto que hay mecanismos para crear modelos que eviten este problema, como añadir interacciones, las soluciones a veces no merecen la pena por su complejidad, ya que el modelo dejaría de ser lineal.

- **Métodos participatorios:** Estas técnicas fomentan la participación de *stakeholders*, pudiendo ser un panel de expertos o el propio público general:

- *Equiponderación* (EW, por sus siglas en inglés): Este método otorga el mismo peso a todos los indicadores. Destaca por su simplicidad y su gran tasa de utilización, de aproximadamente un 46.88 % sobre una muestra de 96 indicadores sintéticos (Gan *et al.*, 2017). En caso de que el análisis multivariante no revele información relevante sobre el conjunto de datos (Nardo *et al.*, 2005), esta técnica se suele elegir para evitar arbitrariedades en la ponderación. No obstante, su elección también lo es, lo cual parece pasar desapercibido para muchos autores (Becker *et al.*, 2017). De hecho, no se suele identificar con los métodos participatorios, como queda reflejado en el trabajo de Gan *et al.* (2017) donde dispone de su propia categoría (*equal weighting*). De hecho, su elección rara vez aparece acompañada de una justificación a su favor, más bien se aportan argumentos que descartan el empleo otras técnicas.

Este método se ha usado para indicadores tan importantes como el IDH y se considera participatorio porque, desde un punto de vista matemático, resulta igual de arbitrario asignar $w_{1r} = w_{2r} = w_{3r} = 0.3$ como establecer $w_{1r} = 0.5$, $w_{2r} = 0.25$ y $w_{3r} = 0.15$ de acuerdo con la opinión pública. De hecho, podría ser que en ocasiones ambos métodos coincidieran en su asignación.

- *Asignación de presupuesto* (BAL, por sus siglas en inglés): Se conoce comúnmente por representar el criterio experto (Saisana y Tarantola, 2002). Básicamente, se le da a un panel compuesto por expertos y versados en cierta materia un presupuesto a distribuir de x puntos. Siguiendo cierta argumentación crítica, deben asignar a cada indicador cierta parte de esos puntos según la importancia relativa que tengan considerando la realidad descrita por el sintético final. La crítica tradicional que recibe esta técnica –obviando la imposibilidad de interpretar el significado de los pesos asignados– es que, a menudo, los expertos priorizarán urgencia sobre importancia relativa (Nardo *et al.*, 2005). Asimismo, se desaconseja su uso cuando aparece un número elevado de indicadores, ya que esto genera un estrés notable sobre los expertos que desvirtúa su asignación (Saisana y Saltelli, 2011).

Asimismo, en pos de lograr resultados útiles que surjan del consenso, el panel consultado debe ser multidisciplinar y representar los diferentes intereses en juego (Soto de la Rosa y Schuschny, 2009). De esta manera, introducir a todos los *stakeholders* puede ser una opción interesante, considerando la opinión experta pero también la del público afectado pero indocto en la cuestión. De hecho, consultar su parecer constituye, *per se*, una técnica de ponderación: *la opinión pública*.

A este respecto, Roeser y Pesch (2016) abogan por afrontar este tipo de retos desde una perspectiva integradora, atendiendo no solo a la información técnica, sino también a aspectos como la emoción y cómo esta proporciona al individuo cierta información sobre los riesgos que, de otra manera, tal vez permaneciera desapercibida. Así, podría concluirse que para estos autores la urgencia no tiene por qué rehuirse, en tanto que podría ser vital en la consideración del problema tanto o más que la importancia relativa.

- *Proceso de jerarquía analítica* (AHP, por sus siglas en inglés): Basada en la decisión multicriterio, permite establecer comparativamente la prioridad de cada indicador (Nardo *et al.*, 2005). Los expertos deben evaluar qué variable es más relevante y en qué medida, estableciendo valores del 1 (idénticas) a 9 (nueve veces más prioritaria). Así, se colocan los indicadores en una matriz de $n \times n$ con los indicadores enfrentados los unos a los otros. Mediante una técnica posterior de autovectores se extraen los pesos, que seguirán una estructura jerárquica en base a las opiniones de los expertos.

De esta manera, los pesos no representan la importancia de cada indicador, sino su capacidad compensatoria (Lee y Chan, 2008). En ocasiones, se dan problemas de inconsistencia si las prioridades se contradicen, los cuales se subsanan repitiendo la asignación de preferencias (Gan *et al.*, 2017). La ratio de inconsistencia de la matriz podría considerarse una medida del error de juicio cometido por cada experto, por lo que este método ayuda a identificarlos y, en todo caso, a enmendarlos (OCDE, 2008). Además, presenta un proceso transparente de asignación de prioridades, algo no muy frecuente en métodos participativos. Pese a ello, su propio funcionamiento lo hace inviable a partir de un cierto número de indicadores (Soto de la Rosa y Schuschny, 2009).

- *Análisis Conjunto* (CA, por sus siglas en inglés): Esta técnica emplea el análisis multivariante tras la asignación de preferencias por parte del público encuestado (Dziechciarz *et al.*, 1999). Básicamente, se presentan varios conjuntos incompletos de datos, cada uno con valores e indicadores diferentes. Así, los encuestados deben elegir qué combinación de indicadores prefieren. De esta manera, a partir de esta decisión, la técnica determina los pesos según la frecuencia de elección de cada indicador para ciertos datos. Al igual que con otros métodos participativos, presenta la temida sustitutabilidad y también resulta problemática su realización con muchos indicadores, puesto que se pueden dar muchas contradicciones e inconsistencia en las preferencias de los encuestados (Gan *et al.*, 2017).

Por último, cabe mencionar que existe un abanico importante de métodos participativos y que, aún hoy en día, sigue surgiendo nuevas técnicas, como la propuesta por Roeser y Pesch (2016). Pese a ello, es notoria la incapacidad de estos métodos para lidiar con una gran cantidad de indicadores, lo cual los deja obsoletos en determinadas circunstancias.

De esta manera, se han analizado los principales métodos de ponderación, algunos de los cuales integran también el proceso de agregación.

2.6.2. Agregación

Finalmente, la última etapa en la elaboración del indicador consiste en agregar los indicadores componentes empleando los pesos obtenidos en la ponderación. La literatura establece tres tipos principales (OCDE, 2008): aritmética, geométrica y multicriterio. Además, otra forma de clasificarlos podría basarse en la cualidad compensatoria de los indicadores, tal y como establece Munda *et al.* (2003). La agregación aritmética, basada en la suma de los valores por su peso correspondiente, toma la siguiente forma:

$$CI_t = \sum_{\forall i} w_i \cdot I_{it} \quad (16)$$

$$\text{con } \sum_{\forall i} w_i = 1 \quad \text{y} \quad 0 \leq w_i \leq 1$$

La agregación aritmética tiene implicaciones muy significativas en la interpretación de los pesos asignados. Básicamente, la fórmula les asigna un significado similar a una “tasa de sustitución” (Saisana *et al.*, 2005) respecto de los otros indicadores, en lugar de reflejar la importancia relativa como tal. Por ejemplo, para I_{1t} respecto de I_{2t} sería igual a $\frac{\Delta I_{1t}}{\Delta I_{2t}} = \frac{w_2}{w_1}$.

Además, esta técnica exige que los indicadores componentes sean “preferentemente independientes” (Nardo *et al.*, 2005, p. 75), lo que rara vez se cumple en la realidad. Para la ecuación 16, desarrollada de la siguiente manera:

$$CI_t = w_1 I_{1t} + w_2 \cdot I_{2t} + w_3 \cdot I_{3t} \quad (17)$$

Ceteris paribus, el aumento de una unidad en I_{1t} siempre habría de aumentar CI_t en w_1 unidades, al menos mientras se cumpliera la mencionada independencia. Sin embargo, esto no suele ocurrir, ya que las variables se retroalimentan. Por lo general, una mayor presencia de I_{2t} alterará cómo I_{1t} influye sobre la realidad descrita por el indicador CI_t , alejando dicho impacto del valor del peso w_{1t} . Así pues, este peso solo indicará bien la influencia de un indicador en una situación perfecta de independencia preferente.

Tal y como se menciona en la sección 2.6.1 sobre ponderación, los modelos lineales pueden introducir interacciones que describan cómo estas variables afectan las unas a las otras (Rua Vieites, 2021). Teniendo en cuenta la posible interacción entre las variables uno y dos, la fórmula resultaría:

$$CI_t = w_1 I_{1t} + w_2 \cdot I_{2t} + w_3 \cdot I_{3t} + w_{4t} \cdot (I_{1t} \cdot I_{2t}) \quad (18)$$

El peso w_{4t} es, por lo pronto, desconocido. Su hallazgo a través de métodos participatorios, estadísticos o incluso empíricos (en situaciones ajenas a la elaboración de indicadores) depende exclusivamente del conjunto de datos y de su comportamiento. Además, su inclusión afecta a toda la interpretación del modelo, en tanto que ahora el aumento de I_{1t} (*ceteris paribus*), conllevaría un aumento de $w_1 + w_2$ sobre CI_t . Pese al aumento en la complejidad del modelo, la ausencia de independencia preferente no sería tan perniciosa para su interpretabilidad, aunque también es cierto que su inclusión en el modelo nunca será perfecta. Esta interacción, aunque documentada (Becker *et al.*, 2017), rara vez se aplica.

Dejando esto a un lado, el método de agregación aritmética presenta un carácter claramente compensatorio (Munda y Nardo, 2003), de manera que un indicador de la realidad socioeconómica de un país podría compensar la destrucción del ecosistema con la mejora de su economía (Greco *et al.*, 2019). Por ello, depende del autor evaluar cuándo esta sustitutabilidad puede desvirtuar el objetivo del indicador y, si así fuera, elegir una técnica distinta.

La alternativa más obvia vendría a ser la agregación geométrica, basada en la ecuación 19:

$$CI_t = \prod_{\forall i} I_{it}^{w_i} \quad (19)$$

Este método, si bien se debe seguir considerando compensatorio (OCDE, 2008), castiga la mentada sustitutabilidad. En este sentido, un valor muy desmejorado en tan solo uno de los indicadores componentes penalizará con fuerza el valor del sintético, lo que alentaría a los *stakeholders* a mantener unas cotas decentes para todos los indicadores. Al mismo tiempo, una leve mejoría en un componente muy mal posicionado se verá muy recompensado por la agregación geométrica, lo cual también empujaría a las partes interesadas a mejorar su peor problema (Nardo *et al.*, 2005). Este método, si bien no tan empleado como el aritmético, está en creciente uso debido a su fácil implantación y sencillez comparado con métodos no compensatorios (Korhonen *et al.*, 2001).

Además de estos dos métodos, también existen las agregaciones impuestas por la fase de normalización, en tanto que se haya asignado una puntuación (*ranking*) o un umbral sobre la media, por ejemplo. Estos métodos, tal y como se ha mencionado, descartan la “distancia” interna entre los datos del mismo indicador, lo que al final implica una pérdida muy importante de información (EC, 2022). De forma similar, existen métodos basados en las posiciones otorgadas por un panel de expertos o el público general, como por ejemplo sería otorgar puntuaciones a los casos en base a cuántas veces los encuestados los colocan en un puesto preferente (OCDE, 2008). Este método, si bien ha sido revisado, resulta ciertamente arbitrario y está íntimamente ligado a los procesos de ponderación.

El último conjunto de métodos versa sobre el enfoque multicriterio no compensatorio (MCA, por sus siglas en inglés). Su objetivo radica en eliminar por completo la posibilidad de compensar unos indicadores con otros, lo cual resulta en ocasiones de útil aplicación en aquellos casos en que todas las dimensiones del sintético deban cuidarse independientemente las unas de las otras.

Este método ha sido usado con éxito en la elaboración del TAI (Cherchye *et al.*, 2008) y se basa en dos principios básicos: la comparación entre los casos, de forma que se explicita claramente las preferencias subjetivas (Gan *et al.*, 2017); y la asignación de un *ranking* que permita su agregación en base a criterios como el de Borda o Condorcet (Greco *et al.*, 2019). A partir del primer precepto, se obtiene una matriz de prelación (Arrow y Raynaud, 1986) que, a ojos del lector, tal vez recuerde a la presentada por el método AHP de ponderación. Sin embargo, si bien en ese caso se comparaba la “importancia relativa” de cada indicador, ahora se analiza qué valor del indicador es mejor (Nardo *et al.*, 2005).

Por ejemplo, partiendo de una equiponderación inicial, se aplicaría la ecuación 20 (OCDE, 2008):

$$e_{jk} = \sum_{\forall i} w_i(Pr_{jk}) + \frac{1}{2} \cdot w_i(In_{jk}) \quad (20)$$

Así, cuando un indicador i tiene un desempeño mejor para un caso j que para k , su peso computa como preferente (Pr_{jk}). En cambio, si dicho rendimiento es indiferente a ambos casos, se suma dividido por dos (se aplica In_{jk}). De esta manera, se potencian aquellos casos en los que se dé cierta preferencia. Así, se obtendrá dicha matriz de prelación a partir de los valores e_{jk} (Arrow y Raynaud, 1986).

En este punto, cabe plantearse qué puntuación otorgar en base a estos valores, para lo cual existen múltiples métodos, siendo la máxima verosimilitud una técnica fiable (Greco *et al.*, 2019). De hecho, este método también se denomina C-K-Y-L (Munda, 2012) y consta

de un número importante de permutaciones, de forma que se elija el resultado con la mayor puntuación global. Entre sus principales desventajas se encuentran la gran exigencia computacional y su poca tasa de utilización debido a su complejidad (Gan *et al.*, 2017), lo que también influye sobre la cantidad de bibliografía disponible desde un punto de vista comparativo. Además, cabe remarcar que proporciona siempre un *ranking* como resultado, lo que también limita sus aplicaciones. Pese a ello, resulta muy útil por su cualidad no compensatoria.

Para concluir, cabe resaltar la presencia de métodos mixtos, que no se ajustan a la clasificación de técnicas compensatorias y no compensatorias. Entre ellos, cabe destacar algunos como el Índice de Mazziotta-Pareto (Mazziotta y Pareto, 2017), un modelo no lineal que penaliza la presencia de un desempeño pobre en cierto indicador y, por tanto, su sustitutabilidad; la penalización por cuello de botella (Greco *et al.*, 2019), construido en torno al peor rendimiento del caso a través de una corrección previa de los valores; y otros como el modelo ZD o variaciones de la ponderación BOD.

Así pues, las técnicas principales de agregación han sido explicadas, incluso aunque algunas de ellas se encuentren integradas con el método de ponderación. Asimismo, la técnica de normalización empleada debería tener un papel esencial a la hora de juzgar qué método posterior utilizar (OCDE, 2008), puesto que ciertos procedimientos podrían ser incompatibles. Asimismo, cabe resaltar la importante frecuencia con que autores eligen el método de agregación lineal que, si bien requiere de unos supuestos inviables de forma realista (la denominada preferencia independiente), se sigue eligiendo por su asequibilidad. De hecho, en la muestra de 96 indicadores sintéticos analizados por Gan *et al.* (2017), el 86.46 % emplearon dicha técnica. No obstante, la simplicidad de estos métodos también los hace más didácticos y comunicativamente útiles (Burgass *et al.*, 2017). Sin embargo, en un contexto donde la sustitutabilidad sea muy perniciosa, los métodos no compensatorios (MCA) muestran una gran utilidad. En cambio, los métodos de agregación geométrica se encuentran a caballo entre sendas realidades, en tanto que permiten cierta compensación pero penalizan de forma notable grandes desajustes.

Básicamente, depende del autor juzgar qué método se ajusta mejor a los indicadores que maneja y los objetivos que pretende reflejar, lo cual ha de introducirse ya desde el propio marco teórico. Por supuesto, los análisis de robustez y sensibilidad deberán evaluar cuán determinante ha sido el método impuesto, lo cual resulta esencial para garantizar la fiabilidad del sintético construido.

2.7. Análisis de robustez

Las decisiones forman parte del núcleo de todo indicador sintético. El establecimiento del objetivo, la determinación de los subgrupos y la elección de cada uno de los criterios empleados para normalizar, ponderar y agregar suponen, en última instancia, fuentes de incertidumbre que menoscaban la fiabilidad y la confianza puesta en los indicadores (Burgass *et al.*, 2017). Cada etapa es susceptible de ser criticada, por lo que la transparencia en las decisiones ha de ser vital para elaborar el indicador y debe tratarse de incluir una variedad de razonamientos suficiente que represente los intereses existentes. Ahora bien, sí existe una forma de evaluar cuán influyentes son los juicios del investigador (OCDE, 2008), especialmente en la parte más práctica de la construcción: el análisis de robustez, que permite alterar cada uno de los pasos para medir cómo reacciona el resultado final. Por supuesto, esto es poco realista para etapas

como el marco teórico o la selección de datos, ya que alterarlos vendría a ser como rediseñar el sintético completo. La falta de este tipo de análisis es común en la práctica y, de hecho, es una de las fuentes de crítica para los detractores de los sintéticos (Greco *et al.*, 2019).

Entrando en materia, la robustez se puede medir a través de la incertidumbre y de la sensibilidad. El primer caso consta de dos pasos: decidir qué partes de la elaboración se alterarán ((Saisana y Tarantola, 2002)) y traducirlos en factores que permitan la comparación y visualización de los cambios conseguidos. El análisis de sensibilidad, al contrario, monitoriza exclusivamente el impacto sobre cada caso sin atender al resultado general del sintético (Greco *et al.*, 2019). No obstante, ambos procesos son importantes para garantizar la robustez de la elaboración.

El análisis de incertidumbre debe evaluar cada paso de la construcción, tratando de identificar cualquier tipo de problema y midiendo su transparencia (Cherchye *et al.*, 2008). Así, se dará una parte más cualitativa que valorará los pasos no matemáticos del indicador en busca de la mencionada transparencia, la consistencia con el criterio experto, la elaboración de matrices de pedigrí (Van Der Sluijs *et al.*, 2005)... A este respecto, el trabajo de Burgass *et al.* (2017) arroja luz sobre todas las fuentes de incertidumbre que se dan en el proceso, así como algunas claves para su identificación y subsanación. Entre ellas, se destaca el involucrar a tantos *stakeholders* como sea posible, analizar estadísticamente la coherencia y robustez de los datos y, finalmente, ser comunicativo y didáctico en la divulgación de resultados, de manera que se minimice el salto interpretativo que ocurre entre autor y lector.

Ahora bien, desde el punto de vista del análisis más estadístico, el manual de la OCDE (2008) establece que, para cada decisión que afecte a los datos, han de establecerse factores de tal forma que cada incertidumbre se corresponda con uno. Así, se obtendría una lista como la siguiente:

$X_1 \rightarrow$ Imputación	$X_2 \rightarrow$ Normalización	$X_3 \rightarrow$ Exclusión de indicador
$X_4 \rightarrow$ Agregación	$X_5 \rightarrow$ Ponderación	$X_6 \rightarrow$ Elección de experto

Habrán factores que describan técnicas empleadas ($X_1 = 1$ sería IM, $X_1 = 2$ sería EM, por ejemplo) y otros que elijan a través de una distribución aleatoria qué indicador descartar del sintético ($X_3 = 10$ implicaría descartar el décimo indicador componente). El último factor escogería al experto que asignaría los pesos dependiendo del método de ponderación obtenido en X_5 (Saisana y Tarantola, 2002). De esta manera, una generación aleatoria de factores proporcionará un sintético diferente, cuya variación debe ser medida. Esto se repite tantas veces como se desee y se puede representar con matrices de gráficos que representen de forma visual los cambios en el sintético final, así como en ciertos casos particulares. Su uso está muy extendido, ya que no exige grandes esfuerzos y resulta fácil de interpretar (Floridi *et al.*, 2011).

El análisis de sensibilidad basado en la varianza es un método muy empleado en la práctica (Becker *et al.*, 2017). Los sintéticos no se consideran modelos lineales, puesto que se da una toma de decisiones que provoca que, aunque la agregación sea lineal, su proceso de elaboración no lo sea. Estas técnicas son muy fiables a la hora de evaluar la elaboración del sintético, puesto que pueden valorar la totalidad de factores involucrados y proporcionan resultados muy sencillos de interpretar (OCDE, 2008). Básicamente, explican cuánta varianza del sintético final se ve afectada por un cambio en uno de los factores, *ceteris paribus*. No obstante, al igual que en el caso de la agregación lineal, cabe la posibilidad de que la existencia de un factor dado

altere cómo otro afecta al resultado final (independencia preferente), lo cual se puede subsanar mediante la utilización de índices de sensibilidad total (Greco *et al.*, 2019).

Estas técnicas emplean programas de computación avanzados, que directamente proporcionan los índices que muestran el impacto sobre la varianza. Además, cabe la posibilidad de fijar externamente ciertos valores, ya sea porque estos sean dictados por alguna autoridad externa o por decisión del autor. Existen, por tanto, mecanismos que permiten la fijación de ciertos factores de acuerdo con su comportamiento sobre el sistema, tal y como describe el manual de la OCDE (2008).

Además, estos análisis basados en la varianza constan de múltiples ventajas, tales como la consideración de interacción entre factores, el estudio de todos los factores de forma simultánea (lo cual se puede representar en gráficas, garantizando así la transparencia y claridad del indicador) y su independencia como modelo, aplicable a los indicadores sintéticos, que no son lineales (Saisana *et al.*, 2005). Por supuesto, existen otros métodos, tales como el *One-at-a-time*, que altera solo una de las variables *ceteris paribus*, métodos regresivos e incluso técnicas derivativas. Sin embargo, dada la no linealidad del proceso de los sintéticos, estas técnicas no se recomiendan.

En resumen, los análisis de varianza permiten esclarecer qué partes de la elaboración han tenido mayor impacto sobre el sintético final o, en su defecto, sobre su varianza (Riedler *et al.*, 2015). De esta manera, si este análisis muestra que el sintético es muy sensible a cierta técnica o indicador, el autor puede replantearse dicha etapa de la construcción, especialmente si dicho resultado no es coherente con lo esperado. Desde el punto de vista de la incertidumbre, todo es susceptible de ser fuente de la misma y generar criticismo hacia el indicador. En este caso, ciertos análisis empíricos han de realizarse *a posteriori*, pero también es esencial considerar y justificar cada decisión de la forma más multidisciplinar y argumentada posible durante el propio proceso de creación, de manera que los aspectos más cualitativos como el marco teórico puedan dotarse de cierto consenso de acuerdo con las partes interesadas (Burgass *et al.*, 2017).

2.8. Divulgación

De acuerdo con Nardo *et al.*, la forma en que se presentan los resultados del indicador no debe minusvalorarse. Al fin y al cabo, buena parte de los objetivos de los indicadores sintéticos radica en la descripción de un problema y su monitorización, de forma que se incentive la puesta en marcha de los *stakeholders* que, de una manera u otra, disponen de cierta capacidad de acción para subsanarlo. De esta manera, la posibilidad de comunicar a las partes interesadas los hallazgos obtenidos es, probablemente, una parte crítica del proceso de elaboración. Esto no debería involucrar solo a las entidades públicas y privadas, sino también al ciudadano medio que motiva su comportamiento. Dada la compleja red social a través de la cual unos y otros se relacionan, la divulgación eficaz del sintético debería potenciar su alcance de transformación social. De hecho, en muchas ocasiones, los indicadores sintéticos han contribuido a la movilización social y a la mejora de la realidad socioeconómica de ciertos países a través de la concienciación y la aplicación de políticas, como en el caso del IDH (PNUD, 1990).

De esta manera, existen múltiples canales de comunicación que permiten difundir el indicador de forma acertada, tales como la elaboración de artículos académicos, la realización

de tesis o trabajos de fin de grado e incluso informes de utilidad pública. El advenimiento de la era tecnológica permite la exploración de nuevos mecanismos de divulgación, que acerquen la realidad a una audiencia más amplia y diversa. Por supuesto, esto dependerá del público que se desee abordar, pero la presentación de resultados a través de *dashboard* interactivos, vídeos u otros recursos visuales debe considerarse para lograr la mayor difusión posible. Asimismo, la utilización de herramientas como tablas, gráficos o *rankings*, que visualmente sean atractivos y didácticos, sigue a la orden del día (Nardo *et al.*, 2005).

Asimismo, en pos de reducir la divulgación como una fuente más de incertidumbre (Burgass *et al.*, 2017) para el sintético –en tanto que puede haber una brecha entre lo presentado por el autor y lo interpretado por el lector–, esta fase debe cuidar al máximo la forma en que se transmiten los resultados. Al fin y al cabo, una de las críticas recibidas por los indicadores sintéticos es su tendencia a simplificar excesivamente realidades complejas y cómo esto puede resultar pernicioso para las mismas. Para evitarlo, una divulgación que aporte información de calidad sobre las conclusiones y también sobre los mecanismos que han permitido llegar a ellas es esencial (Profit *et al.*, 2010), puesto que la actitud de los lectores cambia de acuerdo con la forma en que se les describe la problemática.

Así, la visualización y presentación de resultados debe contar con la posibilidad de acceder a los detalles y pormenores de la elaboración del sintético. La transparencia, al igual que en otros momentos del proceso de construcción, debe brillar en fase de divulgación. Tras la publicación, todos los objetivos marcados por el autor se completarán, lo que da también fin al propio proceso de elaboración del indicador sintético en cuestión.

De esta manera, el capítulo 2 referido al estado de la cuestión se da por concluido, de forma que queda establecido un nuevo marco teórico general que facilite la labor de futuros investigadores y autores de indicadores sintéticos.

Capítulo 3

Indicador sintético 360 Smart Vision sobre el mercado laboral español

LA construcción del indicador es, por supuesto, parte crítica del proyecto. La situación socioeconómica en España, en un contexto de incertidumbre y recuperación tras la irrupción de la pandemia del COVID-19 –tal y como fue presentado en el capítulo 1–, centrará el interés del sintético. De esta manera, se pretende dotar a las entidades públicas y privadas de una herramienta de análisis que les permita abordar de forma más eficaz cambios para la transformación social. No obstante, un indicador sintético que agrupe todos los fenómenos subyacentes a la recuperación socioeconómica sería prácticamente inabarcable. Por ello, el indicador a construir describirá una realidad que, si bien está más acotada, también reviste gran complejidad. Dicha realidad es la siguiente:

La salud del mercado laboral español

Tal y como se ha explicado previamente, el establecimiento del marco teórico en torno al cual se articularán todas las futuras decisiones del autor da comienzo a su elaboración.

3.1. Marco teórico aplicado

El establecimiento del marco teórico es esencial para evitar ambigüedades y situar a lector y autor en el mismo contexto. En primer lugar, se han de definir los conceptos descritos por el indicador, para posteriormente describir las categorías o subgrupos que lo compondrán. Tras estos determinar estos parámetros, se ha de presentar la estructura anidada del sintético, de forma que el lector comprenda fácilmente el impacto de cada una de las categorías sobre el indicador final. Finalmente, se describirá el criterio que establezca qué datos son considerados relevantes para la elaboración. Por supuesto, habrán de satisfacer los requisitos mínimos de calidad establecidos por la Comisión Europea (Eurostat, 2017) y ya mentados en la sección 2.1. De esta manera, tal y como se ha mencionado, el primer paso vendría a ser demarcar el concepto del mercado de trabajo.

Como todo mercado, el mercado laboral es fruto de la unión entre oferta y demanda, en este caso referidas al factor productivo del trabajo (Expansión, 2022). Este factor productivo podría denominarse también capital humano, en tanto que cada persona dispone de una serie de habilidades comercializables que además pueden ser adquiridas mediante el aprendizaje (Acemoglu y Autor, 2011). De esta manera, su formación y sus características serán claves

a la hora de dotarle de cierta productividad que, en cierta manera, repercutirá sobre su salario. Este salario sería la contraprestación proporcionada por un empleador a cambio de sus horas trabajadas. Por supuesto, como con todo fenómeno humano, el equilibrio entre oferta y demanda es imperfecto, por lo que se da la aparición de desempleo y salarios indignos (Mtessigwa, 2013). Esto puede deberse a multitud de factores, tales como ineficiencias en el mercado, competitividad empresarial, duración de la jornada laboral, asimetría en la información de que disponen los agentes, circunstancias estructurales que deriven en fenómenos estacionales...

Las políticas de empleo surgen, por tanto, con la intención de abordar dichas ineficiencias y apoyar a aquellas personas que, por un motivo u otro, no forman parte del esquema funcional del mercado laboral. En este sentido, cabe destacar la existencia de los agentes sobre los cuales versa el mentado equilibrio y muy importantes para comprender su funcionamiento: empleados y empleadores, tanto públicos como privados, que son los que forman parte de la población ocupada. Asimismo, la población activa desempleada, en edad y con voluntad de trabajar, compone también el mercado laboral. Su existencia se debe a las mencionadas ineficiencias, por lo que su protección es uno de los deberes del Estado (artículo 41 de la Constitución Española).

Así pues, el fenómeno referido al mercado de trabajo es complejo. Existen innumerables elementos influyendo sobre él, algunos comprendidos y otros ignotos o discutidos por los académicos. Ahora bien, considerando el objetivo del proyecto en lo referido a este concepto, existen varios fenómenos descritos que el indicador sintético debe recoger. Por ello, a lo largo del presente artículo, la salud del mercado laboral español aludirá a la situación en que *los agentes del mercado ven satisfechas sus necesidades de forma sostenible en el tiempo*. En este sentido, dado que se pretende elaborar un indicador sintético que estudie el mismo fenómeno (mercado laboral) en un espacio dado (España) en diferentes momentos del tiempo, este sintético será una serie temporal.

Tras la definición del término en cuestión, el siguiente paso es la determinación de los subgrupos. En este sentido, si bien una importante cantidad de factores afectan al mercado laboral tal y como está definido, sí se pueden identificar algunos elementos esenciales para su buen funcionamiento, aludiendo en primer lugar a las tres masas sociales descritas anteriormente: el salario y los costes laborales, contraprestación pagada por los empleadores a los empleados; el capital humano, como la prestación productiva que proporciona todo empleado basado en su formación y su capacitación; y, también, la protección a los desempleados, que salvaguarda el bienestar de estos agentes del mercado laboral cuyas necesidades no son atendidas por el mismo. Así pues, cada una de estas dimensiones alude a los agentes involucrados en el mercado laboral y cómo sus necesidades se ven cubiertas.

Además, una dimensión adicional será esencial para que estas necesidades sean sostenibles en el tiempo: la caracterización del empleo, en tanto que se describa la estructura contractual y el número de afiliados, lo cual habrá de ser complementado por el fenómeno del desempleo. Ambos conceptos contribuyen a explicar cuán competitivo y sostenible es el mercado laboral en el tiempo. A fin de cuentas, el estudio de los salarios y la capacitación de los agentes involucrados no goza de impacto suficiente sin un análisis sobre la cantidad de ellos que pueden participar del mercado laboral. Unas tasas adecuadas garantizan además la sostenibilidad de las ayudas al desempleo, así como que la capacitación y los salarios repercutan sobre los agentes de forma beneficiosa. De esta manera, el mercado laboral queda plenamente acotado.

De esta definición se deducen, por tanto, una serie de subgrupos cuyo bienestar debe ser perseguido por todos los agentes, de manera que se puedan abordar como objetivos independientes sobre los cuales se mida el impacto de políticas o iniciativas privadas:

- **Desempleo:** Esta categoría identificará el fenómeno del desempleo en España desde una perspectiva completa, incluyendo por tanto realidades debidas al COVID-19 tales como el auge de los ERTes y su implicación en las tasas de paro convencionales. Así, se caracterizará la población activa desempleada que surge fruto de las ineficiencias del mercado laboral. Este subgrupo es, al mismo tiempo, un importante reflejo de su situación y también una de las consecuencias de otros fenómenos tales como la capacitación del factor humano o la estructura salarial en España.
- **Empleo:** Por otro lado, la caracterización del empleo se torna esencial para comprender las complejidades del mercado laboral. De esta manera, se integrarán indicadores referidos a la cantidad de población ocupada, la productividad efectiva de dicha población y también el tipo de contratos imperante. Al igual que el desempleo, la situación del empleo depende de otros factores del mercado, como por ejemplo la capacitación, los salarios mínimos o los incentivos que proporcione el Estado al empleo, siendo por tanto una consecuencia de la salud del mercado laboral.
- **Habilidades y capacitación:** Esta dimensión alude a los estudios y la formación de la población activa en España, en tanto que representa la posible prestación que un empleado cualquiera puede proporcionar a su empleador. Este factor goza de especial relevancia a la hora de configurar los mercados laborales, ya que es una de las causas de una economía competitiva y sana (Acemoglu y Autor, 2011). Además, su posible mejoría tendría un importante impacto social, por lo que goza de especial relevancia.
- **Protección a los desempleados:** La forma en que el Estado apoya a aquellos cuyas necesidades no son debidamente satisfechas por el mercado laboral es, sin duda, digna de estudio y análisis. Al fin y al cabo, el desempleo es un gran peligro social, puesto que está íntimamente ligado a la pobreza. De hecho, algunos autores lo vinculan a la llamada *trampa de la pobreza* (Rosas y Jiménez-Bandala, 2018). Básicamente, en un contexto de pocos recursos económicos y escasa formación (tal y como podría ser un desempleado), las posibilidades de reinserirse en el mercado laboral se reducen de forma notable (Arnal *et al.*, 2013).

Por ello, el Estado busca paliarlo mediante las ayudas al desempleo, a la vez que garantiza unas condiciones de vida mínimas. De esta manera, cuán bien son protegidas las personas desocupadas forma parte de las dinámicas del mercado laboral, en tanto que facilita su futura reinserción e incluso la adquisición de nuevas competencias. Por ello, es otra de las causas que influye sobre la salud del mercado laboral.

- **Salarios y costes laborales:** La contraprestación otorgada por empleados a empleados constituye esta dimensión. Multitud de factores influyen sobre ella, como por ejemplo la capacitación, la existencia del salario mínimo interprofesional o determinados derechos laborales. Aun así, este fenómeno tiene una influencia notable sobre la salud del mercado

de trabajo, puesto que la obtención del salario es la principal meta de los empleados y uno de los costes más importantes para empleadores, lo que desde luego desempeña un impacto más que notorio sobre el desempleo (Marimon, 1997). Por ello, además de constituir una de las causas de la salud del mercado laboral, también es uno de los principales objetivos que gobiernos y agentes sociales buscan mejorar. Hallar por tanto un subindicador que describa y permita abordar el fenómeno reviste gran importancia.

Estos subgrupos, categorías u objetivos describirán el comportamiento de la *salud del mercado de trabajo*, garantizando por tanto que los agentes sociales dispongan de información pormenorizada, transparente y trazable sobre la elaboración del sintético final.

Tras la identificación de los subgrupos, el estado de la cuestión (sección 2.1) compele a asumir un criterio de selección de datos según unos preceptos claros de calidad. Tal y como se ha explicado previamente, los datos procederán de la base de datos de 360 Smart Vision (Deloitte, 2021), un *dashboard* de origen público-privado que aspira a proporcionar indicadores fiables sobre la situación socioeconómica española, en especial tras la irrupción del coronavirus. Sus datos provienen, en muchos casos, de fuentes estadísticas oficiales, tales como el SEPE, el INE, el Ministerio de Economía o Eurostat (perteneciente a la Comisión Europea). Asimismo, también se emplean datos obtenidos por la empresa *Jobandtalent*, una compañía centrada en la inserción laboral y el emparejamiento de empresarios y futuros empleados. Los datos empleados por la plataforma ya pasan un primer control de calidad, en tanto que organismos públicos deben seguir unos estrictos protocolos de recolección de información y a las entidades privadas, centradas en el ánimo de lucro, deben disponer de bases de datos fiables y contrastadas para ejercer una buena praxis en sus operaciones.

De esta manera, partiendo de la extensa red de indicadores proporcionados por 360 Smart Vision, se emplearán aquellos que dispongan de las siguientes características: pertinencia, de forma que aporten información para la caracterización del subgrupo en que se enmarquen; circunscripción espacio-temporal, en tanto que describan el contexto pandémico y la posterior recuperación en España; y cierta completitud temporal (precepto ya incluido en el EPI (Wolf *et al.*, 2022)), excluyendo a aquellos indicadores componentes que superen cierto umbral de información ausente.

El resto de criterios mencionados en la sección 2.1 se ven cumplidos por su inclusión en el *dashboard*, en tanto que este cuida con sumo detalle la precisión y fiabilidad de sus indicadores, a la par que muestra claridad en sus descripciones y libre acceso gracias a su publicación en internet. Por ello, aunque las matrices de pedigrí no sean necesarias ni factibles (por la gran cantidad de indicadores a analizar), sí cabría hacer un análisis preliminar de los datos de los indicadores para estudiar su coherencia interna. Al fin y al cabo, la responsabilidad de garantizar la fiabilidad en la elaboración del indicador recae sobre este proyecto, por lo que han de tomarse medidas a tal efecto. Por estas razones, los criterios presentados permitirán cumplir los preceptos de la buena práctica estadística y, sobre todo, los del sentido común.

Así, la estructura anidada del indicador ha sido descrita y los criterios de selección debidamente presentados. Como conclusión, la siguiente figura 6 sintetiza la composición del sintético *salud del mercado laboral español*, a la vez que describe los cinco objetivos marcados y sus principales indicadores componentes:



Figura 6. Estructura anidada del sintético junto a los subgrupos e indicadores componentes principales

Los indicadores de los subgrupos han sido resumidos en base a características comunes, de manera que la visualización de la figura 6 sea más sencilla y didáctica. Cada indicador de la plataforma 360 Smart Vision es enmarcado en uno de los subgrupos presentados, que además representan fenómenos abordables por los agentes involucrados. De esta manera, se pretende incentivar la transformación social y la monitorización del impacto de las políticas e iniciativas implantadas *ad hoc*. No obstante, una lista detallada de los indicadores componentes se presenta en la tabla 24 presente en el anexo B, de forma que se muestran también los segregados por sexo y por sector de ocupación que tendrán cierto impacto sobre el indicador.

De esta manera, quedando el marco teórico totalmente descrito, se da paso a la recolección de los datos provenientes de la plataforma 360 Smart Vision.

3.2. Datos de 360 SMART VISION

La estructuración, recolección y selección de datos a partir de la plataforma ocupará buena parte del presente análisis. La estructuración alude al proceso mediante el cual se ordena la información para ser tratada posteriormente mediante técnicas de imputación de información ausente y demás fases de la elaboración. La información relativa a la descripción de cada indicador individual se encuentra recogida en la plataforma 360 Smart Vision (2021).

En este caso, estructuración y selección de datos ocurren de forma simultánea, puesto que la aplicación de ciertos criterios de selección es inmediata, como por ejemplo descartar aquellos indicadores componentes cuya medición se circunscriba a las comunidades autónomas o a otros países europeos. Este proceso de selección queda reflejado en el código 1 del anexo D, que depura de esta manera aquellos datos que estén circunscritos al mismo espacio geográfico. En cuanto a su carácter temporal, el indicador tiene por objetivo describir la recuperación del mercado laboral durante la pandemia del coronavirus, por lo que enero de 2020 supondrá el albor del indicador sintético.

Esto se debe a dos razones: primeramente, el caso cero de COVID-19 registrado en España ocurrió el 31 de enero de 2020 (RTVE, 2022), por lo que el indicador proporcionará una contextualización prepandémica (enero y febrero de 2021) para establecer comparativas en la serie temporal y garantizar la fiabilidad del indicador. En segundo lugar, la mayoría de indicadores proporcionados por la plataforma dan comienzo con la pandemia, puesto que 360 Smart Vision nace con el objetivo de representar la realidad socioeconómica en España tras el coronavirus. Por ello, tratar de construir un indicador anterior reduciría considerablemente la base de datos, lo que implicaría la exclusión de variables relevantes y la consecuente disminución de la fiabilidad del sintético.

Asimismo, garantizar una misma circunscripción temporal implica adaptar las diferentes “granularidades” de las variables, de manera que todas se muevan en las mismas frecuencias. Para ello, es especialmente relevante la tabla 24 presentada en el anexo B, donde se muestra la frecuencia de sus datos desde la mensual hasta la diaria. La práctica totalidad de los mismos presentan una distribución mensual, por lo que este será el intervalo de tiempo elegido para describir el sintético. Al fin y al cabo, elegir una granularidad trimestral desecharía información muy valiosa y reduciría enormemente la utilidad del sintético, mientras que establecerla semanal o diaria implicaría la aparición de una cantidad ingente de datos ausentes. Por ello, estas dos opciones han de descartarse.

Además, a la par que se aplican los preceptos de selección, el código 1 del anexo D también depura los datos originales, cuya estructura inicial dista de ser la deseada para cualquier tratamiento posterior. Una instantánea de la misma se presenta a través de la figura 7:

Mercado de Trabajo	Desempleo	Solicitudes / Usuarios Activos por Mes Transporte	España	Total	2019-09-01	0.6120815
Mercado de Trabajo	Desempleo	Solicitudes / Usuarios Activos por Mes Transporte	España	Total	2019-08-01	0.4227799
Mercado de Trabajo	Desempleo	Solicitudes / Usuarios Activos por Mes Transporte	España	Total	2019-07-01	0.6278149
Mercado de Trabajo	Desempleo	Solicitudes / Usuarios Activos por Mes Transporte	España	Total	2019-06-01	0.5540507
Mercado de Trabajo	Desempleo	Solicitudes / Usuarios Activos por Mes Transporte	España	Total	2019-05-01	0.7842574
Mercado de Trabajo	Desempleo	Solicitudes / Usuarios Activos por Mes Transporte	España	Total	2019-04-01	0.7845697
Mercado de Trabajo	Desempleo	Solicitudes / Usuarios Activos por Mes Transporte	España	Total	2019-03-01	1.0639464
Mercado de Trabajo	Desempleo	Solicitudes / Usuarios Activos por Mes Transporte	España	Total	2019-02-01	0.85199
Mercado de Trabajo	Desempleo	Solicitudes / Usuarios Activos por Mes Transporte	España	Total	2019-01-01	1.2650093
Mercado de Trabajo	Desempleo	Tasa de paro ajustada Eurostat	España	Total	2021-06-01	0.151
Mercado de Trabajo	Desempleo	Tasa de paro ajustada Eurostat	España	Total	2021-05-01	0.154
Mercado de Trabajo	Desempleo	Tasa de paro ajustada Eurostat	España	Total	2021-04-01	0.156
Mercado de Trabajo	Desempleo	Tasa de paro ajustada Eurostat	España	Total	2021-03-01	0.154
Mercado de Trabajo	Desempleo	Tasa de paro ajustada Eurostat	España	Total	2021-02-01	0.157
Mercado de Trabajo	Desempleo	Tasa de paro ajustada Eurostat	España	Total	2021-01-01	0.161
Mercado de Trabajo	Desempleo	Tasa de paro ajustada Eurostat	España	Total	2020-12-01	0.162
Mercado de Trabajo	Desempleo	Tasa de paro ajustada Eurostat	España	Total	2020-11-01	0.164
Mercado de Trabajo	Desempleo	Tasa de paro ajustada Eurostat	España	Total	2020-10-01	0.163
Mercado de Trabajo	Desempleo	Tasa de paro ajustada Eurostat	España	Total	2020-09-01	0.164
Mercado de Trabajo	Desempleo	Tasa de paro ajustada Eurostat	España	Total	2020-08-01	0.159
Mercado de Trabajo	Desempleo	Tasa de paro ajustada Eurostat	España	Total	2020-07-01	0.154
Mercado de Trabajo	Desempleo	Tasa de paro ajustada Eurostat	España	Total	2020-06-01	0.153
Mercado de Trabajo	Desempleo	Tasa de paro ajustada Eurostat	España	Total	2020-05-01	0.152
Mercado de Trabajo	Desempleo	Tasa de paro ajustada Eurostat	España	Total	2020-04-01	0.155
Mercado de Trabajo	Desempleo	Tasa de paro ajustada Eurostat	España	Total	2020-03-01	0.15
Mercado de Trabajo	Desempleo	Tasa de Paro Trimestral	Unión Eur	Total	2021-06-30	7.1
Mercado de Trabajo	Desempleo	Tasa de Paro Trimestral	España	Total	2021-06-30	15.1

Figura 7. Muestra de los datos originales proporcionados por la plataforma

De esta manera, el código debía realizar una doble labor: organizar y seleccionar. Organizar los datos consiste en clasificarlos por los subgrupos descritos en la sección 3.1, de manera además que se homogeneice el vector temporal de cada una de las categorías. Por ello, el código mentado recorre el conjunto de datos en bruto y genera matrices de datos donde el primer punto temporal es el mencionado enero de 2020, mientras que el último lo proporciona el indicador componente con la información más reciente posible. Esto contribuirá a mantener la mayor cantidad de datos posibles, en tanto que no se descarten por no pertenecer a un vector de fechas elegido arbitrariamente. De esta manera, se generan cinco matrices de datos con los indicadores

componentes ordenados, algunos de los cuales se han de descartar de acuerdo a la aplicación de los criterios de calidad impuestos en el marco teórico.

Para ello, se aplica el criterio de pertinencia mencionado anteriormente, puesto que se eliminan tres indicadores irrelevantes para la caracterización de los subgrupos y del sintético final, a saber: “jornada media pactada”, cuya aplicación efectiva es poco fiable y, además, no presenta variaciones relevantes (un 0.1 % anual entre 1996 y 2003 de acuerdo con Benito *et al.* (2004)); “solicitudes / usuarios activos” y “vacantes disponibles / usuarios activos”, que son indicadores proporcionados por la empresa *Jobandtalent* que reflejan información de interés para sus operaciones, pero no suficientemente pertinente para la caracterización del mercado laboral. Estas decisiones están respaldadas por el criterio experto de los doctores economistas que colaboran en la plataforma 360 Smart Vision. La aplicación de este principio también conlleva cierta controversia, puesto que el uso de una cantidad tan importante de variables puede llevar a la redundancia de la estructura de datos y, finalmente, ser problemática. No obstante, dado que no existe una relación causal perfecta *ex-ante* entre ninguno de los indicadores presentes, desecharlos podría considerarse poco justificado, por lo que una aplicación conservadora del criterio de pertinencia prevalece.

Además, rememorando lo explicado en la sección 2.2, existen autores que han introducido umbrales que garantizan un mínimo de datos presentes para cada indicador (Sainani, 2015), lo que promovería la aplicación del precepto de la completitud temporal. Este criterio también se implementó en el EPI (Wolf *et al.*, 2022). Así pues, se ha decidido implantar un mínimo de un 70 % de los casos para considerar a los indicadores aptos para su inclusión, sin incluir a los que presenten una estructura de datos trimestrales, puesto que la ausencia de información forma parte de la propia naturaleza del propio indicador. De acuerdo con Horton y Kleinman (2007), esto solo debe realizarse si los indicadores son tangenciales al propósito del proyecto, por lo que se ha estudiado qué indicadores de la tabla 24 del anexo B serían desechados de acuerdo con este umbral.

En primer lugar, serían excluidos los indicadores relativos a los “ERTEs” y derivados, que no reflejan estrictamente el fenómeno del desempleo. Al fin y al cabo, las personas en ERTE no trabajan, pero tampoco se consideran desempleados (Gómez Bengoechea, 2020). Su exclusión podría ser problemática si no existieran otras variables, como la denominada “Impacto en el mercado laboral (ICADE)”, que en cierto modo integran los ERTEs como mecanismos causados por una situación disfuncional del mercado de trabajo. Efectivamente, ninguno de estos indicadores relativos a los ERTEs cumple el porcentaje umbral establecido en el 70 %, ya que la serie cesó su conteo en enero de 2021. Por ello, dado que incumple el mínimo de datos existentes y no son esenciales para caracterizar el *Desempleo*, su exclusión se torna efectiva.

En segundo lugar, destacan los indicadores “Ocupación crisis 1976 / 1991 / 2007”, claramente no aplicables a un indicador sintético sobre la pandemia. Aunque la comparativa con las tasas de ocupación de la presente crisis puedan resultar interesantes, su inclusión no casa con el objetivo del proyecto de crear un indicador sintético que refleje una serie temporal que comience en 2020, puesto que todo son datos ausentes. Por ello, además de incumplir el umbral establecido, sus datos están muy lejos de ser oportunos, por lo que también se excluyen.

Así, en pos de salvaguardar también el criterio de claridad y accesibilidad al proceso realizado, se incluirán instantáneas o tablas que conforman las principales etapas en la producción del sintético. En líneas general, solo ejemplos característicos o didácticos se

presentarán en la memoria, de forma que el resto de tablas o recursos permanezcan en el anexo. Así, se fomenta la accesibilidad libre a los datos sin generar un perjuicio a la claridad del proyecto. Por ello, una instantánea de la tabla que agrupa los indicadores presentes en la categoría *Empleo* se presenta a continuación:

'Fecha'	'Altas Afiliados'	'Bajas Afiliados'	'Contratos Registrados'	'Contratos Registrados a Tiempo Parcial'	'Contratos Registrados CVE'	'Contratos Registrados Indefinidos'	'Contratos Registrados Interinidades'	'Contratos Registrados Temporales'	'Ocupación crisis 2020'	'Productividad por hora efectivamente trabajada'	'Productividad por puesto de trabajo a tiempo completo'	'Total Afiliados'
'Aug-2021'	84384	100604	1407563	473494	1712542.3	118985	59640	1288578	NaN	NaN	NaN	19473724
'Jul-2021'	105761	95055	1838250	675631	1634109.8	165500	76907	1672750	NaN	NaN	NaN	19591728
'Jun-2021'	109591	109614	1798047	657589	1558935.0	172866	68264	1625181	98.522	-6.986	0.740	19500277
'May-2021'	99488	87156	1545308	518699	1513358.0	156148	63912	1389160	97.745	-4.876	-0.308	19267221
'Apr-2021'	85426	75258	1356845	441024	1397362.9	164080	62841	1192765	96.969	-2.765	-1.356	19055298
'Mar-2021'	78293	81871	1404107	453042	1475289.6	207191	62590	1196916	96.193	-0.655	-2.404	18920902
'Feb-2021'	74705	76258	1212284	339408	1459781.0	132431	63754	1079853	96.423	-1.403	-2.873	18850112
'Jan-2021'	93691	97518	1302429	358990	1451485.2	124191	66557	1178238	96.652	-2.151	-3.343	18829480
'Dec-2020'	87074	88323	1355147	401285	1381692.3	111822	66636	1243325	96.882	-2.899	-3.812	19048433
'Nov-2020'	76004	84164	1449810	419228	1512946.5	128189	71882	1321621	96.374	-2.816	-3.676	19022002
'Oct-2020'	96898	89649	1551357	528183	1352808.7	152319	76952	1399038	95.866	-2.734	-3.539	18990364
'Sep-2020'	143821	120663	1632484	542413	1408332.0	163209	74131	1469275	95.358	-2.652	-3.403	18862713
'Aug-2020'	71456	76187	1118663	353195	1419365.0	96275	53965	1022388	94.636	-0.576	-3.732	18792376
'Jul-2020'	87149	79819	1536122	531513	1312696.0	141105	67421	1395017	93.913	1.500	-4.062	18785554
'Jun-2020'	74456	80720	1159602	362772	1003079.0	114393	47439	1045209	93.190	3.576	-4.392	18624337
'May-2020'	55230	45520	850617	195131	832856.0	76692	32902	773925	94.983	2.438	-4.105	18556129
'Apr-2020'	47802	49994	673149	134515	717819.0	59042	38898	614107	96.776	1.300	-3.819	18458667
'Mar-2020'	80012	118714	1256510	397997	1327930.0	145393	64815	1111117	98.570	0.162	-3.533	19006760
'Feb-2020'	102118	91422	1594763	556975	1896847.0	178193	72302	1416570	NaN	NaN	NaN	19250229
'Jan-2020'	86435	123249	1764837	566811	1810956.0	178978	76702	1585859	NaN	NaN	NaN	19164494

Figura 8. Instantánea de la tabla de estructuración y selección de datos en el subgrupo *Empleo*

El resto de tablas se encuentran en el anexo C, disponibles para el lector. Los indicadores mostrados en la figura 8 ya han pasado, además, los controles debidos de fiabilidad y consistencia, principalmente por dos motivos: en primer lugar, como es sabido, por proceder de la plataforma 360 Smart Vision, garantía de buen proceder en la toma de datos; y, en segundo lugar, por la realización de un análisis de valores atípicos que revelaría la eventual presencia de errores en la información proporcionada, de forma que el código 1 del anexo D notifica el mes en que dicho valor extremo tiene lugar. Por supuesto, el análisis revela una cantidad notable de *outliers* en los meses posteriores a la irrupción del COVID-19, lo cual es coherente con el fenómeno descrito por los mismos. al fin y al cabo, la presencia de valores atípicos no siempre corresponde a problemas en los datos. Este análisis ratifica la solidez de la plataforma 360 Smart Vision, pues todos los indicadores muestran datos consistentes y coherentes en el tiempo.

De esta manera, todos los preceptos han sido presentados y aplicados. A continuación, se presenta la figura 9 que sintetiza las acciones emprendidas en la presente sección:

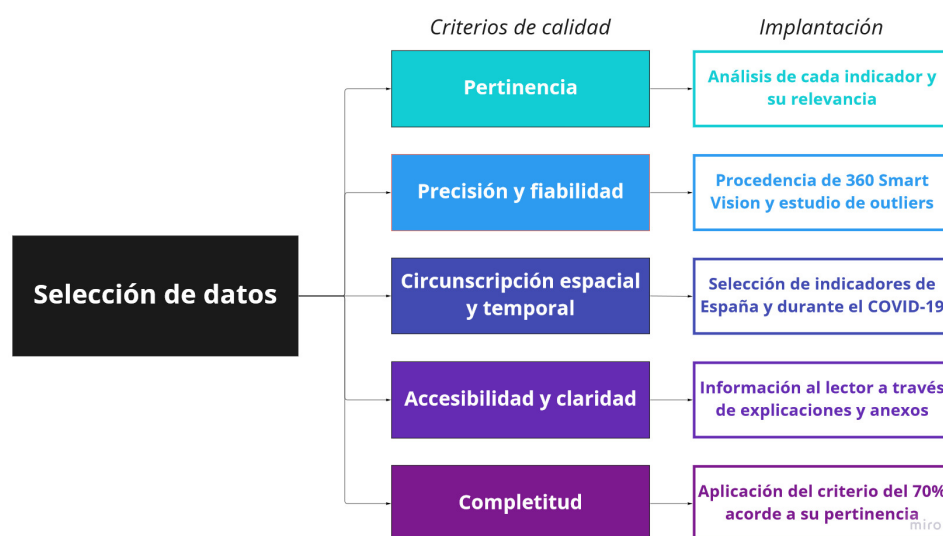


Figura 9. Resumen de los criterios de calidad aplicados y su implementación

De esta manera, se da por concluida la etapa de selección de datos del proceso de elaboración del indicador sintético, tal y como establece el manual de la OCDE (2008).

3.3. Imputación de información ausente

Tal y como se ha explicado a lo largo del estado de la cuestión, la información ausente es una fuente de incertidumbre importante a la hora de diseñar un indicador. El primer paso, tal y como es descrito por Nardo *et al.* (2005), es la determinación del tipo de ausencias que se producen en los datos, a saber: MCAR, MAR o NMAR. Las características de cada una aparecen en la sección 2.3. Como resumen, resulta esencial que las ausencias se deban a fenómenos MCAR o MAR, para los que los métodos de imputación producen resultados adecuados.

Sin embargo, antes de la realización del análisis MCAR de Little, que revela si se puede asegurar que los datos presentan la problemática distribución NMAR, resulta importante tratar los datos con frecuencia trimestral o diaria. En el segundo caso, no se trataría estrictamente de una imputación de datos faltantes, ya que su granularidad es mayor y, por tanto, para cada mes hay 30 medidas aproximadamente. Tal y como aparece en la tabla 24 del anexo B, solo las variables referidas al número de afiliados existentes o su variación diaria sufren de esta condición (“Total afiliados”, “Bajas afiliados” y “Altas afiliados”).

Debido a la naturaleza de estos indicadores, el uso de instantáneas no es recomendable, aunque sea eficiente. De hecho, pese a que el empleo de medias mensuales sea común en la literatura (Illes y Lombardi, 2013), el empleo del primer dato de cada mes también está recogido (Hernández Díaz y Marín Jaramillo, 2016). Sin embargo, incluso aunque los indicadores en cuestión suelen mostrar la principal variación el primer día de cada mes –por la creación de nuevos contratos, la conclusión de otros...–, en ocasiones esto también genera grandes desajustes en la estructura de datos. Al fin y al cabo, hay inicios de mes, como en mayo o en enero, en los que el número tanto de bajas como de altas de afiliados es irrisorio, debido al carácter festivo del día o incluso el periodo en que dicha fecha se enmarque. De esta forma, estos *outliers* inferiores compartirían serie con valores elevados típicos de meses normales, por lo que el indicador obtenido sería ciertamente errático. Por estas razones, se ha optado por el uso de medias mensuales que homogeneizan en cierta medida la serie de datos, de manera que no se introduzcan más atípicos en un conjunto de datos que, debido a la pandemia de COVID-19, ya está repleto de ellos.

En el caso de “Tasa de paro trimestral” y otras de granularidad trimestral, reveladas por la importante cantidad de datos ausentes y la propia descripción de los indicadores de 360 Smart Vision, se ha decidido realizar una interpolación convencional para su imputación. Este paso ya queda reflejado en la tabla 8 de la sección anterior, puesto que tiene lugar en la misma sección del código 1 del anexo D. Una ausencia de este tipo se ajusta a la naturaleza NMAR, puesto que existe una causa conocida e inherente a la propia variable que causa la existencia de datos faltantes. Así, tal y como enuncia Buuren (2018), las técnicas típicas de imputación no son adecuadas para tratar la muestra, ya que generan datos sesgados y alejados de la realidad.

Por ello, también se ha considerado la posibilidad de desechar las variables, pero algunas de ellas gozan de especial relevancia para caracterizar categorías como, por ejemplo, *Desempleo* y “Tasa de paro trimestral”. Las otras variables afectadas son “Ocupación crisis 2020”, “Productividad por hora efectivamente trabajada” y “Productividad por puesto de trabajo a

tiempo completo”. Si bien son pertinentes para el sintético, el sesgo de la interpolación será mínimo considerando que solo ocurre en cuatro indicadores. De esta manera, la interpolación lineal convencional imputa los valores mensuales que estaban ausentes, lo que implica un 66,6 % del conjunto total para esas variables, aproximadamente. Además, esta decisión no debería resultar especialmente problemática, ya que la interpolación no cambia tendencias, lo que salvaguarda el objetivo principal del indicador sintético y la información proporcionada por las variables originales.

Una vez estos pasos han sido realizados, procede llevar a cabo la imputación de datos ausentes. Para ello, el primer paso consiste en el análisis MCAR de Little (OCDE, 2008) de cada una de las matrices de datos correspondientes a cada categoría o subgrupo, las cuales ya han sido depuradas por la estructuración del código y la aplicación del criterio de completitud. Esta prueba, tal y como se menciona en la sección 2.3, arroja información sobre el tipo de ausencias. Dado que a menudo MCAR y MAR son prácticamente indistinguibles (Gómez *et al.*, 2017) –de hecho, la naturaleza MCAR se considera ideal, poco realista–, un test de Little permitirá dilucidar si la implantación de técnicas de imputación generará datos fiables o no. Por ello, se presentan en la tabla 7 los p-valores obtenidos a partir de los análisis en IBM SPSS de los subgrupos *Desempleo*, *Empleo* y *Habilidades y capacitación*, por ser los que presentan algún dato ausente:

	Desempleo	Empleo	Habilidades y capacitación
<i>p-valor</i>	0.653	0.185	0.352

Tabla 7. Prueba MCAR de Little de las categorías *Desempleo*, *Empleo* y *Habilidades y Capacitación*

De esta manera, todos los p-valor se encuentran por encima del nivel de significación más conservador (0.1 en este caso), por lo que la hipótesis nula de que los datos son de naturaleza MCAR no se debe descartar. Por tanto, los subgrupos cumplen los requisitos para ser sometidos a las técnicas de imputación. Afortunadamente, los subgrupos *Protección a los desempleados* y *Salarios y costes laborales* se encuentran completos. De esta manera, se ha de estudiar qué técnicas implementar en las categorías afectadas, para lo que conviene realizar un análisis preliminar de datos perdidos, el cual proporciona la tabla 8 para las variables con información ausente:

	N	μ	σ	Perdidos	
				Cant.	%
<i>Impacto en el mercado laboral ICADE</i>	16	0.217	0.035	4	20.0
<i>Tasa de paro ajustada Eurostat</i>	16	0.157	0.004	4	20.0
<i>Tasa de Paro Trimestral</i>	16	15.626	0.583	4	20.0
<i>Ocupación crisis 2020</i>	16	96.191	1.508	4	20.0
<i>Productividad por hora efectivamente trabajada</i>	16	-1.346	2.734	4	20.0
<i>Productividad por puesto de trabajo a tiempo completo</i>	16	-2.976	1.456	4	20.0
<i>Habilidades y capacitación</i>	16	-	-	2	11.1

Tabla 8. Análisis preliminar de las variables con datos ausentes

La tabla 8 también agrupa los indicadores de la categoría *Habilidades y capacitación* en su subgrupo, puesto que todos ellos presentan las mismas ausencias. De hecho, los datos de esta categoría proporcionados por la plataforma 360 Smart Vision aparecen por primera vez en marzo de 2020, con la irrupción de la pandemia. Asimismo, el análisis preliminar muestra un porcentaje de datos ausentes igual o inferior al 20 % (ratio perdidos/disponibles), por lo que sería un porcentaje manejable por las técnicas de imputación.

En cuanto al empleo de técnicas, la más básica y simplista la eliminación del caso (*case deletion*). Aunque no se considera una imputación propiamente, sí es una forma de lidiar con los datos ausentes, puesto que elimina los meses que no presenten información para la totalidad de las variables. Considerando que las matrices estudiadas disponen de 20 casos aproximadamente, este conjunto de técnicas resulta inapropiado, puesto que se reduciría dicho número de forma drástica. Esto mermaría la potencia estadística del sintético y quebraría la frecuencia mensual del sintético, menoscabando su utilidad. Por ello, este conjunto de métodos ha de descartarse.

Posteriormente, cabe plantearse la implantación de técnicas de imputación simple, donde el método de EM goza de la mejor reputación (Gold y Bentler, 2000); o múltiple, con el empleo de métodos MCMC que generan resultados muy robustos (OCDE, 2008). Dado que la muestra no es NMAR, cualquiera de estos procedimientos resulta apto, pero sus resultados serán diferentes dependiendo de la categoría concreta sobre la que se apliquen. Por ello, si bien el método de EM resulta el más eficiente y eficaz por su fácil implantación con IBM SPSS desde el punto de vista teórico, se ha optado por un análisis en torno al NRMSE de los datos (ecuación 2 de la sección 2.3) que se cerciore de emplear la técnica óptima de imputación de datos. Este procedimiento, ya explicado en el estado de la cuestión y presente en el manual de la OCDE (2008), se ha implantado al final del código 1 del anexo D.

En primer lugar, se han de generar matrices de datos completas, es decir, se debe realizar la eliminación de casos con información faltante. Tras esto, han de suprimirse un porcentaje de casos representativos y conocidos de la muestra, acorde a la ratio de ausentes establecido en la tabla 8, por lo que se han eliminado los datos referidos a dos meses, lo cual implica una ratio de perdidos de en torno a un 15 % por categoría. No obstante, esto no aplica a dos variables de cada subgrupo, ya que hay métodos que requieren la existencia de información parcial para imputar con éxito. La elección de las dos variables cuyos datos completos permanecerán en el conjunto se obtiene a partir de un número generado con una distribución uniforme, de manera que las matrices obtenidas para la comparación de técnicas sean aleatorias. De esta manera, la figura 10 representa la tabla de *Empleo* creada para IBM SPSS en aras de comparar las técnicas de imputación de información ausente:

Fecha	Altas Afiliados	Bajas Afiliados	Contratos Registrados	Contratos Registrados a Tiempo Parcial	Contratos Registrados CVE	Contratos Registrados Indefinidos	Contratos Registrados Interinidades	Contratos Registrados Temporales	Ocupacion crisis 2020	Productividad por hora efectivamente trabajada	Productividad por puesto de trabajo a tiempo completo	Total Afiliados
Jun-21	109591.18	109613.6	1798047	657589	1558935	172866	68264	1625181	98.521553	-6.9856	0.7403	19500277.41
May-21	99487.524	87156.38	1545308	518699	1513358	156148	63912	1389160	97.745435	-4.8755	-0.307833333	19267221
Apr-21		75257.65						1192765				
Mar-21	78293.478	81870.91	1404107	453042	1475289.574	207191	62590	1196916	96.1932	-0.6553	-2.4041	18920901.87
Feb-21	74705	76258.15	1212284	339408	1459781	132431	63754	1079853	96.423	-1.403	-2.873	18850112
Jan-21	93691	97517.68	1302429	358990	1451485	124191	66557	1178238	96.652	-2.151	-3.343	18829480
Dec-20	87074.474	88322.68	1355147	401285	1381692.35	111822	66636	1243325	96.881839	-2.8986	-3.8121	19048433.32
Nov-20	76003.905	84163.52	1449810	419228	1512946.466	128189	71882	1321621	96.373999	-2.816333333	-3.675666667	19022001.57
Oct-20		89648.57						1399038				
Sep-20	143821	120663.1	1632484	542413	1408332	163209	74131	1469275	95.358318	-2.6518	-3.4028	18862712.8
Aug-20	71455.667	76186.86	1118663	353195	1419365	96275	53965	1022388	94.635622	-0.5759	-3.7324	18792376.14
Jul-20	87149.174	79818.87	1536122	531513	1312696	141105	67421	1395017	93.912926	1.5	-4.062	18785554.3
Jun-20	74455.955	80719.91	1159602	362772	1003079	114393	47439	1045209	93.19023	3.5759	-4.3916	18624336.68
May-20	55230.15	45519.85	850617	195131	832856	76692	32902	773925	94.983364	2.438066667	-4.105433333	18556128.85
Apr-20	47802.3	49994.4	673149	134515	717819	59042	38898	614107	96.776498	1.300233333	-3.819266667	18458666.8
Mar-20	80012.318	118714.1	1256510	397997	1327930	145393	64815	1111117	98.569633	0.1624	-3.5331	19006759.59

Figura 10. Instantánea del subgrupo *Empleo* tras la eliminación de casos con información ausente

De esta manera, se imputan los datos empleando IBM SPSS, a través de imputaciones simples de EM y de regresión, así como la imputación múltiple (IM) empleando MCMC. Dado que los análisis de White-Saphiro y Kolgomorov-Smirnov no permiten descartar la normalidad de los datos para aproximadamente la mitad de las variables, la técnica de EM asumirá normalidad y la regresión empleará variantes normales para corregir la estimación y así no subestimar la variabilidad de los datos, lo cual mejora su robustez (Musil *et al.*, 2002). De esta forma, se generan tres archivos por cada técnica empleada, de manera que el ejemplo de esta categoría *Empleo* se presenta en el anexo C. Como ejemplo de los mismos, la figura 11 representa la matriz resultante de la imputación múltiple para la categoría *Empleo*:

Fecha	Altas Afiliados	Bajas Afiliados	Contratos Registrados	Contratos Registrados a Tiempo Parcial	Contratos Registrados CVE	Contratos Registrados Indefinidos	Contratos Registrados Interinidades	Contratos Registrados Temporales	Ocupacion crisis 2020	Productividad por hora efectivamente trabajada	Productividad por puesto de trabajo a tiempo completo	Total Afiliados
Jun-21	109591	109614	1798047	657589	1558935.0	172866.0	68264.0	1625181.0	98.522	-6.986	0.740	19500277
May-21	99488	87156	1545308	518699	1513358.0	156148.0	63912.0	1389160.0	97.745	-4.876	-0.308	19267221
Apr-21	100718	75258	1302429	488740	1210453.8	111822.0	57712.8	1192765.0	96.173	-2.555	-1.896	19196120
Mar-21	78293	81871	1404107	453042	1475289.6	207191.0	62590.0	1196916.0	96.193	-0.655	-2.404	18920902
Feb-21	74705	76258	1212284	339408	1459781.0	132431.0	63754.0	1079853.0	96.423	-1.403	-2.873	18850112
Jan-21	93691	97518	1302429	358990	1451485.2	124191.0	66557.0	1178238.0	96.652	-2.151	-3.343	18829480
Dec-20	87074	88323	1355147	401285	1381692.3	111822.0	66636.0	1243325.0	96.882	-2.899	-3.812	19048433
Nov-20	76004	84164	1449810	419228	1512946.5	128189.0	71882.0	1321621.0	96.374	-2.816	-3.676	19022002
Oct-20	105693	89649	1536122	508128.6	1270615.7	141105.0	56719.0	1399038.0	95.252	-0.851	-2.371	18989150
Sep-20	143821	120663	1632484	542413	1408332.0	163209.0	74131.0	1469275.0	95.358	-2.652	-3.403	18862713
Aug-20	71456	76187	1118663	353195	1419365.0	96275.0	53965.0	1022388.0	94.636	-0.576	-3.732	18792376
Jul-20	87149	79819	1536122	531513	1312696.0	141105.0	67421.0	1395017.0	93.913	1.500	-4.062	18785554
Jun-20	74456	80720	1159602	362772	1003079.0	114393.0	47439.0	1045209.0	93.190	3.576	-4.392	18624337
May-20	55230	45520	850617	195131	832856.0	76692.0	32902.0	773925.0	94.983	2.438	-4.105	18556129
Apr-20	47802	49994	673149	134515	717819.0	59042.0	38898.0	614107.0	96.776	1.300	-3.819	18458667
Mar-20	80012	118714	1256510	397997	1327930.0	145393.0	64815.0	1111117.0	98.570	0.162	-3.533	19006760

Figura 11. Instantánea de la categoría *Empleo* tras la imputación múltiple en IBM SPSS

Así, el código 2 del anexo D se encarga de implementar la fórmula del NRMSE para cada categoría, obtenida como el sumatorio de los NRMSE de sus indicadores componentes. Así, su valor no solo dependerá de la fiabilidad de la técnica, sino también del número de indicadores presentes en cada categoría. De esta manera, se obtiene la tabla de errores proporcionados por cada técnica de imputación según la categoría imputada:

NRMSE			
	Desempleo	Empleo	Habilidades y capacitación
<i>Regresión</i>	1.900	0.475	0.471
<i>EM</i>	2.385	0.179	0.396
<i>IM</i>	2.279	0.139	0.394

Tabla 9. Comparativa de los NRMSE para los métodos de imputación disponibles en IBM SPSS

De esta manera, el análisis muestra cómo la técnica de imputación múltiple es la más robusta y fiable para la mayoría de las categorías. Esto refrenda la literatura, aunque una excepción ocurre para el primer subgrupo. En dicho caso, la regresión constituye el método menos sesgado, probablemente debido a la corrección empleada que mejora sensiblemente su potencia. Asimismo, el método EM, que presenta resultados casi idénticos al IM en la tercera categoría y mejora a este en la primera, lo cual implica un rendimiento general aceptable, aunque no destaque significativamente en ningún subgrupo. Esto puede deberse a la poca cantidad de datos y la presencia de *outliers* debido al COVID-19, que altere la normalidad asumida en la estimación, lo cual se ve refrendado por los análisis de Kolmogorov-Smirnov y Saphiro-Wilk que permiten descartar la hipótesis nula de normalidad para un gran número de indicadores.

Por todo ello, los resultados son coherentes con la bibliografía estudiada, considerando especialmente que no existe una regla de oro en la aplicación de una técnica u otra (OCDE, 2008). Para evitarlo, el empleo del análisis en torno al NRMSE aporta una importante información que complementa la teoría de cada método. Por lo tanto, dado que se ha realizado un análisis empírico sobre la fiabilidad de cada método, este será usado como criterio a emplear para la imputación de datos, puesto que además es coherente con lo esperado. Por ello, la regresión se empleará para imputar *Desempleo*, mientras que la imputación múltiple será la herramienta elegida para las otras dos categorías. Así, se obtienen las tablas completas e imputadas de cada subgrupo, ejemplo de las cuales se presenta en la figura 12:

'Fecha'	'Altas Afiliados'	'Bajas Afiliados'	'Contratos Registrados'	'Contratos Registrados a Tiempo Parcial'	'Contratos Registrados CVE'	'Contratos Registrados Indefinidos'	'Contratos Registrados Interinidades'	'Contratos Registrados Temporales'	'Ocupacion crisis 2020'	'Productividad por hora efectivamente trabajada'	'Productividad por puesto de trabajo a tiempo completo'	'Total Afiliados'
'Aug-2021'	84384	100604	1407563	473494	1712542.3	118985	59640	1288578	97.937	-4.868	-0.413	19473724
'Jul-2021'	105761	95055	1838250	675631	1634109.8	165500	76907	1672750	97.541	-6.130	-0.345	19591728
'Jun-2021'	109591	109614	1798047	657589	1558935.0	172866	68264	1625181	98.522	-6.986	0.740	19500277
'May-2021'	99488	87156	1545308	518699	1513358.0	156148	63912	1389160	97.745	-4.876	-0.308	19267221
'Apr-2021'	85426	75258	1356845	441024	1397362.9	164080	62841	1192765	96.969	-2.765	-1.356	19055298
'Mar-2021'	78293	81871	1404107	453042	1475289.6	207191	62590	1196916	96.193	-0.655	-2.404	18920902
'Feb-2021'	74705	76258	1212284	339408	1459781.0	132431	63754	1079853	96.423	-1.403	-2.873	18850112
'Jan-2021'	93691	97518	1302429	358990	1451485.2	124191	66557	1178238	96.652	-2.151	-3.343	18829480
'Dec-2020'	87074	88323	1355147	401285	1381692.3	111822	66636	1243325	96.882	-2.899	-3.812	19048433
'Nov-2020'	76004	84164	1449810	419228	1512946.5	128189	71882	1321621	96.374	-2.816	-3.676	19022002
'Oct-2020'	96898	89649	1551357	528183	1352808.7	152319	76952	1399038	95.866	-2.734	-3.539	18990364
'Sep-2020'	143821	120663	1632484	542413	1408332.0	163209	74131	1469275	95.358	-2.652	-3.403	18862713
'Aug-2020'	71456	76187	1118663	353195	1419365.0	96275	53965	1022388	94.636	-0.576	-3.732	18792376
'Jul-2020'	87149	79819	1536122	531513	1312696.0	141105	67421	1395017	93.913	1.500	-4.062	18785554
'Jun-2020'	74456	80720	1159602	362772	1003079.0	114393	47439	1045209	93.190	3.576	-4.392	18624337
'May-2020'	55230	45520	850617	195131	832856.0	76692	32902	773925	94.983	2.438	-4.105	18556129
'Apr-2020'	47802	49994	673149	134515	717819.0	59042	38898	614107	96.776	1.300	-3.819	18458667
'Mar-2020'	80012	118714	1256510	397997	1327930.0	145393	64815	1111117	98.570	0.162	-3.533	19006760
'Feb-2020'	102118	91422	1594763	556975	1896847.0	178193	72302	1416570	96.059	-5.056	-1.188	19250229
'Jan-2020'	86435	123249	1764837	566811	1810956.0	178978	76702	1585859	95.898	-0.329	-3.216	19164494

Figura 12. Representación del subgrupo *Empleo* tras la aplicación de la imputación múltiple

Así, con la obtención del resto de matrices completas de datos (anexo C), se da por finalizado el apartado de imputación para dar paso al análisis multivariante.

3.4. Análisis multivariante

Existen multitud de métodos que permiten esclarecer las características estadísticas de los datos, tal y como se menciona en la sección 2.4. Tal y como aseveran Floridi *et al.* (2011), las conclusiones del análisis multivariante debe informar al autor sobre la estructura de los datos, de forma que contribuyan a las decisiones relativas a la ponderación y agregación posteriores. Por ello, esta se trata de una etapa crítica en todo proceso de elaboración de indicadores sintéticos.

La primera de las técnicas estudiadas está constituida por las regresiones auxiliares, método que describe las correlaciones existentes entre las variables. Existen diferentes acercamientos para su realización, tales como el uso de matrices de gráficos de regresión auxiliar, lo cual es inviable para grandes conjuntos de datos; o, también, el estudio de matrices de correlación, donde aparecen los coeficientes de Pearson mentados en la sección 2.4. Este último método, si bien no arroja conclusiones definitivas, sí permite comprender *grosso modo* la multicolinealidad de los datos.

Por ello, el primer análisis multivariante realizado consiste en la elaboración de las matrices de correlación de todas las categorías. Debido a la facilidad que presenta el programa *Gretl* para su elaboración, se ha utilizado de forma excepcional para este propósito. De esta manera, la figura 13 presenta la matriz de correlación de la categoría *Empleo* a modo de ejemplo y empleando colores cuya saturación destaca los coeficientes elevados:

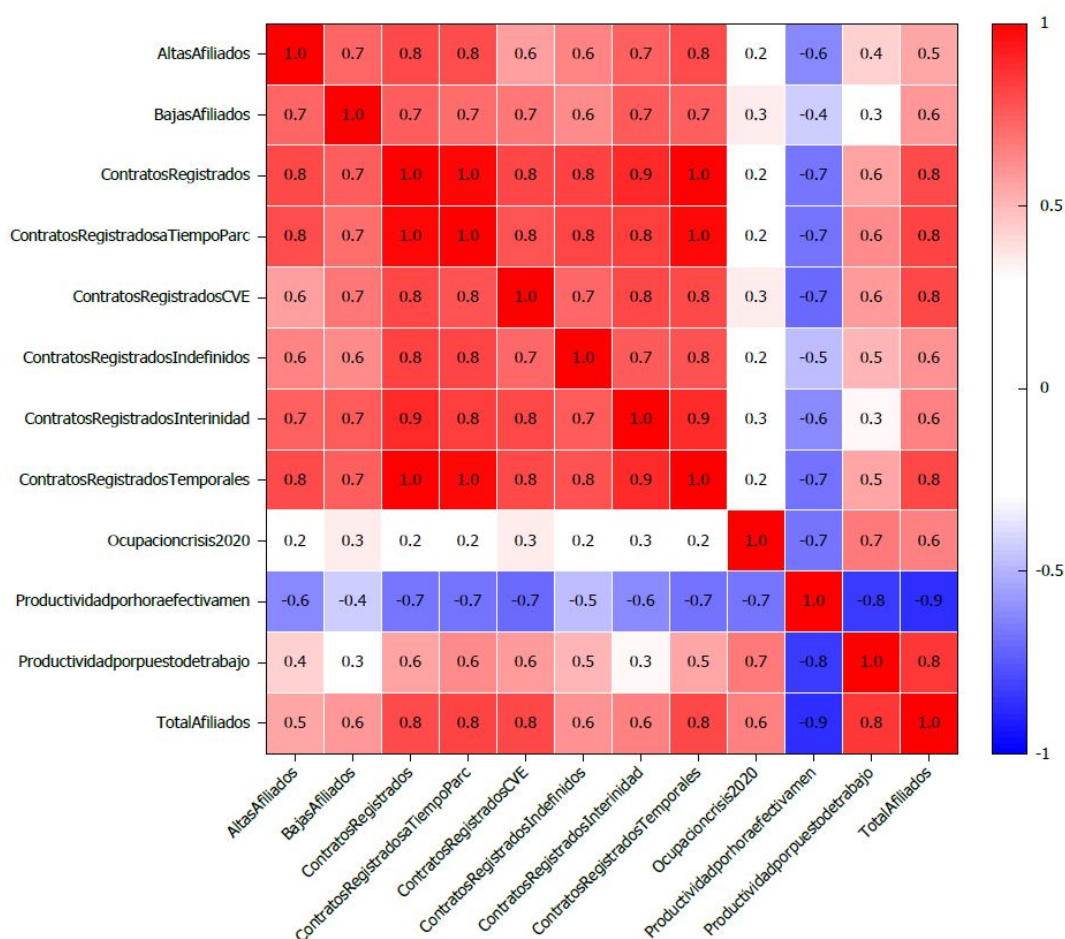


Figura 13. Matriz de correlaciones de los indicadores de la categoría *Empleo*

Las figuras 31, 32, 33 y 34 presentes en el anexo C completan este análisis para el resto de las categorías. En ellas, al igual que en la figura 13, los coeficientes de Pearson tienden a tomar valores muy elevados, dada la similitud entre muchos de los indicadores empleados. Su segmentación por sectores, sumada también a su gran número, provoca una estructura de datos altamente correlada, donde varios indicadores presentarían, supuestamente, el mismo comportamiento a lo largo de la serie temporal. Este fenómeno, si no se aborda correctamente, podría influir de forma muy negativa sobre la ponderación y agregación de los datos, en tanto que se podría estar realizando un doble conteo de un mismo fenómeno representado por varios indicadores componentes. Aunque esto no se aplique a la totalidad de las variables estudiadas, sí parece darse una correlación preocupante entre la mayoría de ellas.

Para corroborar la existencia de esta problemática, se ha de realizar el análisis de componentes principales, ampliamente detallado en la sección de 2.4 de análisis multivariante. Para que su realización sea fiable, se ha de someter la muestra de datos a un análisis KMO de adecuación para PCA (Razali y Wah, 2011). Afortunadamente, el programa IBM SPSS lo realiza simultáneamente, por lo que simplemente habría de aplicarse la técnica y posteriormente ratificar que, como mínimo, el valor de KMO es superior a 0.6 (Reisi *et al.*, 2014).

Asimismo, otro aspecto a analizar del PCA es la cantidad mínima de datos requeridos para garantizar la potencia de la técnica, para lo cual no existe consenso. Típicamente, un mínimo de diez casos suele establecerse, pero existen ratios de casos y variables que aumentan las garantías. En este sentido, los conjuntos de datos analizados cumplen el mínimo establecido por el manual

de la OCDE (2008). Además, también se presenta la problemática de la elección del número de componentes principales, existiendo a tal efecto varios criterios (Peres-Neto *et al.*, 2005): el criterio de Kaiser, el del bastón roto o el del porcentaje de varianza explicada. Sin embargo, en tanto que no es necesario extraer factores para el análisis multivariante, se tomará un criterio básico como el de Kaiser a modo de referencia, es decir, que cada componente extraído explique, como mínimo, lo mismo que un indicador cualquiera. Esto implica que su autovalor ha de ser superior a 1.

Asimismo, también será pertinente una rotación *varimax* como la mencionada en la sección 2.4, de manera que aumente la interpretabilidad de la técnica. A continuación, se presenta la tabla 10 de varianza explicada de la categoría *Empleo*, donde aparecen todos los componentes que explicarían el 100 % de la varianza, así como los tres extraídos por el programa según el criterio de Kaiser:

Varianza total explicada									
<i>Comp.</i>	Autovalores iniciales			Extracción original			Rotación		
	Total	% var.	% acu.	Total	% var.	% acu.	Total	% var.	% acu.
1	8.309	69.238	69.238	8.309	69.238	69.238	6.534	54.449	54.449
2	1.692	14.100	83.338	1.692	14.100	83.338	3.367	28.889	83.338
3	0.616	5.137	88.475						
4	0.470	3.920	92.395						

Tabla 10. Matriz de varianza total explicada por los componentes obtenidos por PCA para *Empleo*

Nótese cómo la rotación solo afecta a los componentes finalmente extraídos, ya que su acumulado de porcentaje de varianza explicada es idéntico al de la extracción original. De esta manera, el análisis PCA reduciría un total de doce indicadores a tres componentes, de manera que los problemas de correlación revelados por el análisis regresivo se solventarían a la vez que se aumentaría notablemente la manejabilidad del conjunto de datos.

No obstante, cabe plantearse entonces la composición de esos tres componentes en función de los indicadores componentes que los forman. De ahí se deducirá si efectivamente hay una reducción de dimensiones relevante y comprensible. A tal efecto, se presenta la tabla 11, en la que se muestran las cargas factoriales de cada indicador respecto de los tres componentes extraídos, una vez se ha realizado la rotación *varimax*:

Matriz de componente rotada		
Indicadores originales	Componentes	
	1	2
<i>Altas Afiliados</i>	0.827	0.177
<i>Bajas Afiliados</i>	0.797	0.171
<i>Contratos Registrados</i>	0.935	0.302
<i>Contratos Registrados a Tiempo Parcial</i>	0.901	0.340
<i>Contratos Registrados CVE</i>	0.751	0.439

<i>Contratos Registrados Indefinidos</i>	0.813	0.224
<i>Contratos Registrados Interinidades</i>	0.903	0.192
<i>Contratos Registrados Temporales</i>	0.926	0.304
<i>Ocupación crisis 2020</i>	0.018	0.881
Productividad por hora efectivamente trabajada	-0.467	-0.816
Productividad por puesto de trabajo a tiempo completo	0.283	0.886
Total Afiliados	0.575	0.775

Tabla 11. Matriz de componentes rotados obtenidos por PCA para la categoría *Empleo*

La matriz de componentes rotados presenta, en el caso de la categoría *Empleo*, una clara agrupación de indicadores relacionados. Se han resaltado aquellas cargas factoriales por encima de 0.5, en pos de aumentar la legibilidad de la tabla. Por ejemplo, los contratos registrados totales y su segmentación formarían, *grosso modo*, el primero de los componentes. El segundo componente vendría dado por la productividad y la cantidad de empleados, mientras que el último reflejaría el indicador “Bajas afiliados”. La agrupación en torno a ciertos componentes no implica que, de forma intrínseca, dos o más indicadores reflejen el mismo fenómeno, sino que describen el mismo comportamiento estadístico. No obstante, en ocasiones, la causa de ese comportamiento similar se encuentra en la existencia de una realidad subyacente representada en dichos indicadores. En este caso, su similitud conceptual apunta a este fenómeno, por lo que la técnica PCA estaría reflejando un serio problema de multicolinealidad entre los indicadores componentes.

Retomando además lo mencionado sobre el test KMO o la esfericidad de Bartlett, el programa ha sido incapaz de realizar el análisis al respecto para algunas categorías: *Desempleo*, *Empleo* y *Habilidades y Capacitación*, excepto para la categoría *Salarios y costes laborales*; que son las que más indicadores presentan. En estos casos, IBM SPSS devolvía un mensaje de error, aludiendo a una matriz de correlaciones no cierta positiva, es decir, de determinante nulo por disponer de autovalores también nulos. De acuerdo con el trabajo referente de Jolliffe (2002), esto una muestra de una extraordinaria correlación presente en los subgrupos, de forma que la técnica PCA es capaz de establecer una dependencia lineal perfecta entre ciertas variables, tantas como autovalores nulos haya en la matriz de correlación en cuestión. De hecho, la matriz de correlaciones presentará un rango inferior a su tamaño, lo cual identificaría la cantidad de dimensiones a reducir para solucionar el problema.

El fenómeno descrito afecta de forma muy negativa a la estabilidad de PCA en la extracción de factores, puesto que puede proporcionar resultados muy poco fiables o robustos (Lorenzo-Seva y Ferrando, 2021). Al fin y al cabo, si se dispone de una matriz de correlaciones “A” mal condicionada (no cierta positiva o de determinante casi nulo), cualquier perturbación o cambio de valor sobre el sistema alterará de forma radical las soluciones del mismo. Por esta razón, los factores extraídos de la técnica PCA no serían robustos.

Esto sería especialmente relevante en caso de decidirse el empleo de la técnica para ponderar y agregar los datos, puesto que habría que tratarlos previamente para obtener una matriz de correlaciones cierta positiva que no influya tan negativamente sobre el proceso. Para ello, Lorenzo-Seva y Ferrando (2021) proponen técnicas de *smoothing* que subsanan el problema

sin alterar demasiado el conjunto existente de datos. Aun así, esto no resulta necesario para la fase de análisis multivariante.

Por el contrario, en las categorías menos “pobladas”, en las cuales parece darse una menor correlación entre indicadores, el análisis KMO y el de esfericidad de Bartlett sí se llevaron a cabo de forma exitosa. La siguiente tabla 12 representa los valores de los test de adecuación para dichas categorías:

Análisis de KMO y Bartlett		
Indicadores	<i>Protección a los desempleados</i>	<i>Salarios y costes laborales</i>
Medida KMO	0.512	0.627
Prueba de esfericidad de Bartlett (<i>p-valor</i>)	0.00	0.00

Tabla 12. Adecuación de *Protección a los desempleados* y *Salarios y costes laborales* para PCA

De esta tabla se deduce que el análisis PCA no debe aplicarse sobre la categoría *Protección a los desempleados*, puesto que no cumple el criterio mínimo de $KMO > 0.6$, establecido ya de forma conservadora en el estado de la cuestión (Reisi *et al.*, 2014). Sin embargo, en el caso del otro subgrupo, la aplicabilidad del análisis es más factible, puesto que sí cumple dicho criterio. Además, presenta un *p-valor* de la prueba Bartlett que implica correlaciones en su conjunto de datos, puesto que este test evalúa la hipótesis nula de igualdad entre la matriz de correlaciones del subgrupo y la identidad (nada correlada, por tanto). Dado que ambos *p-valor*es son nulos, puede rechazarse dicha H_0 . Esto, sumado a la medida KMO, permite la aplicación de la técnica PCA sobre *Salarios y costes laborales* en una eventual ponderación y agregación.

El proceso de análisis PCA se realiza para todas las categorías, con la excepción de la categoría *Protección a los desempleados*, cuyas dimensiones hacen que tenga poco sentido la aplicación de la técnica PCA. Además, esta categoría también incumple el criterio de KMO mayor de 0.6. Al fin y al cabo, solo tres indicadores forman parte de dicho subgrupo, por lo que emplear una técnica de reducción de dimensiones, incluso en un contexto donde dos de ellas están muy correladas (ver imagen 33 en la página 135), apenas aportaría ninguna información nueva. Por si esto fuera poco, la comprensión de los datos no reviste complejidad como puede ocurrir en otras categorías, por lo que aplicar PCA sobre la susodicha es innecesario.

En lo referente al resto de subgrupos, sus matrices de varianza total explicadas y de componentes rotados se encuentran en el anexo B. Así, la tabla 26 muestra cómo los 21 indicadores de la categoría *Desempleo* se convertirían en 5 componentes que describirían el 89.4 % de la varianza original. Asimismo, la tabla 25 revela el mismo fenómeno que la tabla 11 mostrada anteriormente, puesto que indicadores de similar índole se agrupan en torno al mismo componente. En el caso de la categoría *Habilidades y capacitación*, la varianza total explicada por los componentes de la técnica PCA es el 96.3 %, logrando reducir 24 indicadores en tan solo 4 componentes. Esto se muestra en las tablas 27 y 28. En lo referente a *Salarios y costes laborales*, la tabla 30 refleja que dos componentes mostrarían el 93.7 % de la varianza descrita por los ocho indicadores originales, cuyas cargas factoriales se presentan en la tabla 29. En este subgrupo, las cargas factoriales de los indicadores reflejan los mismos fenómenos de agrupamiento que en otras categorías.

Para concluir el análisis multivariante, se realizará un estudio sobre el coeficiente alfa de Cronbach. Este coeficiente proporciona una medida de la consistencia de un grupo de datos, de manera que un valor cercano a la unidad reflejaría la existencia de un fenómeno que explique todo el conjunto (OCDE, 2008). En este caso, en tanto que cada categoría pretende representar una realidad determinada alineada con objetivos de transformación social (*Desempleo, Empleo, Salarios y costes laborales...*), cabe esperar que los coeficientes sean elevados. A tal efecto, la tabla muestra sus valores para las cinco categorías estudiadas:

Coeficiente alfa de Cronbach	
<i>Desempleo</i>	0.846
<i>Empleo</i>	0.853
<i>Habilidades y capacitación</i>	0.827
<i>Protección a los desempleados</i>	0.384
<i>Salarios y costes laborales</i>	0.711

Tabla 13. Coeficientes alfa de Cronbach para las diferentes categorías

Los coeficientes se mantienen por encima del valor de 0.7 establecido como umbral de fiabilidad en la sección 2.4 del estado de la cuestión, exceptuando la cuarta categoría. Este subgrupo presenta, por tanto, mucha menos correlación y consistencia por constar de solo tres indicadores, pese a que dos de ellos sí presentan un coeficiente de Pearson reseñable (ver matriz de correlación de la figura 33 presente en el anexo C). Además, también se han estudiado los coeficientes de Cronbach si se eliminaran variables tal y como se menciona en el manual de la OCDE (2008). Estos análisis se presentan en el anexo B, puesto que no arrojan resultados relevantes para el estudio multivariante. Retomando el comentario sobre los coeficientes de Cronbach generales, de la tabla 13 se concluye que cuatro de los subgrupos son fiables y efectivamente representan algún fenómeno subyacente, mientras que la categoría de *Protección a los desempleados* no cumple estos estándares. Esto podría mejorarse con la inclusión de más variables que caractericen dicho fenómeno de forma pertinente.

A partir de los estudios multivariantes realizados, queda atestiguada una gran correlación existente entre los indicadores de las diferentes categorías. Dada la importante cantidad de indicadores considerados para la elaboración del sintético, este fenómeno resulta coherente y, en cierto modo, previsible. A fin de cuentas, la eliminación de indicadores siempre ha de realizarse con sumo cuidado, puesto que se corre el riesgo de desechar información relevante y pertinente (Horton y Kleinman, 2007). A partir del análisis multivariante, en cambio, el autor dispone de más información para tomar decisiones que salvaguarden la fiabilidad del sintético final. Esta nueva información acerca de la estructura de los datos se revela a partir del presente análisis, por lo que cualquier decisión previa estaría infundada y sería peligrosa. De hecho, Burgass *et al.* (2017) afirman que la elaboración de un sintético debería ser iterativa, en tanto que cada etapa de su creación revela nuevos aspectos importantes que, de uno u otro modo, podrían alterar la visión y las decisiones previas tomadas en el proceso.

Como colofón, se ha realizado un análisis *cluster* de cada una de las categorías, empleando como método de agrupación en grupos el enlace promedio a través de la distancia euclídea al cuadrado y una estandarización convencional. Los dendogramas en cuestión se presentan en el anexo C y han de utilizarse como complemento de los análisis PCA para comprender las interacciones entre las variables. A través del susodicho estudio, se refrenda las conclusiones

obtenidas por las técnicas de Cronbach, de PCA y de regresiones auxiliares en lo referente a la cercanía conceptual entre ciertos indicadores. Por ello, todos los análisis multivariantes apuntan en la misma dirección: el conjunto de datos estudiado refleja efectivamente fenómenos subyacentes con correlaciones muy elevadas, lo cual debe ser cuidadosamente abordado en la etapa de ponderación y agregación.

De esta manera, la estructura de datos ha sido detallada y se procede a realizar la normalización de los indicadores originales.

3.5. Normalización

La normalización es una etapa crucial del proceso de elaboración de indicadores. Tal y como se asevera en la sección 2.5 del estado de la cuestión, este procedimiento permite a los autores dotar de una escala común a todos los indicadores componentes que conforman los diferentes subgrupos. En este sentido, la interpretabilidad aumenta si se usa la misma técnica en todos ellos, ya que si no esto modificaría totalmente la forma en que finalmente se agregan. Además, considerando que el proyecto versa sobre un indicador sintético compuesto por una serie temporal, algunas técnicas no serán aplicables.

No obstante, primero cabe plantearse si la realización de las transformaciones explicadas en dicha sección es pertinente, en tanto que suavice la estructura de los datos y rebaje el impacto de posibles *outliers*. Estas transformaciones previas, tal y como asegura la OCDE (2008), son en cierto modo arbitrarias y deben responder a una clara y patente necesidad. Al fin y al cabo, la representación del fenómeno subyacente podría verse muy alterada, de manera que la transformación suponga mayor perjuicio que, por ejemplo, una posible asimetría. Existen, por tanto, dos transformaciones principales: la logarítmica y la “saturación” de valores típicos, las cuales deben ser evaluadas de acuerdo a los análisis univariantes presentados en el anexo B.

En primer lugar, la transformación logarítmica podría ser una opción para aquellas variables con curtosis superiores a 3.5 y asimetrías mayores que 1 en valor absoluto (Hudrliková, 2013). En dichos análisis, dos de las categorías presentan alguna variable que cumple estos criterios (ver tabla 36). Entonces, ¿qué impide la aplicación del método? Para comprender bien si es pertinente su realización, cabe retomar el ejemplo presentado en la página 44 de la sección 2.5. Si a lo largo de una serie temporal de datos se pretende medir la aparición de efemérides meteorológicas, lo lógico y coherente es que aparezcan valores atípicos en el registro estudiado. En dicho caso, la distribución de los datos presentará *outliers* que generarán asimetrías y coeficientes de curtosis muy significativas que, en otra situación, podrían revelar la existencia de mediciones erróneas o indeseadas, pero no en el caso mencionado. Este ejemplo es análogo respecto de la realidad medida por el indicador sintético que se desea elaborar, puesto que el mercado de trabajo sufrió mucho durante la pandemia. Cabe esperar, por tanto, valores atípicos tanto de pico, previamente a la pandemia, como de valle, en el momento más severo de la misma.

Esta misma explicación se aplica a la posible saturación de los *outliers*. En tanto que se pretende medir su impacto, homogeneizar el conjunto de datos en dichos casos iría en contra del objetivo del sintético. Asimismo, la alta fiabilidad de los indicadores componentes, corroborada tanto por los criterios de 360 Smart Vision como por el análisis empírico de los valores atípicos mencionado en la sección 3.2, hace que la existencia de mediciones erróneas sea nula o prácticamente hipotética. Por ello, no se realizará ni el truncamiento propuesto por

Freudenberg (2003) ni la transformación logarítmica explicada en el manual de la OCDE (2008), puesto que iría en contra del objetivo del trabajo y no existe prueba alguna que revele problemas de fiabilidad en los datos. Sin embargo, sí cabe plantearse esta posibilidad para el análisis de robustez que se realizará más adelante, puesto que la decisión de no transformar los datos es en sí una fuente de incertidumbre para el sintético (2017).

Una vez esto ha sido explicado, se han de evaluar los diferentes métodos de normalización para decidir cuál es más apropiado de acuerdo con el objetivo del sintético. Después de todo, algunos de ellos son idóneos para la elaboración de *rankings*, pero no para series temporales. El objetivo del sintético no es representar en qué fechas la *salud del mercado de trabajo* tiene un mejor desempeño, sino elaborar un vector temporal que permita comparar su evolución manteniendo las distancias entre los valores. Si no se hiciera esto, simplemente se obtendría una serie temporal con un orden de puntuación o, incluso, con apelativos como “bien”, “mal” o “regular” según su rendimiento. Por ello, dado que garantizar la distancia entre los datos es fundamental, la técnica de puntuación se descarta (ver página 44 de la sección 2.5). Asimismo, el método de distancia a una referencia comprime demasiado los datos, tal y como se atestiguó en la figura 4. Por ello, estos métodos no resultan apropiados ni especialmente ventajosos para la elaboración de un indicador sintético que presenta una estructura de serie temporal.

Por estas razones, cabe plantearse por tanto las otras dos técnicas principales de normalización: la estandarización, que coloca todos los valores en torno a una media nula y varianza unidad; y el método Mín-Máx (Saisana *et al.*, 2005), que supone la división de cada término entre el rango posible de valores previa resta del mínimo registrado en la serie (ver ecuación 8). Los *outliers*, aunque tienen un gran impacto sobre ambas técnicas, han de permanecer por el mencionado alineamiento con el objetivo del sintético. Llegados a este punto, entonces, ¿qué criterio permite dilucidar la elección de método?

La estandarización, si bien su utilización está muy presente en las referencias, puede presentar valores nulos y, por supuesto, mostrará valores negativos. Esto puede afectar de forma notoria a la agregación, impidiendo la aplicación de técnicas geométricas. Esto se justifica mediante la ecuación 19, en la que los valores negativos estarían potenciados a cierto peso w_i , previsiblemente entre 0 y 1, lo cual implicarían raíces de un número negativo. Esta incompatibilidad también aparece en el manual de la OCDE (2008). Dado que el empleo de números complejos implicaría complicaciones innecesarias, el método Mín-Máx se elige como técnica de normalización, de acuerdo con la siguiente fórmula:

$$I_{it} = \frac{x_{it} - \min(x_i)}{\max(x_i) - \min(x_i)} \quad (21)$$

Este método proporciona valores cuyo rango varía desde 0 (mínimo de la serie) hasta 1 (máximo de la serie). Esta escala, al contrario que otras técnicas, favorece la etapa posterior de ponderación y agregación, ya que bajo ciertos preceptos el resultado final mantendría la escala mencionada, cuya interpretación es muy sencilla. Asimismo, la técnica no impone ninguna restricción relevante en cuanto a la estructura de datos originales, con la salvedad del impacto que los valores atípicos tienen sobre ella. Para medir esto, se realizará una versión “saturada” de la normalización, de manera que se cuantifique la influencia de los *outliers* sobre el resultado final, al menos en lo referente a la técnica de normalización.

De esta manera, el código 3 presente en el anexo D aplica la técnica mencionada. Una vez realizado, surge la siguiente problemática: no todos los indicadores tienen el mismo sentido, es

decir, algunos describen fenómenos que reflejan el mal funcionamiento del fenómeno descrito, mientras que otros lo representan en positivo. Por ejemplificar, la categoría *Empleo* dispone de una variable *Altas afiliados*, que representa la buena salud del fenómeno; y también *Bajas afiliados*, que refleja lo contrario. De esta manera, en aras de aumentar la interpretabilidad de los datos, resulta beneficioso orientarlos en el mismo sentido. Esto se realiza, en el método Mín-Máx, de acuerdo con la ecuación 11, de manera que simplemente se resten a uno cada uno de los valores. Así, se invierte la escala en esos casos. Esto ya viene reflejado en el código mencionado.

De esta manera, se presenta la tabla 14 con los valores ya normalizados, de manera que sean representativos del conjunto de categorías normalizadas presentadas en el anexo C:

Fecha	Altas Afiliados	Bajas Afiliados	Contratos Registrados	Contratos Registrados a Tiempo Parcial	Contratos Registrados CVE	Contratos Registrados Indefinidos	Contratos Registrados Interinidades	Contratos Registrados Temporales	Ocupacion crisis 2020	Productividad por hora efectivamente trabajada	Productividad por puesto de trabajo a tiempo completo	Total Afiliados
Aug-21	0.381	0.291	0.630	0.626	0.844	0.405	0.607	0.637	0.882	0.201	0.775	0.896
Jul-21	0.604	0.363	1.000	1.000	0.777	0.719	0.999	1.000	0.809	0.081	0.789	1.000
Jun-21	0.644	0.175	0.965	0.967	0.713	0.768	0.803	0.955	0.991	0.000	1.000	0.919
May-21	0.538	0.464	0.749	0.710	0.675	0.655	0.704	0.732	0.847	0.200	0.796	0.714
Apr-21	0.392	0.617	0.587	0.566	0.576	0.709	0.680	0.547	0.703	0.400	0.592	0.527
Mar-21	0.318	0.532	0.627	0.589	0.642	1.000	0.674	0.551	0.558	0.599	0.387	0.408
Feb-21	0.280	0.605	0.463	0.379	0.629	0.495	0.700	0.440	0.601	0.529	0.296	0.345
Jan-21	0.478	0.331	0.540	0.415	0.622	0.440	0.764	0.533	0.644	0.458	0.204	0.327
Dec-20	0.409	0.449	0.585	0.493	0.563	0.356	0.766	0.594	0.686	0.387	0.113	0.521
Nov-20	0.294	0.503	0.667	0.526	0.674	0.467	0.885	0.668	0.592	0.395	0.140	0.497
Oct-20	0.511	0.432	0.754	0.728	0.539	0.630	1.000	0.741	0.497	0.403	0.166	0.469
Sep-20	1.000	0.033	0.823	0.754	0.586	0.703	0.936	0.808	0.403	0.410	0.193	0.357
Aug-20	0.246	0.605	0.382	0.404	0.595	0.251	0.478	0.386	0.269	0.607	0.128	0.295
Jul-20	0.410	0.559	0.741	0.734	0.505	0.554	0.784	0.738	0.134	0.803	0.064	0.288
Jun-20	0.278	0.547	0.418	0.422	0.242	0.374	0.330	0.407	0.000	1.000	0.000	0.146
May-20	0.077	1.000	0.152	0.112	0.098	0.119	0.000	0.151	0.333	0.892	0.056	0.086
Apr-20	0.000	0.942	0.000	0.000	0.000	0.000	0.136	0.000	0.667	0.785	0.112	0.000
Mar-20	0.335	0.058	0.501	0.487	0.517	0.583	0.724	0.469	1.000	0.677	0.167	0.484
Feb-20	0.566	0.409	0.791	0.781	1.000	0.804	0.894	0.758	0.533	0.183	0.624	0.699
Jan-20	0.402	0.000	0.937	0.799	0.927	0.810	0.994	0.918	0.503	0.630	0.229	0.623

Figura 14. Instantánea del subgrupo *Empleo* tras la normalización y reorientación de los datos

Así, la etapa de normalización llega a su fin tras la evaluación de la aplicabilidad de cada técnica y la posterior implementación realizada. La próxima etapa, clave en la elaboración del sintético, versa sobre la ponderación y la agregación.

3.6. Evaluación intermedia de resultados

Llegados a este punto, se dispone de una estructura de datos ordenada, normalizada y preparada para ser ponderada y agregada. Cada categoría goza de un número determinado de indicadores, los cuales han sido seleccionados garantizando que la máxima cantidad de información se mantuviera presente, de manera que no se desechara ninguna de manera arbitraria. Como consecuencia de ello, los indicadores presentan correlaciones muy elevadas entre sí, tal y como se presentó en la sección de 3.4 a través de las matrices de correlación, el coeficiente alfa de Cronbach, el análisis *cluster* y, por supuesto, la aplicación del criterio de adecuación KMO para PCA.

Por lo tanto, este paso no resulta sencillo. Elegir una técnica que no se vea afectada por tamaña correlación resulta arduo y, además, la cantidad presente de indicadores hace difícil la implementación de técnicas participatorias de ponderación, al menos en el nivel de subagregación inicial. Además, ciertas técnicas como DEA o BOD, que integran ponderación y agregación, no son idóneas para la realización de series temporales y además implican una sustitutabilidad muy elevada, la cual no se desea implementar sobre el sintético final. La gran

cantidad de indicadores, además, genera una complejidad importante a la hora de evaluarlos, al igual que pasaría con los métodos de MCA y regresivos. Asimismo, la técnica PCA no debería aplicarse, pues para algunas categorías no gozaría de la suficiente fiabilidad (*Desempleo, Empleo, Habilidades y capacitación y Protección a los desempleados*).

Sin embargo, tal y como se mencionó en la metodología del trabajo (ver sección 1.6), el proceso de elaboración ha de estar en constante realimentación para lograr resultados óptimos, es decir, que no hay motivo alguno para ser inflexibles llegados a este punto. Si la dimensionalidad de los datos es un problema, no hay razón alguna para no tratar de resolverla tomando las medidas oportunas. De hecho, en este momento se dispone de una valiosa información sobre la estructura de datos que, en el momento de la elaboración del marco teórico, no se poseía. De hecho, la problemática presentada ya se recoge en las referencias estudiadas:

La exclusión de indicadores individuales relevantes y la inclusión de otros irrelevantes en la matriz de correlación factorizada afectará, normalmente de forma sustancial, a los factores que se obtendrán. [...], conocer sus estructuras factoriales por adelantado ayuda a elegir correctamente qué indicadores incluir y produce el mejor análisis factorial. Este dilema crea el problema del huevo y la gallina (OCDE, 2008, p. 66).

Por lo tanto, si bien los problemas descritos son desconocidos al comienzo de la elaboración, su reciente hallazgo permite abordarlos desde una perspectiva integral. De hecho, ya en la sección 3.2 se menciona la problemática referida a la redundancia en la estructura de datos (ver página 67), causante de los problemas de correlación. Básicamente, aunque no se pueda demostrar una relación causal previa entre dos variables ($y = ax + b$), las matrices de correlación no ciertas positivas revelan que sí se comportan como tal. A fin de cuentas, correlación no siempre implica causalidad.

Para solventar el problema, existen varias formas de afrontarlo. En primer lugar, un aumento de la muestra de datos es la solución óptima, ya que solo mejora la robustez del análisis y, al mismo tiempo, no implica alterar la estructura de datos preexistentes. No obstante, dado que la cantidad de variables es exuberante tanto para las técnicas participatorias como estadísticas (por la consiguiente correlación), se ha llevado a cabo un análisis que permita reducir las dimensiones de las matrices de correlación hasta que estén debidamente condicionadas, momento en el cual se aplicará la técnica PCA, muy útil para los sintéticos con una cantidad notable de datos y elevadas correlaciones. Después de todo, la reducción en las dimensiones solo eliminará las dependencias lineales directas, pero no la elevada correlación del conjunto. Además, este método proporciona una gran flexibilidad en la agregación posterior, permitiendo la aplicación de una gran variedad de técnicas. Este proceder, si bien puede resultar algo tosco por la posible pérdida de información, no es tal si existe un análisis cuantitativo y cualitativo que refrende la redundancia de información (Lorenzo-Seva y Ferrando, 2021).

De esta manera, la reducción en las dimensiones tiene dos vertientes: análisis empírico, en tanto que el código 4 del anexo D proporciona un valor para la sobredimensionalidad de cada categoría (diferencia entre el tamaño de la matriz de correlación y su rango), el cual coincidiría con el número mínimo de indicadores redundantes (Jolliffe, 2002); y un análisis cualitativo, que sopesa los resultados arrojados por el programa y las relaciones causales de las variables desde un punto de vista teórico. De esta manera, el análisis dual proporcionará resultados fiables.

La tabla 14 presenta el cálculo de la diferencia entre el tamaño y el rango de la matriz de correlación para cada categoría, como una muestra de su posible sobredimensionalidad:

Número de indicadores a eliminar				
<i>Desempleo</i>	<i>Empleo</i>	<i>Habilidades y capacitación</i>	<i>Protección a los desempleados</i>	<i>Salarios y costes laborales</i>
2	1	17	0	0

Tabla 14. Número de indicadores linealmente dependientes en cada categoría

Llama poderosamente la atención el número 17, ya que la categoría en cuestión dispone de 24 indicadores, lo que implica una colinealidad demasiado importante. Asimismo, los subgrupos *Protección a los desempleados* y *Salarios y costes laborales* no presentan los problemas mentados. El primero, por su baja dimensionalidad, no es apto para PCA en ningún caso, ya que tampoco cumple el criterio de KMO ni el del sentido común. En cambio, el segundo, es apto en todos los sentidos. Por ello, las tres primeras categorías centrarán el análisis de ahora en adelante. En el caso de las dos primeras categorías, el código fuente empleado proporciona también qué indicador, de ser eliminado, contribuiría en mayor medida a la estabilidad y robustez de la matriz de correlación en cuestión, lo cual permitiría posteriormente la aplicación de PCA en los subgrupos.

Así, el programa revela que, en el caso de *Desempleo*, los indicadores a eliminar serían “Tasa de paro” y “Tasa de paro mujeres”. En el caso de *Empleo*, en cambio, solo sería “Contratos registrados”. Comenzando por este último subgrupo, la eliminación de dicho indicador ya solucionaría el problema y, de hecho, el valor de KMO sería aproximadamente 0.6 tras su exclusión. Esto se atribuye a la segmentación de los contratos registrados en otros indicadores según la temporalidad, por ejemplo, de forma que se describiría el fenómeno de los contratos totales a partir de ellos. No obstante, un análisis cualitativo de los indicadores, también corroborado a partir del estudio multivariante, revela que variables de dicha categoría tales como “Altas afiliados”, “Bajas afiliados” y “Total afiliados” son, en cierto modo, redundantes. Al fin y al cabo, dos son variables de tipo flujo y la última de existencias. Por lo tanto, con dos de ellas bastaría para caracterizar la situación de los afiliados en España: tanto cuántos hay como su evolución. Por ello, la eliminación de uno de estos indicadores será positiva para la estabilidad y robustez del conjunto, de forma que se conserven la variable *stock* “Total afiliados” y el indicador “Bajas afiliados”. Estos pequeños cambios afectan de forma notoria a la prueba KMO, que ahora arroja un resultado de 0.652. Esto es garante de la robustez de la nueva estructura de datos.

En el caso de *Desempleo*, la solución es más compleja. Los resultados arrojados por el programa son coherentes con el análisis preliminar de las relaciones entre variables, ya que el paro general viene, en cierto modo, determinada por el paro presente en las diferentes industrias. Por supuesto, dado que no todas las actividades económicas están representadas, esta relación no es perfecta, razón por la cual se había mantenido el indicador en un comienzo. Lo mismo se aplica al paro por sexos. Al fin y al cabo, el paro total ya está determinado por los paros sectorizados, por lo que con tan solo uno de los sexos se caracterizaría al otro. La relación no es del todo perfecta, pero el análisis cualitativo sí avala la exclusión de uno de ellos para salvaguardar la robustez de los datos. Sin embargo, dado que el colectivo de mujeres es más vulnerable en el mercado laboral (Alonso y Salmerón, 2003), mantener sus indicadores resulta más pertinente a la hora de generar un sintético que busque la mayor mejoría social, puesto que

así no se invisibiliza esta información a través de otros indicadores.

Pese a los avances logrados con estas exclusiones, la categoría *Desempleo* sigue requiriendo de un ajuste adicional para ser suficientemente robusto y adecuado para PCA. Llegados a este punto, en que cantidad de información y robustez final se contraponen, la flexibilización de los criterios originales de selección de datos debe considerarse. Por ello, la pertinencia de ciertos indicadores tales como la ratio “Solicitudes / Vacantes”, proporcionada de forma segmentada en sectores distintos a los del paro, queda en entredicho, pues se contrapone con la ratio general que también está presente en el conjunto de datos. Además, en este caso la fuente de información es un organismo privado, en contraposición con la totalidad del resto de variables. Por ello, ha de ponerse especial atención en su comportamiento estadístico y teórico. Al fin y al cabo, estas variables representan una muestra de cuánta gente está buscando trabajo a través métodos alternativos al SEPE, desde la iniciativa privada. Por ello, tal vez la presencia de 8 indicadores que, además, presentan elevadas correlaciones, sea excesiva para representar ese fenómeno. De hecho, dejando la más característica, la ratio general, la robustez del conjunto aumenta de forma notoria. Además, se obtendría un $KMO = 0.639$ tras este cambio, lo que hace a la categoría apta para PCA.

En el caso de la tercera categoría, un análisis más detallado habrá de realizarse. Este subgrupo, está formado por otros 8 subconjuntos que representan la capacitación social y el desempleo asociado a la misma: tasa de paro tras realizar FP, tras realizar la ESO, tras Primaria... así con cada fase del sistema educativo, proporcionando también la tasa de paro general. Para cada uno de estos subconjuntos, se dispone de también de la tasa segmentada para hombres y para mujeres. Así, se dan 24 indicadores que, desechando uno de los sexos, darían 16. Asimismo, si se considera que la tasa de paro general viene dada por una relación directa con las otras variables, en tanto que están segregadas según los estudios de los individuos, se obtendrían 14 indicadores divididos en 7 subgrupos. Sin embargo, tal y como se muestra en la tabla 14, solo habría 7 indicadores linealmente independientes en la muestra de 24 originales. Por ello, manteniendo tan solo la tasa de paro total según la capacitación, la matriz de correlación finalmente es definida positiva. La implementación de esta reducción de las categorías se encuentra en el código 5 presente en el anexo D.

Pese a ello, los problemas de robustez aún no concluyen. Para solventarlos, cabe la posibilidad de reorientar la categoría, de manera que siga alineada con el objetivo pero lo refleje de forma ligeramente distinta. Al fin y al cabo, la categoría *Habilidades y capacitación* busca representar el impacto del conocimiento y la formación sobre el mercado laboral. Sin embargo, este impacto no es homogéneo en todo el mercado, es decir, dejando de lado la variable estudiada (competencias adquiridas o nivel de educación), existen otras variables afectando a este fenómeno. Variables que, si bien son exógenas, deben ser consideradas. Una de ellas, y de la que además se disponen datos, es el sexo del individuo. La crisis del coronavirus supone un mayor riesgo de desempleo para las mujeres (Alon *et al.*, 2020), que además tienden a ser remuneradas con menor salario y suelen padecer graves problemas de conciliación que, en ocasiones, perjudica su estatus laboral. Por ello, el peor escenario, el que más mueve a la transformación social y que además refleja cómo la capacitación impacta sobre la salud del mercado laboral, debe implicar el análisis de la tasa de paro de la mujer. Así, conformando la categoría con los indicadores de nivel de educación completado para las mujeres, se dota a la misma de una adecuación aceptable $KMO = 0.576$ para el análisis PCA, de manera que se tenga en cuenta la perspectiva de género para abordar los mayores retos planteados por el

COVID-19.

Los nuevos valores obtenidos para los análisis KMO y de esfericidad de Bartlett se presentan a continuación:

Análisis de KMO y esfericidad de Bartlett			
Indicadores	<i>Desempleo</i>	<i>Empleo</i>	<i>Habilidades y capacitación</i>
Medida KMO	0.652	0.671	0.576
Bartlett (<i>p-valor</i>)	0.00	0.00	0.00

Tabla 15. Nueva adecuación de *Desempleo*, *Empleo* y *Habilidades y capacitación* para PCA

De esta manera, se han solucionado de forma tanto cualitativa como cuantitativa los problemas referidos a la robustez de los indicadores sin alterar ninguno de los objetivos ni pervertir los criterios de selección, ya que tanto los análisis cualitativos como cuantitativos han impulsado toda decisión aplicada. Los resultados posteriores, por tanto, gozarán de menor incertidumbre tras esta evaluación intermedia de resultados. Además, las otras dos categorías, más limpias y robustas, permanecen inalteradas. Sin embargo, sí resulta importante considerar que los cambios realizados en las tres categorías en cuestión son significativos en el porvenir del indicador sintético, por lo que el análisis de robustez (sección 3.8) final comparará los resultados que se hubieran obtenido en caso de haber aplicado PCA sobre el conjunto total de los indicadores. Así, se observará el impacto de esta evaluación de resultados y sus decisiones, así como el acierto o el error derivados de las mismas.

De esta manera, se procede a realizar la fase de ponderación y agregación.

3.7. Ponderación y agregación

Tal y como se menciona en la anterior sección, la ponderación y la agregación, dado que se maneja una importante cantidad de variables (42 tras la reformulación anterior, divididas en las cinco categorías) y que la sustitutabilidad de los indicadores es importante, la mayoría de las técnicas de ponderación y agregación han de descartarse. Los motivos ya han sido explicados y, una vez se ha obtenido una estructura adecuada y robusta para la realización de la técnica PCA, aplicarla resulta sencillo y apropiado para dotar al sintético final de interpretabilidad y potencia estadística.

De esta manera, la ponderación y agregación iniciales serán determinadas por esta técnica. Al igual que en otras secciones, se explicará el proceso de elaboración de los factores para el área de *Empleo*, de forma que el resto de tablas o figuras se encuentran en sus respectivos anexos. El primer paso, al igual que en la fase de análisis multivariante, consiste en rotar la matriz de componentes para darle mayor legibilidad e interpretar los fenómenos que refleja cada componente. Así, en el caso de la mencionada categoría, se obtiene la matriz rotada de componentes directamente de IBM SPSS, representada en la tabla 16:

Matriz de componente rotada	
Indicadores	Componente

	1	2
<i>Bajas Afiliados</i>	-0.802	-0.162
<i>Contratos Registrados a Tiempo Parcial</i>	0.892	0.337
<i>Contratos Registrados CVE</i>	0.793	0.411
<i>Contratos Registrados Indefinidos</i>	0.830	0.207
<i>Contratos Registrados Interinidades</i>	0.915	0.177
<i>Contratos Registrados Temporales</i>	0.919	0.300
<i>Ocupacion crisis 2020</i>	0.039	0.874
<i>Productividad por hora efectivamente trabajada</i>	-0.454	-0.824
<i>Productividad por puesto de trabajo a tiempo completo</i>	0.281	0.889
<i>Total Afiliados</i>	0.593	0.763

Tabla 16. Matriz de componentes rotada con rotación *varimax* de la categoría *Empleo*

Posteriormente, para mejorar aún más su interpretabilidad, se aplica la ecuación 13 para obtener las cargas factoriales finales, resultantes de la ratio entre la varianza explicada por un indicador i en el factor j y la varianza total explicada por dicho factor. De esta manera, se obtienen sus valores, los cuales dan lugar a la siguiente matriz representada por la tabla 17:

<i>Matriz</i>	Rotada		Cuadrática	
Indicadores	<i>Componentes</i>		<i>Componentes</i>	
	1	2	1	2
<i>1</i>	-0.802	-0.162	0.127	0.008
<i>2</i>	0.892	0.337	0.157	0.035
<i>3</i>	0.793	0.411	0.124	0.051
<i>4</i>	0.830	0.207	0.136	0.013
<i>5</i>	0.915	0.177	0.165	0.010
<i>6</i>	0.919	0.300	0.166	0.027
<i>7</i>	0.039	0.874	0.000	0.232
<i>8</i>	-0.454	-0.824	0.041	0.206
<i>9</i>	0.281	0.889	0.016	0.240
<i>10</i>	0.593	0.763	0.069	0.177

Tabla 17. Comparación de las matrices rotada y rotada cuadrática normalizada

Así, una vez facilitada la comprensión de la dimensionalidad de los datos, solo falta determinar qué pesos se aplicarán a los indicadores para posteriormente agregarlos en torno al indicador que reflejará la situación de cada categoría. Para lograr esto, existen dos posibles caminos. En primer lugar, cabe la posibilidad de hallar los indicadores intermedios mencionados en la página 49 de la sección 2.6 (OCDE, 2008). Para calcularlo, bastaría con emplear las σ_{ij} mencionadas en la ecuación 13 y que aparecen ya en la matriz rotada cuadrática y normalizada de la tabla 17.

En este punto, el manual de la OCDE (2008) recomienda usar solo aquellos σ_{ij} que sean medianamente relevantes, es decir, mayores de 0.1 por ejemplo. Para facilitar su legibilidad y comprensión, se han resaltado sobre la tabla 17 los valores que cumplan dicho criterio. Sin embargo, esto redundaría en una varianza explicada aún menor, por lo que resulta más apropiado considerar todos los indicadores para formar los intermedios, ya que así la varianza explicada se conserva y prácticamente no se altera la interpretación de los componentes. Después de todo, aquellos no resaltados apenas afectarán, tanto para bien como para mal. Cabe resaltar, además, que estos indicadores intermedios no son los componentes ni los factores que se obtendrían a partir de la técnica. Más bien, se consideran un estadio intermedio entre el subindicador de categoría y los indicadores que la componen.

Así, aplicando la fórmula mencionada para cada valor de los indicadores a lo largo del tiempo, se obtienen los componentes con la misma estructura de serie temporal. En el caso de *Empleo*, la tabla 18 refleja los valores que toman los indicadores intermedios que formarían, posteriormente, el subindicador de la categoría:

<i>Fecha</i>	Indicador intermedio 1	Indicador intermedio 2
'Aug-2021'	0.583	0.687
'Jul-2021'	0.812	0.695
'Jun-2021'	0.737	0.749
'May-2021'	0.651	0.654
'Apr-2021'	0.600	0.566
'Mar-2021'	0.637	0.511
'Feb-2021'	0.520	0.456
'Jan-2021'	0.503	0.428
'Dec-2020'	0.531	0.436
'Nov-2020'	0.603	0.430
'Oct-2020'	0.661	0.416
'Sep-2020'	0.623	0.384
'Aug-2020'	0.438	0.338
'Jul-2020'	0.629	0.354
'Jun-2020'	0.389	0.283
'May-2020'	0.241	0.313
'Apr-2020'	0.176	0.352
'Mar-2020'	0.488	0.569
'Feb-2020'	0.746	0.557
'Jan-2020'	0.739	0.533

Tabla 18. Indicadores intermedios obtenidos para la categoría *Empleo*

Llegados a este punto, todas las categorías han sido elaboradas en torno a los intermedios gracias al código 6 presente en el anexo D, salvo el subgrupo *Protección a los desempleados*, formado por solo tres indicadores. En este caso, la aplicación de PCA resultaba carente de

sentido, lógica y fiabilidad; por lo que ha de elegirse otro criterio de ponderación para los tres indicadores que lo forman: “Beneficiarios asistenciales”, “Prestaciones contributivas al desempleo parcial” y “Prestaciones contributivas al desempleo total”, presentando estas dos últimas una correlación elevada (0.9 de acuerdo con la matriz de correlación ilustrada por la figura 33 en la página 135 del anexo C). La cercanía entre estas dos variables también queda patente en el análisis *cluster* realizado de ellas (ver 38). Por lo tanto, ha de plantearse la posibilidad de incurrir en un doble conteo (Greco *et al.*, 2019) en caso de realizar una equiponderación.

Para dilucidar este fenómeno, por desgracia, no existe herramienta analítica. Estadísticamente hablando, la correlación es indistinguible de la causalidad. Dependiendo del campo de aplicación, estos fenómenos tienden a disociarse o, al contrario, a aparecer siempre unidos. Sin embargo, en un caso concreto, ha de ser el autor el que dilucide si la causalidad entre dos indicadores es, de hecho, la causa de la correlación existente entre ellos. ¿Representan diferentes fenómenos o solo uno desde enfoques diferentes? En estos momentos de tribulación e incertidumbre, el criterio experto puede ser útil para salir de dudas o para que, al menos, ayude al autor a tomar la decisión.

En este sentido, los “Beneficiarios asistenciales”, las “Prestaciones contributivas al desempleo parcial” y las “Prestaciones contributivas al desempleo total”, si bien se adscriben a un mismo marco u objetivo (cuán bien el Estado protege a sus desempleados), representan realidades distintas, sobre todo en lo referente al tipo de paro que los produce. Por ello, aunque dos de los indicadores presenten elevada correlación, esta no se ha de considerar suficientemente relevante como para alterar una posible equiponderación. Esta técnica, lejos de ser elegida por ser una opción predeterminada, es escogida de forma totalmente deliberada y, de hecho, en consonancia con el criterio experto de los doctores economistas que colaboran en el panel 360 Smart Vision. Por lo tanto, podría incluso argüirse que se trata de una aplicación de la técnica BAL (Saisana y Tarantola, 2002).

Retomando la técnica PCA, antes se había mencionado la existencia de dos métodos para obtener los subindicadores de categoría. El segundo, aún no explicado, versa sobre la obtención directa del subindicador sin pasar por los intermedios. Para ello, combinando las ecuaciones 13 y 23, se obtendría el peso del indicador i sobre el total a partir de la ecuación 22:

$$\alpha_i = \sigma_{ij} \cdot w_j = \sum_{\forall j} \frac{\lambda_{ij}^2}{\sum_{\forall i} (\lambda_{ij})^2} \cdot \frac{\phi_j}{\sum_{\forall j} \phi_j} = \sum_{\forall j} \frac{\lambda_{ij}^2}{\sum_{\forall j} \phi_j} \quad (22)$$

Empleando α_i , se puede lograr la agregación directa en el subindicador de categoría mediante una simple agregación aritmética. No obstante, sigue siendo necesario considerar los pesos de los intermedios para corroborar que se toman de forma acertada, no simplemente por la variabilidad de los datos. Al fin y al cabo, dichos pesos han de valorarse conjuntamente con la tabla rotada cuadrática y normalizada, en un ejercicio de análisis crítico donde se considere la naturaleza conceptual de los indicadores que forman cada intermedio y el peso que recibe dicho indicador. Por supuesto, otra opción es considerar los pesos individuales (α_i) y compararlos entre sí, aunque la cantidad de indicadores puede complicar el análisis. Los pesos de los intermedios se obtienen a través de la fórmula 23 presentada a continuación:

$$w_j = \frac{\phi_j}{\sum_{\forall j} \phi_j} = \frac{\sum_{\forall i} (\lambda_{ij})^2}{\sum_{\forall j} \sum_{\forall i} (\lambda_{ij})^2} = \frac{\text{Varianza explicada por el indicador intermedio } j}{\text{Varianza total explicada por todos}} \quad (23)$$

Así, los pesos de los indicadores intermedios de *Empleo* se presentan en la tabla 19:

Pesos de los indicadores intermedios w_j	
<i>Intermedio 1</i>	<i>Intermedio 2</i>
0.6069	0.3931

Tabla 19. Pesos de los indicadores intermedios de *Empleo* obtenidos de la varianza explicada

La categoría *Empleo* presenta dos intermedios con pesos de 60.1 % y 39.9 %. Considerando la tabla 17, el primero responde a las variables referentes a los contratos registrados además del número de bajas de afiliados, mientras que el segundo alude más bien a la ocupación o número de empleados y su productividad. De esta manera, aunque los pesos proporcionados por la técnica PCA pueden responder solo a la variabilidad estadística, es labor del autor evaluar si la importancia relativa de los indicadores intermedios corrobora dicho comportamiento estadístico. En este caso, el primer intermedio alude al flujo en el empleo, es decir, cuánta gente nueva registra un contrato o se va de la seguridad social. En cambio, el otro intermedio atañe a variables de tipo *stock*, pues reflejan cuántos empleados hay y características como su productividad. De acuerdo con el criterio experto de doctores economistas de la Universidad Pontificia de Comillas, los pesos son coherentes con lo esperado y podrían haberse dado así si la técnica de ponderación fuera criterio experto (BAL). Al fin y al cabo, el primer intermedio refleja mayor información, por estar segregado según el tipo de contrato firmado. Es por ello lógico sobreponderar en cierta medida el primero frente al segundo.

De esta manera, el código 6 presente en la página 168 del anexo D obtiene todos los pesos y resultados de la técnica PCA (ver anexo B), cuyas matrices rotadas son en primer lugar proporcionadas por IBM SPSS. El código elabora las categorías finales empleando tanto los indicadores intermedios (según su w_j) como los indicadores individuales (según su α_i). Así, a estos últimos se les asignan los pesos mostrados en la tabla 20, los cuales son obtenidos a partir de la fórmula 22 y la matriz de rotación de la tabla 16:

Indicadores	α_i
<i>Bajas Afiliados</i>	0.0801
<i>Contratos Registrados a Tiempo Parcial</i>	0.1087
<i>Contratos Registrados CVE</i>	0.0953
<i>Contratos Registrados Indefinidos</i>	0.0876
<i>Contratos Registrados Interinidades</i>	0.1038
<i>Contratos Registrados Temporales</i>	0.1116
<i>Ocupación crisis 2020</i>	0.0915
<i>Productividad por hora efectivamente trabajada</i>	0.1058
<i>Productividad por puesto de trabajo a tiempo completo</i>	0.1039
<i>Total Afiliados</i>	0.1116

Tabla 20. Pesos de los diferentes indicadores para la subagregación en la categoría *Empleo*

En el caso de *Empleo*, los pesos de cada uno de los indicadores son ciertamente similares, lo cual ratifica lo descrito por la tabla 19. Ambos mecanismos de creación de los subindicadores

de categoría son válidos, aunque es importante recordar que los indicadores intermedios no son equivalentes a los factores o los componentes de la técnica PCA, ya que conceptualmente no son lo mismo. Por tanto, su utilidad es tal que se solo emplean para interpretar la idoneidad de la ponderación. Al fin y al cabo, la obtención de los subindicadores de las categorías se puede llevar a cabo directamente a partir de los pesos de la tabla 20. Una vez obtenidos para todas las categorías (acúdase al anexo B para el resto de categorías), la agregación realizada será aritmética, ya que el propio concepto de PCA implica que no existe correlación entre los componentes y, por tanto, se evita el doble conteo. Aun así, esto implica cierta sustituibilidad entre los indicadores originales, la cual es inevitable a menos que se aplique una técnica no compensatoria, la cual entonces debe recalcular los pesos empleando otro mecanismo. Esta sustituibilidad se corregirá en la agregación final, donde tomará forma geométrica, de manera que se palie este efecto sobre el sintético final sin aumentar en demasía la complejidad de la elaboración. Asimismo, el hecho de que los pesos en PCA sumen la unidad facilita su interpretación, así como contribuye a mejorar la transparencia del proceso de elaboración.

No obstante, primeramente se deben agregar los indicadores ponderados de cada una de las categorías, para lo cual se lleva a cabo la agregación aritmética mencionada, cuya forma desarrollada se presentó a través de la ecuación 17, ya sea empleando w_j y los factores intermedios o α_i con cada indicador individual. En el caso de la categoría *Protección a los desempleados*, cada indicador será equiponderado. Así, se obtiene la tabla 21 de categorías:

'Fecha'	Desempleo	Empleo	Habilidades y Capacitación	Protección a los desempleados	Salarios y costes laborales
'Jun-2021'	0.655	0.757	0.517	0.089	0.633
'May-2021'	0.558	0.653	0.484	0.198	0.769
'Apr-2021'	0.467	0.572	0.452	0.230	0.809
'Mar-2021'	0.425	0.571	0.419	0.248	0.412
'Feb-2021'	0.344	0.476	0.355	0.308	0.318
'Jan-2021'	0.386	0.480	0.291	0.374	0.333
'Dec-2020'	0.477	0.497	0.227	0.279	0.427
'Nov-2020'	0.472	0.530	0.223	0.184	0.463
'Oct-2020'	0.451	0.579	0.218	0.157	0.450
'Sep-2020'	0.473	0.591	0.214	0.220	0.461
'Aug-2020'	0.452	0.386	0.392	0.359	0.429
'Jul-2020'	0.502	0.534	0.571	0.500	0.369
'Jun-2020'	0.407	0.349	0.749	0.747	0.370
'May-2020'	0.423	0.245	0.781	0.952	0.147
'Apr-2020'	0.296	0.205	0.812	0.786	0.220
'Mar-2020'	0.714	0.505	0.843	0.137	0.543
'Feb-2020'	0.752	0.674	0.397	0.066	0.650
'Jan-2020'	0.797	0.666	0.454	0.065	0.639

Tabla 21. Subindicadores de las diferentes categorías que forman la *Salud del mercado de trabajo*

De esta manera, la primera subagregación ha sido completada. Tras obtener el subindicador de cada categoría, se puede monitorizar individualmente cada uno de los objetivos del sintético. Además, a partir de esta tabla se puede, por tanto, obtener un gráfico que permitan comprender la evolución de los objetivos del indicador, categoría a categoría:



Figura 15. Evolución de los subgrupos que conforman la *Salud del mercado de trabajo*

Esta gráfica presenta, desde luego, innumerables lecturas. Cada categoría u objetivo se comporta de una manera independiente, aunque existe una clara tendencia a la baja durante los meses de pandemia, a excepción de *Protección a los desempleados*, como era de esperar. De la misma manera, la figura 16 representa un gráfico apilado que permite comprender qué ocurriría en caso de agregar las diferentes categorías de forma aritmética. Al fin y al cabo, las áreas dibujadas se presentan superpuestas, como si su importancia fuera idéntica. De esta manera, la figura mencionada se presenta a continuación:

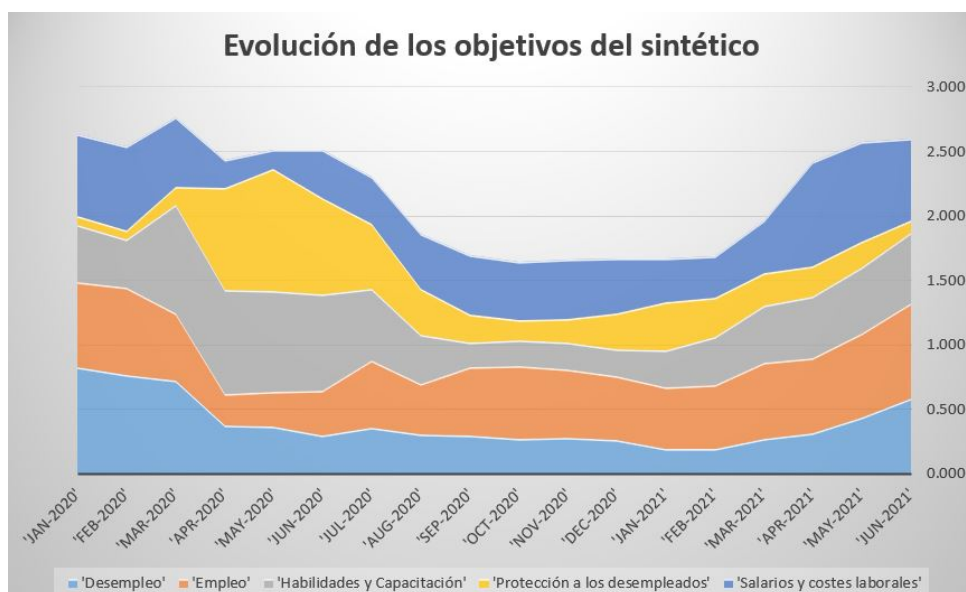


Figura 16. Gráfico apilado de la evolución de los subgrupos del indicador sintético

Si bien las gráficas descritas son muy visuales y didácticas, el análisis de la tendencia de cada una de ellas puede resultar complejo. Para solventarlo, se presentan también las figuras individualizadas en la página 139 presente en el anexo C. A modo de ejemplo, se presenta la figura relativa a la categoría *Habilidades y capacitación*:

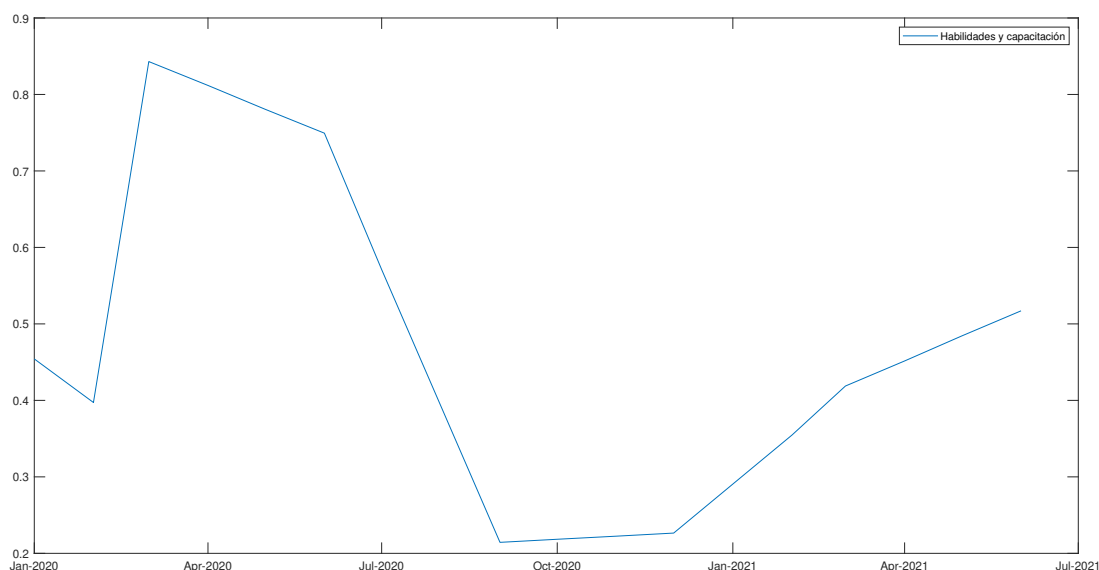


Figura 17. Evolución individualizada de la categoría *Habilidades y capacitación*

Esta categoría debe ser analizada de forma especial, en tanto que caracteriza el estado de las mujeres en el mercado laboral acorde a su nivel educativo. De hecho, dado que este colectivo se ha visto más damnificado por la pandemia del coronavirus (Marina Gómez García *et al.*, 2021), la visibilización de su situación se torna importante. En este caso, la figura 17 presenta una tendencia algo extraña previa a la pandemia, ya que la categoría mejora excesivamente su desempeño hasta marzo para ya posteriormente desplomarse con la irrupción del coronavirus. Esto se atribuye a dos factores: en primer lugar, porque en esos meses se centraban la mayoría de datos ausentes de la categoría, por lo que su imputación aporta cierto sesgo; y en segundo lugar, por el impulso de ciertas iniciativas como los ERTes, que contribuyen a mitigar el impacto de la pandemia sobre unas tasas de paro (según el nivel educativo) que no consideran a las personas en ERTE como desempleados. No obstante, una vez la pandemia se afianza en la sociedad, se ve claramente cómo la tasa de paro por nivel de educativo aumenta, empeorando por tanto el desempeño del indicador. El comportamiento de la tasa de paro masculina, si bien no incluida, es similar, aunque con valores ligeramente superiores.

De esta manera, una vez obtenidos los diferentes subgrupos, se ha de decidir qué técnica de ponderación y agregación se emplea para la elaboración del indicador sintético final. Para ello, existen varias técnicas principales: equiponderación, BOD, BAL, AHP o incluso emplear la opinión pública. Este paso, por ser también crítico en la elaboración, ha de ser debidamente estudiado en la sección 3.8, donde habrá de compararse diferentes métodos de ponderación. Asimismo, también debe considerarse el método de agregación, ya que esto puede influir en la forma en que se consideran los pesos proporcionados por la ponderación.

Así, dado que ya se ha empleado el método aritmético en la subagregación, parece necesario reducir la compensatoriedad del conjunto, lo cual implica la implementación de un MCA o, en su defecto, una agregación geométrica que reduce de forma notable el carácter compensatorio de la aritmética. Esta fórmula, desde un enfoque sencillo, comprensible y transparente, proporciona

un indicador sintético donde un aumento de una unidad en un subindicador no se compensa por la disminución de otro, al menos en las mismas cantidades. Por ello, la agregación geométrica se llevará a cabo según criterio experto, es decir, la técnica BAL (ver 52 para más información). De esta manera, retomando la figura 6 de la sección 2.1, las categorías *Desempleo* y *Empleo* representan las consecuencias observables de una serie de fenómenos estructurales y coyunturales que afectan al mercado laboral, algunos intrínsecos al propio mercado y otros provenientes de otros sectores sociales (finanzas, economía...). Dado que dichas categorías miden este impacto, juntas deberían representar el 50 % del peso final y, como gozan de una importancia similar, serán equiponderadas en torno a dicho valor: $W_1 = 0.25$ y $W_2 = 0.25$.

Por otro lado, el resto de categorías representa parte de las causas de la situación del mercado laboral, como pueden ser los costes del trabajo (salarios) o la productividad (debida a una mayor formación o nivel educativo). Asimismo, la *Protección a los desempleados* supone la garantía que el Estado proporciona a sus ciudadanos frente a los desajustes o ineficiencias del mismo, salvaguardando un bienestar mínimo e impidiendo que caigan en situaciones de exclusión social que reduzcan su potencial reintegración en el mercado. Por ello, dado que esta última categoría resulta más difusa en cuanto a ser una causa directa, recibirá un $W_4 = 0.1$. En cambio, *Habilidades y capacitación* y *Salarios y costes laborales*, dada la importancia que tienen como causas intrínsecas de la salud del mercado laboral, serán equiponderadas sobre el 40 % del peso final, es decir, $W_3 = 0.2$ y $W_5 = 0.2$. De esta forma, aplicando la técnica de agregación geométrica explicada en la ecuación 19, se obtiene el indicador sintético final:

'Fecha'	'Índice Mercado de Trabajo'
'Jun-2021'	0.509
'May-2021'	0.519
'Apr-2021'	0.462
'Mar-2021'	0.387
'Feb-2021'	0.321
'Jan-2021'	0.313
'Dec-2020'	0.330
'Nov-2020'	0.332
'Oct-2020'	0.327
'Sep-2020'	0.340
'Aug-2020'	0.372
'Jul-2020'	0.449
'Jun-2020'	0.427
'May-2020'	0.362
'Apr-2020'	0.380
'Mar-2020'	0.549
'Feb-2020'	0.493
'Jan-2020'	0.510

Tabla 22. Indicador sintético final: *La salud del mercado laboral español*

Estos datos son el resultado de un largo proceso de elaboración en que toda decisión ha influido sobre la serie temporal obtenido. Tal y como era de esperar, la salud del mercado laboral desciende muy fuertemente a partir de marzo de 2020, momento en el cual irrumpe con fuerza la enfermedad en España (Instituto de Salud Carlos III, 2022). En ese momento, la actividad económica y laboral se ralentiza, lo cual también fue promovido por el gobierno español y otros agentes sociales. La imposibilidad de ir a trabajar, ya fuera por enfermedad o por cese de actividad, fue usual. No obstante, algunas medidas, tales como los ERTES o el IMV, contribuyeron a mantener el mercado laboral saneado hasta cierto punto, en detrimento del déficit público. Esto, sumado al levantamiento del estado de alarma en junio de 2020, parece haber implicado una mejoría estival en el mercado laboral español. Sin embargo, tras los meses de junio y julio, los contagios fueron retornando lentamente al panorama pandémico español.

Por estas razones, en la tabla 22 se observa cómo la salud del mercado laboral tarda varios meses en volver a remontar. Si bien esta evolución debería ser analizada con detalle por economistas y expertos en la materia, la aparición de un segundo estado de alarma, en conjunción con un mercado ya exhausto y pesimista frente a la incertidumbre generada por las nuevas olas de coronavirus, supuso el peor momento para el mercado laboral en su conjunto. No obstante, la aparición de las vacunas y el posterior descenso de los casos, que además ya no revestían la misma severidad, permitió que a partir de marzo y abril de 2021 el mercado laboral retomase cierto dinamismo. Aunque ciertamente muchos otros factores influyen sobre el sintético, los hechos acaecidos en esos momentos sí contribuyen a dotar de argumentos a los valores obtenidos. Todos estos puntos se presentan a continuación en la figura 18:



Figura 18. Evolución de la salud del mercado laboral español tras la agregación final

De esta manera, tras la obtención del resultado final, la etapa de ponderación y agregación se da por concluida. El próximo y último paso debe ser, de acuerdo con lo explicado, el análisis de robustez.

3.8. Análisis de robustez

El análisis de robustez, tal y como se explica en la sección 2.7 de la página 56, aspira a cuantificar el impacto de las decisiones tomadas a lo largo del proyecto sobre el resultado final del indicador sintético. Por ello, el primer paso para hacerlo ha de ser identificar dichas decisiones según la etapa en que tuvieron lugar. Por supuesto, tal y como describen Burgass *et al.*, las elecciones del autor forman parte del proceso desde el propio marco teórico, de manera que cada etapa se considera una causa de incertidumbre. Por ello, resulta indispensable acotar qué etapas son susceptibles de análisis y cuántas técnicas alternativas podrían aplicarse para medir el mencionado impacto. Irónicamente, este marco de análisis también conlleva su propio sesgo y, por tanto, incertidumbre. Al fin y al cabo, si se intentan evaluar todas las técnicas mencionadas en el estado de la cuestión para cada etapa, el análisis de robustez sería inabarcable. Por ello, se han de elegir una serie de alternativas para cada una de las fases de carácter crítico para el proceso de elaboración. A tal efecto, la tabla 23 refleja las etapas críticas como variables *input* del proceso, de manera que se presente además la técnica elegida en cada caso y una alternativa que permita la comparación:

X_1	<i>Normalización</i>	Mín-Máx
		Mín-Máx saturado
X_2	<i>Evaluación de resultados</i>	Categorías reformuladas
		Categorías originales
X_3	<i>Ponderación final</i>	Equiponderación
		Criterio experto
X_4	<i>Agregación final</i>	Aritmética
		Geométrica

Tabla 23. Técnicas aplicadas en cada etapa del proyecto y sus alternativas

De esta manera, la tabla describe un árbol de decisiones cuyos impactos, de medirse de forma efectiva, proporcionarían 16 resultados distintos. Cada decisión, por tanto, conforma un factor del modelo, por ejemplo X_1 se refiere a la normalización, de manera que $x_1 = 1$ implicaría normalización Mín-Máx tal y como aparece en el proyecto. Al contrario, $x_1 = 0$ se asociaría a la alternativa considerada, en este caso normalización Mín-Máx previa saturación de los *outliers*. Así, los factores toman la unidad cuando se refieren a la técnica empleada en la elaboración del indicador sintético y valor nulo si responden a su alternativa. De esta manera, considerando que $X = [X_1 X_2 X_3 X_4]$, cada posible resultado vendrá dado por la combinación de los valores que tome cada factor, donde por ejemplo $X = [1 \ 1 \ 1 \ 1]$ reflejaría el indicador sintético obtenido en la sección 3.7. En ocasiones, los análisis de robustez abordan las etapas del proceso de elaboración de forma independiente, es decir, cuantifican cómo afecta una técnica de normalización sin considerar cómo influiría cambiar también, por ejemplo, los métodos de ponderación o agregación. Evaluar cada etapa por separado es, cuando menos, rudimentario. En el otro extremo, en cambio, se encuentran las técnicas de análisis de sensibilidad basadas en la varianza (OCDE, 2008), las cuales requieren de herramientas avanzadas de estudio estadístico. Por ello, su ejecución reviste importante complejidad y su objetivo, si bien se alcanza de forma inmejorable, también se puede lograr a través de otros mecanismos.

Por estas razones, ha de realizarse un análisis de incertidumbre que muestre los diferentes resultados que se obtendrían de alterar cada decisión, así como un análisis fundamental de sensibilidad. Estos dos estudios describirán, por tanto, la robustez del indicador sintético obtenido. Asimismo, para comprender fácilmente qué decisiones se evaluarán, la figura 19 presenta un árbol de técnicas y alternativas que podrían haberse tomado a lo largo del proyecto:

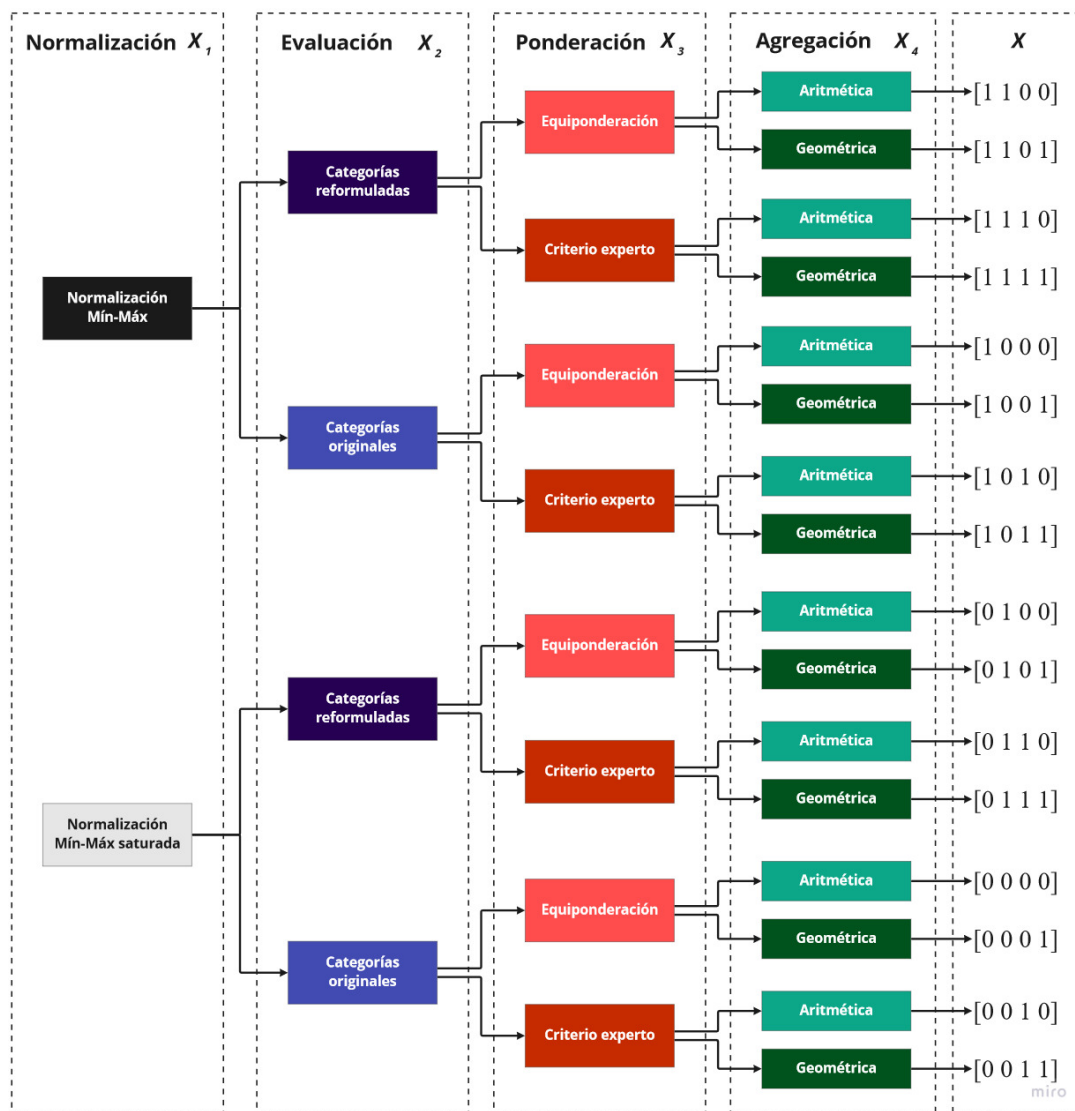


Figura 19. Árbol de decisiones potenciales para cada etapa del proceso

A través de la figura 19 se percibe claramente cómo cada decisión puede influir sobre el impacto que tiene una posterior. Al fin y al cabo, una modificación en el método de agregación final influye de forma diferente si la normalización o la ponderación previas cambian también. Por ello, a continuación se presentan los resultados comparados para cada itinerario de decisiones posibles. La figura 20 representa la evolución de las categorías u objetivos del sintético en el caso de alterar la normalización Mín-Máx por una Mín-Máx con saturación previa, es decir, en la cual los valores atípicos han sido reemplazados por los valores máximos o mínimos dentro del rango presentado en la fórmula 6. Asimismo, también aparece la modificación del proceso de evaluación de resultados, de manera que se muestren tanto las categorías reformuladas (con reducción en el número de indicadores) como aquellas formadas por la totalidad de los indicadores originales. Dicha figura se presenta a continuación:

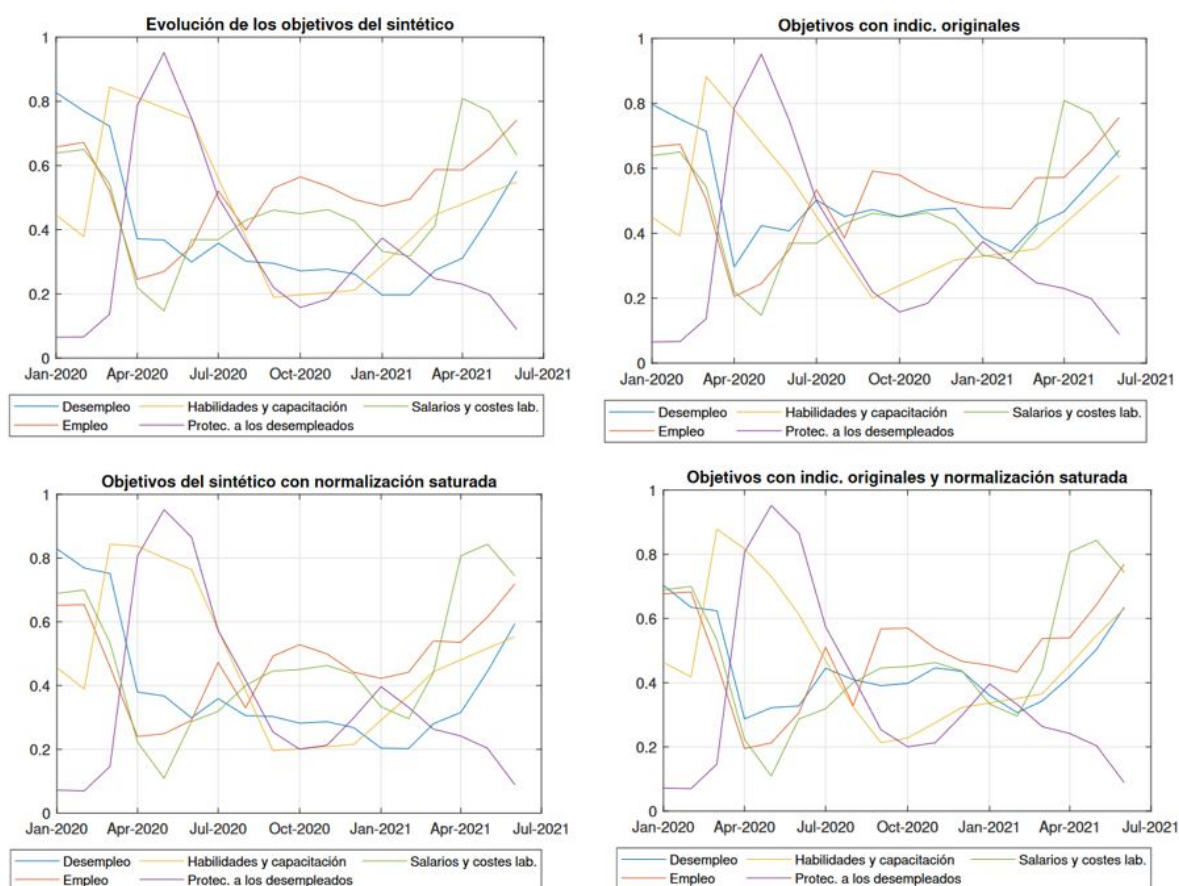


Figura 20. Evolución de los objetivos del sintético según las diferentes técnicas de la elaboración

A partir del comportamiento de cada gráfica, se pueden extraer varias conclusiones sobre el impacto que tiene cada técnica sobre las categorías estudiadas. La saturación de *outliers* o valores atípicos goza de un efecto de compresión sobre las funciones dibujadas, tal y como se observa en el comportamiento de *Desempleo* en las ilustraciones de las categorías con indicadores originales. En lugar de alcanzar su máximo en 0.8, este valor pico inicial se comprime hasta 0.7 con la saturación, lo que deja entrever que los datos del primer trimestre de 2020 son muy positivos en comparación con el resto de meses. De hecho, que la categoría saturada presente una serie temporal en general inferior a su equivalente sin saturación indica la presencia de *outliers* inferiores. Al fin y al cabo, esta transformación consiste en acercar el valor atípico al resto de valores de forma que estos puedan tomar una distribución más uniforme entre el 0 y el 1. Así, un valor atípico muy elevado generará una serie normalizada que, sin saturar, presentará valores muy cercanos al cero, pues la unidad se verá caracterizada por dicho *outlier*. Al saturarlo en el valor establecido por los límites de la fórmula 6, toda la serie crecerá hacia el 1, mientras que comparativamente ese valor se verá perjudicado. Por ello, además de la compresión mencionada, la saturación de valores atípicos desplaza la serie temporal en el sentido del *outlier* saturado.

De esta manera, la saturación afecta de forma dual a los resultados. Pese a ello, en estos casos, la compresión y el desplazamiento no ocurren de forma evidente en todas las categorías ya que, aunque dentro de ellas la saturación sí afecte a algunos indicadores, la ponderación y subagregación de los susodichos camufla dicho fenómeno, a menos que ocurra de forma simultánea en varios indicadores componentes de la misma categoría.

Asimismo, la figura 20 también permite comparar la forma que tendrían los objetivos fundamentales del indicador si se mantuvieran las categorías originales con todas las variables. Por supuesto, los subgrupos *Protección a los desempleados* y *Salarios y costes laborales* no presentan cambio alguno, ya que no resultaban problemáticas en su momento (sección 3.4). En cuanto a las demás, *Empleo* prácticamente no se ve alterada por la eliminación de dos variables redundantes. De forma similar, la categoría *Habilidades y capacitación*, la más modificada, tampoco presenta un cambio notable de comportamiento. Por el contrario, el subgrupo *Desempleo*, como consecuencia de la reformulación, muestra un valle más profundo y crítico durante los meses pandémicos.

Considerando las decisiones tomadas en la sección 3.6, se deduce de los hechos descritos que la plataforma privada que proporciona datos relativos al empleo presentó unas tendencias menos pesimistas que los organismos públicos durante la pandemia, puesto que la eliminación de parte de sus indicadores ha rebajado el valor general de la serie. De cualquier forma, el impacto del cambio no es del todo manifiesto, por lo que las decisiones tomadas relativas a la evaluación intermedia de resultados y a la normalización no deberían tener mucha influencia sobre el sintético final.

Una vez analizado la influencia de las decisiones sobre las categorías, cabe plantearse el mismo fenómeno sobre el sintético final, donde también ha de valorarse el impacto de la agregación geométrica empleada. Los diferentes resultados que se obtendrían se presentan en la figura 21 ilustrada a continuación:

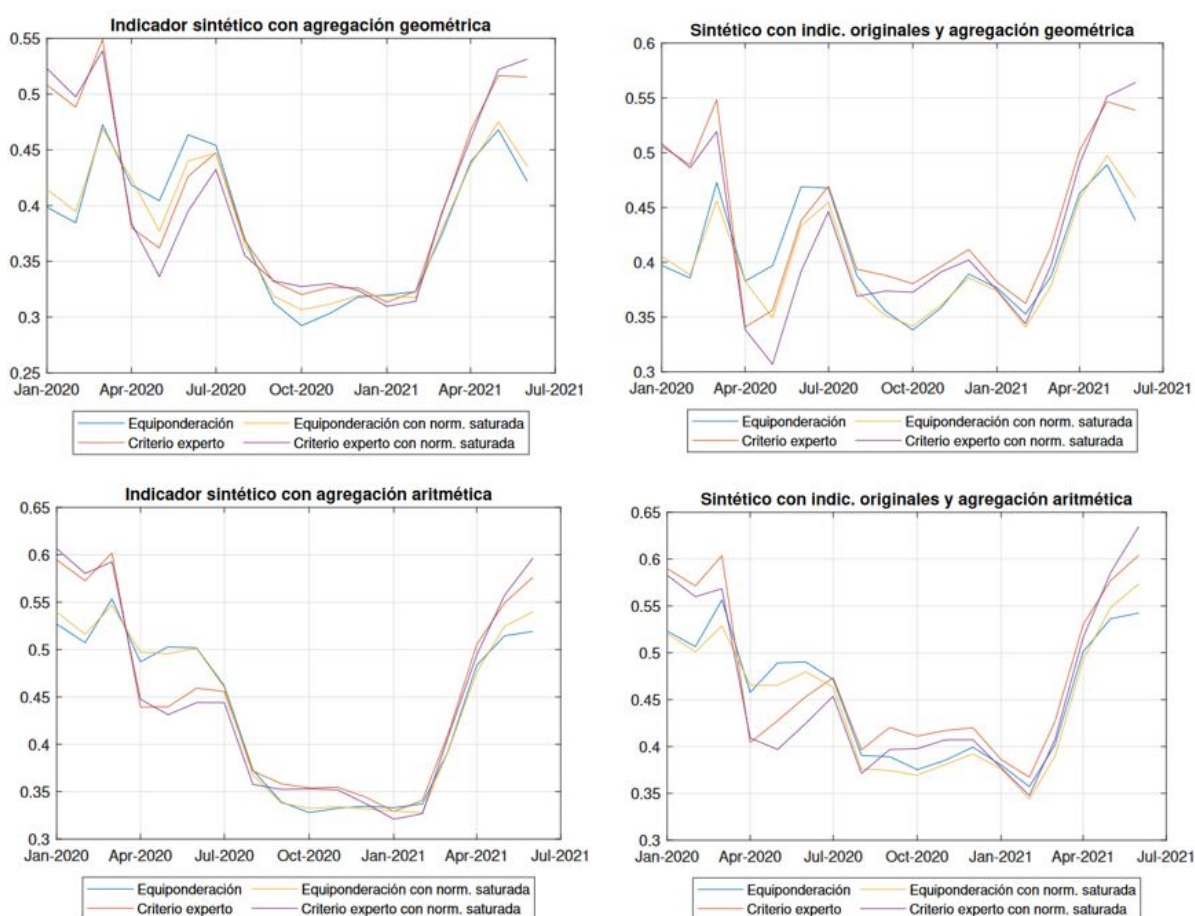


Figura 21. Evolución del indicador sintético según las diferentes técnicas de la elaboración

Cada gráfica, diferente según el método de agregación aplicado y la reformulación o no de las categorías, representa cuatro funciones que reflejan también las diferentes técnicas de ponderación y normalización consideradas. En total se disponen 16 resultados distintos, cada uno representativo de una de las ramas presentadas en la figura 19. De esta manera, diversas conclusiones se pueden extraer para cada uno de los factores X_i considerados como decisiones en el proceso de elaboración:

- **Agregación:** La agregación geométrica genera un conjunto de resultados que, en comparativa, presenta valores de desempeño inferiores. Este método, en tanto que multiplica los valores en lugar de sumarlos, penaliza con fuerza que alguna o varias categorías no rindan adecuadamente. Esto se ve claramente en la escala de cada gráfica, donde aquellas dotadas de agregación aritmética disponen de valores máximos más altos. Asimismo, la agregación geométrica genera un dibujo ligeramente distinto en los meses de abril a julio de 2020, donde los indicadores sintéticos recuperan algo de vigor y alcanzan valores similares a los proporcionados por la agregación aritmética. Esto se atribuye a la recuperación de *Salarios y costes laborales* y *Empleo* que, tras marcar mínimos, se recuperan ligeramente y dejan de lastrar con tanta fuerza al sintético final.

Por ello, aunque incluso las otras categorías se vean mermadas durante esos meses, el funcionamiento de la agregación geométrica premia la mejora de aquella área con peor desempeño. De hecho, esto ocurre al contrario en la agregación aritmética, donde el carácter compensatorio puro de la técnica equilibra la subida de unas categorías con la bajada de otras. En resumen, la elección de la técnica de agregación goza de un impacto notorio sobre los resultados obtenidos, de forma que altera la evolución descrita por el sintético debido especialmente a la sobrepenalización de los malos desempeños a lo largo de la serie temporal.

- **Ponderación:** Representada por las funciones azules y rojas, la aplicación de una u otra técnica influye con fuerza sobre el resultado final. La equiponderación genera un sintético con menores extremos, puesto que proporciona mayor compensación entre categorías como *Desempleo* y *Protección a los desempleados*, que tienden a tomar comportamientos opuestos; ya que, por ejemplo, cuando uno va bien el otro no es necesario. Esta técnica de ponderación genera que, especialmente en los primeros meses de pandemia, las medidas de choque contra el desempleo y los costes laborales compensen los malos rendimientos del paro y el empleo.

Al mismo tiempo, los indicadores contruidos con esta técnica no alcanzan los valores proporcionados por el criterio experto ni en los meses prepandémicos ni en los posteriores al levantamiento del último estado de alarma (abril de 2021). En cambio, el criterio experto presenta un empeoramiento y una recuperación más acusadas al principio y al final de la serie temporal, en tanto que rebaja la compensatoriedad de las categorías que muestran los efectos de la salud del mercado laboral (consecuencias) y las causas que los generan (ver 6). Por estas razones, el método de ponderación de las categorías es fundamental para abordar la elaboración del sintético de forma acertada y rigurosa, puesto que la equiponderación no tiene en cuenta el comportamiento y las interacciones entre las susodichas.

- **Reformulación de las categorías:** Los sintéticos producidos a partir de las categorías reducidas presentan menores diferencias cuando se emplean técnicas distintas de normalización, ponderación y agregación. Las funciones muestran comportamientos similares, es decir, que la reformulación proporciona robustez al indicador sintético. Esta fiabilidad vendría a ser comprendida como la capacidad del modelo de proporcionar resultados similares ante la existencia de perturbaciones, las cuales serían en este caso la variación de los factores X_i .

Por ello, estos hechos se atribuyen directamente a lo mentado en la sección 3.6, donde se reducen ciertas categorías conflictivas en pos de la fiabilidad y robustez del sintético. Aunque el efecto no es tampoco algo incuestionable, sí se percibe cómo las funciones que emplean la totalidad de los indicadores originales presentan comportamientos ligeramente más erráticos, con pequeños cambios en el comportamiento y forma de las funciones. Por esta razón, se ratifica la decisión tomada en la sección 3.6, en tanto que no altera excesivamente los resultados pero sí proporciona robustez al sintético obtenido.

- **Normalización:** Como colofón, la normalización funciona de forma similar a la explicada para la figura 20. Sus implicaciones, por tanto, son complejas, ya que la existencia de valores atípicos contribuye a modificar la serie en el sentido de dicho *outlier*, aunque la multiplicidad de su existencia puede dificultar este análisis, a la par que la posterior ponderación diluye los indicadores componentes en los factores de la técnica PCA. Por ello, evaluar con tino las causas de cada movimiento relativo de las funciones saturadas y no alteradas se torna complicado. Sin embargo, dicha saturación no afecta de forma especialmente relevante a los resultados, especialmente en las categorías reformuladas y, por tanto, construidas de forma robusta. Las dos funciones generadas por las variaciones de la técnica se mantienen muy cercanas en el tiempo describiendo, además, el mismo tipo de comportamiento. Por lo tanto, si bien cuantificar este impacto es harto difícil debido a la cantidad de procesos posteriores aplicados, la figura 21 no presenta grandes variaciones para uno u otro método.

De esta manera, se concluye el estudio cualitativo de las fuentes de incertidumbre como parte del análisis de robustez. Las técnicas que impactan en mayor medida de forma cualitativa sobre el resultado del sintético son la agregación y la ponderación, puesto que alteran los comportamientos y los valores tomados por las series temporales consideradas. La reformulación de las categorías proporciona resultados ligeramente más cercanos entre sí, mientras que la normalización no afecta de forma notoria en la mayoría de los casos, especialmente en aquellos cercanos al sintético que sí es producto de las decisiones del proyecto. Este vendría a estar representado por la línea roja de la gráfica superior izquierda de la figura 21, aunque fue debidamente presentado a través de la figura 18 de la página 94.

En conclusión, cada decisión goza de su propio papel en la determinación del sintético final y, hasta el análisis de robustez, resulta inviable conocer con exactitud cómo influyen dichas elecciones. Por ello, resulta imprescindible mantener la transparencia y el buen hacer en cada etapa del proceso, de tal manera que se eviten las “cajas negras” que arrojan resultados a partir de un mecanismo desconocido, al menos a ojos del lector. Evitar este tipo de decisiones, puesto que estos son frecuentes causas de errores de interpretación por la sobresimplificación de su naturaleza. Esto, sumado a que pueda darse algún desacierto o excesiva arbitrariedad en su creación, puede resultar perjudicial para la realidad que aspiran a transformar (Grupp y Mogege,

2004).

Finalmente, la combinación de todos los resultados posibles proporciona la figura 22 con diagramas de caja y bigotes para cada fecha estudiada y comparadas con el indicador sintético obtenido en el proceso de elaboración:

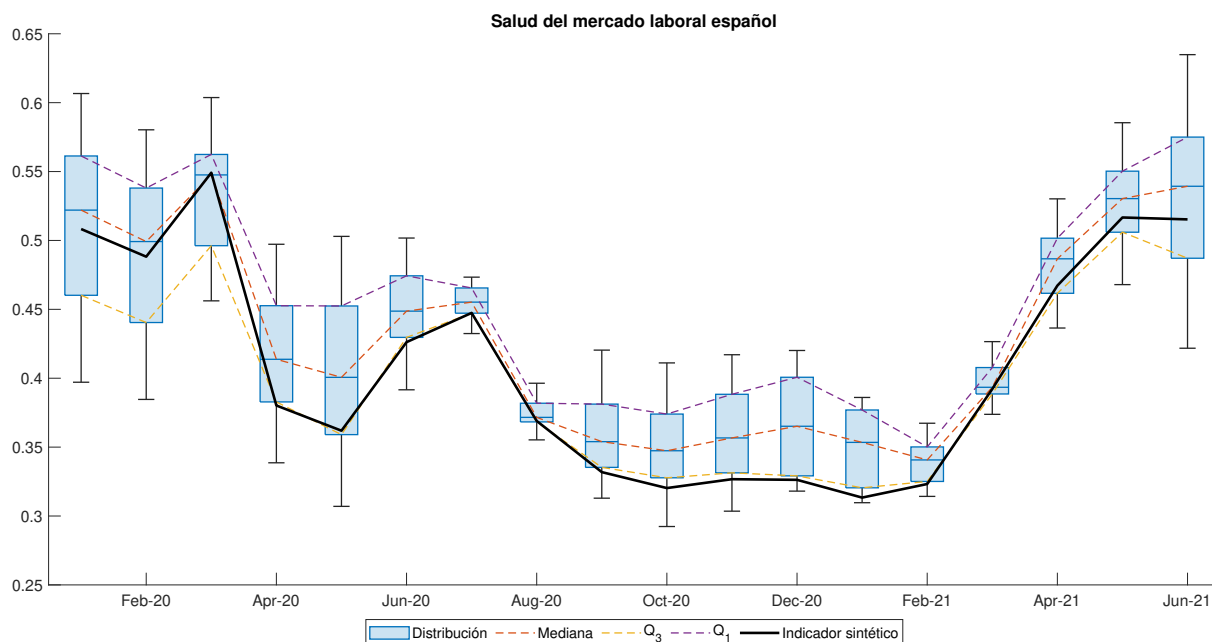


Figura 22. Evolución del sintético elaborado frente a la distribución de sus variaciones

En esta figura se observan claramente las características definitorias del sintético elaborado, en tanto que emplea una agregación geométrica con criterio experto y la técnica de normalización Mín-Máx. Comparativamente, la agregación geométrica provoca que el resultado obtenido se encuentre bajo la mediana en la práctica totalidad la serie temporal, puesto que esta técnica perjudica los malos desempeños. Esto es importante porque se alinea con los objetivos del sintético, en tanto que ninguna de las categorías debería ser maltratada u olvidada por los agentes sociales para dotar de salud al mercado laboral español. Asimismo, el criterio experto proporciona la forma de la gráfica, tal y como se ha explicado anteriormente. Esta ponderación reduce aún más la sustituibilidad de ciertas categorías, aunque facilita la compensación llevada a cabo por aquellas con mayor peso (*Desempleo* y *Empleo*). No obstante, el sintético obtenido cumple con lo esperado y los estándares de fiabilidad y robustez analizados a lo largo de la sección.

El análisis de sensibilidad basado en la varianza no es estrictamente necesario tras el estudio de incertidumbre realizado a lo largo de la sección. De esta manera, dado que el indicador sintético construido ha quedado totalmente caracterizado, se da por concluido su proceso de elaboración considerando también los elementos visuales aportados como parte de la labor de divulgación del proyecto.

Capítulo 4

Conclusiones

Tras la obtención de resultados, la elaboración del indicador sintético ha llegado a su fin, proporcionando en el proceso importantes enseñanzas y fundamentos para futuros trabajos. De esta manera, existen una serie de conclusiones extraíbles del proyecto, comenzando desde el propio inicio del trabajo:

- La organización de información otorga a la misma de un valor añadido, en tanto que facilita su comprensión y permite monitorizar el impacto de la aplicación de iniciativas, ya sea en el mundo público o privado. De esta manera, la carencia de un sintético que evaluara y agrupara la multitud de indicadores referentes al mercado de trabajo suponía un lastre para cualquier interesado en examinar su situación. Especialmente, para aquellos que no gozan de herramientas potentes de análisis estadístico o *Big Data*, ya que no cuentan con criterios claros y fiables para el descarte o la agregación de tanta información. El indicador *salud del mercado laboral español* surge para mitigar esta problemática.
- La elaboración de indicadores sintéticos está repleta de decisiones en todas y cada una de las etapas, desde el establecimiento del propio marco teórico donde se dota de estructura al futuro sintético hasta el análisis de robustez que tiene por objetivo abordar las fuentes de incertidumbre y fiabilidad del proceso. Por ello, resulta de vital importancia llevar a cabo un ejercicio de transparencia y claridad, de manera que el lector comprenda cada elección y los motivos que la impulsaron. Asimismo, el autor de los indicadores debe velar por introducir en sus valoraciones tantas perspectivas y enfoques como sea posible, de manera que todas las partes que pudieran interesarse por el sintético vean representadas sus inquietudes o intereses. De esta manera, el indicador elaborado encontrará mayor respaldo social y gozará de enfoques que enriquezcan su construcción.
- El capítulo 2 presenta una recopilación crítica y actualizada de la situación actual de las técnicas aplicables a las diferentes etapas de la elaboración. Este capítulo cumple con uno de los objetivos del proyecto, en tanto que proporciona un nuevo complemento para la construcción de indicadores a fecha de junio de 2022. Al fin y al cabo, las técnicas empleadas de análisis, imputación, ponderación... gozan de un potencial de aplicación mucho mayor que el concerniente a los indicadores sintéticos, ya que están presentes en cualquier situación donde se requiera cierto estudio estadístico. Por ello, las técnicas utilizadas en la construcción de indicadores pueden experimentar avances en ocasiones no reflejados por la bibliografía de los sintéticos, razón por la cual resultaba fundamental realizar una revisión renovada de su situación. Después de todo, el manual de la OCDE

(2008) pronto cumplirá 15 años.

- Posteriormente, las enseñanzas halladas en dicho capítulo se aplican en la elaboración efectiva del indicador del proyecto (capítulo 3), de forma que a lo largo de su creación se dan numerosas referencias cruzadas a partes explicadas o comentadas con anterioridad en la parte de diseño teórico (capítulo 2). Básicamente, en primer lugar se asienta el marco conceptual sobre la construcción de sintéticos y, posteriormente, se realiza la aplicación experimental de los hallazgos de cada etapa conceptual. Este diseño teórico se considera uno de los objetivos fundamentales del trabajo, puesto que aporta tanto valor como lo hace el indicador construido *per se*.
- Las diferentes técnicas empleadas responden siempre a criterios de aplicabilidad, pertinencia y sentido común. En aquellas etapas en las que no existe un único itinerario – análisis multivariante, por ejemplo– se han propuesto multitud de métodos que satisfagan eventualmente la curiosidad del lector y caractericen los datos tanto como sea posible. Para salvaguardar la trazabilidad de las elecciones y su transparencia, se han argumentado todas las decisiones tomadas, aunque en ocasiones la prelación entre uno u otro método sea difusa. Pese a que resulta inviable elaborar un sintético sin ningún tipo de arbitrariedad, decisiones relevantes como la aplicación de normalización Mín-Máx, la reformulación de las categorías y las técnicas de ponderación y agregación han sido evaluadas posteriormente en la sección 3.8. De esta manera, si bien la idoneidad de cada decisión no se puede demostrar fehacientemente, se proporciona al lector de las herramientas para juzgar por sí mismo la adecuación de las mismas.
- Los resultados obtenidos son coherentes con los eventos acaecidos durante los meses pandémicos. La mejoría del mercado laboral, coincidente con la reactivación parcial del turismo en los meses estivales en 2021 y con la vacunación general de la sociedad, sucedía a una depresión de dimensiones otrora nunca vistas, donde los indicadores individuales de las diferentes categorías presentan en su mayoría los mínimos de la serie considerada. No obstante, también es cierto que la aplicación de medidas de choque por parte de los agentes sociales contribuyeron a mitigar los efectos de la pandemia sobre el mercado laboral, razón por la cual el desplome del indicador no es totalmente abrupto tras marzo de 2020. Pese a ello, la aplicación de estas medidas sigue afectando a muchas empresas y particulares, así como al propio Estado que ha visto cómo la deuda pública sigue creciendo de forma importante. Todos estos fenómenos, si bien explican parte del comportamiento del indicador sintético elaborado, han de ser evaluados por un experto en la cuestión, ya que el comportamiento de una realidad tan compleja como el mercado laboral español no debe resolverse en unas simples líneas.

Futuras líneas de investigación

Por todo ello, se propone el análisis del comportamiento del indicador como unas de las futuras líneas de investigación, de manera que se profundice debidamente en los factores propios y ajenos al mercado laboral que han contribuido a la caracterización del fenómeno en cuestión. Esto requerirá de un análisis multidimensional y holístico al alcance de expertos economistas.

Asimismo, la aplicación del proceso de elaboración del sintético a otras áreas de la plataforma 360 Smart Vision puede resultar muy edificante para la sociedad, especialmente si se complementa con el análisis del experto antes mencionado. Después de todo, la creación de sintéticos para las 8 áreas presentes en el *dashboard* (consumo, sanidad, mercados financieros, sostenibilidad, actividad económica...) puede ayudar a caracterizar esos factores que influyen sobre el mercado de trabajo u otras realidades. Por ello, estas dos líneas de investigación podrían retroalimentarse entre sí, proporcionando por tanto un análisis completo de la realidad socioeconómica española que podría basarse en la elaboración de un sintético general para todo el panel 360 Smart Vision. No obstante, un indicador tan ambicioso requerirá de un grupo de trabajo multidisciplinar, puesto que se considerarían fenómenos sociales muy dispares no abordables por una única persona. A fin de cuentas, al igual que en el indicador elaborado en este proyecto, incluir multitud de enfoques resulta enriquecedor para el propio sintético creado, ya que puede reflejar de forma más acertada los diferentes intereses sociales.

De esta manera, tras proponer futuras líneas de investigación y presentar los resultados al público general, el proyecto cumple con todos los objetivos propuestos. A partir de este momento, la sociedad española y europea disponen de una herramienta que permite monitorizar las acciones emprendidas por los agentes sociales sobre el mercado laboral español, de manera que se comience a dilucidar qué iniciativas son más exitosas y efectivas en su labor. Asimismo, también se proporciona al mundo académico de un complemento actualizado para la elaboración de futuros indicadores sintéticos. El objetivo dual del proyecto, por tanto, es alcanzado a través de dos impactos valiosos para la sociedad en su conjunto.

Por todo ello, el indicador sintético logra asentar las bases de futuros trabajos a la par que proporciona información útil a los agentes con voluntad de transformación social. De esta manera, se da por concluido el proyecto.

Bibliografía

- Abdella, G. M., Kucukvar, M., Ismail, R., Abdelsalam, A. G., Onat, N. C., y Dawoud, O. (2021). A mixed model-based Johnson's relative weights for eco-efficiency assessment: The case for global food consumption. *Environmental Impact Assessment Review*, 89:106588.
- Abell, C. L. B. (2010). Amartya Sen y el desarrollo humano. *Revista Memorias*. Disponible en: <https://docplayer.es/46613332-Amartya-sen-y-el-desarrollo-humano.html> [Consultado 17-05-2022].
- Acemoglu, D. y Autor, D. (2011). Handbook of labor economics. *Elsevier*, 4. Disponible en: <https://economics.mit.edu/faculty/dautor/data/acemoglu> [Consultado 20-06-2022].
- Alam, M., Dupras, J., y Messier, C. (2016). A framework towards a composite indicator for urban ecosystem services. *Ecological Indicators*, 60:38–44.
- Alejandro Martínez Vélez (2022). El Gobierno permite ampliar hasta diez años el plazo de amortización de los préstamos avalados por el ICO. *El Español*. Disponible en: https://www.elespanol.com/invertia/empresas/banca/20220621/gobierno-permite-ampliar-amortizacion-prestamos-avalados-ico/681932208_0.html [Consultado 25-06-2022].
- Alon, T., Doepke, M., Olmstead-Rumsey, J., y Tertilt, M. (2020). The impact of the coronavirus pandemic on gender equality. p. 6.
- Alonso, L. E. y Salmerón, L. T. (2003). Trabajo sin reconocimiento o la especial vulnerabilidad de las mujeres jóvenes en el mercado laboral. *Cuadernos de Relaciones Laborales*, p. 37.
- Arnal, M., Finkel, L., y Parra, P. (2013). Crisis, desempleo y pobreza: análisis de trayectorias de vida y estrategias. *Cuadernos de Relaciones Laborales*, 31(2):281–311.
- Arrow, K. y Raynaud, H. (1986). Social Choice and Multicriterion Decision-Making. *The MIT Press*. Disponible en: <https://ideas.repec.org/b/mtp/titles/0262511754.html> [Consultado 17-06-2022].
- Becker, W., Paruolo, P., Saisana, M., y Saltelli, A. (2017). Weights and Importance in Composite Indicators: Mind the Gap. En Ghanem, R., Higdon, D., y Owhadi, H., editores, *Handbook of Uncertainty Quantification*, pp. 1187–1216. Springer International Publishing, Cham.
- Benito, S. R., Navarro, J. L. M., Ortiz, L. P., y Nestares, Y. C. R. (2004). La negociación colectiva en España: Análisis económico. p. 31.
- Bennett, D. A. (2001). How can I deal with missing data in my study? *Australian and New Zealand Journal of Public Health*, 25(5).
- Burck, J., Bals, C., y Bohnenberger, K. (2011). The Climate Change Performance Index 2012; Der Klimaschutz-Index. Ergebnisse 2012.

- Burgass, M. J., Halpern, B. S., Nicholson, E., y Milner-Gulland, E. (2017). Navigating uncertainty in environmental composite indicators. *Ecological Indicators*, 75:268–278.
- Buuren, S. v. (2018). *Flexible Imputation of Missing Data, Second Edition*. CRC Press.
- Calafati, R. O. (2017). Estrategias para el tratamiento de datos faltantes ("missing data") en estudios con datos longitudinales. Disponible en: <http://openaccess.uoc.edu/webapps/o2/bitstream/10609/64085/6/romancalafatiTFG0617memoria.pdf> [Consultado 03-01-2022].
- Caligiuri, G., Paulsson, G., Nicoletti, A., Maseri, A., y Hansson, G. K. (2000). Evidence for Antigen-Driven T-Cell Response in Unstable Angina. *Circulation*, 102(10):1114–1119. Publisher: American Heart Association.
- Carlos E. Cué (2022). Coronavirus: El Gobierno ordena la hibernación de la economía para evitar el colapso del sistema sanitario. *El País*. Disponible en: <https://elpais.com/espana/2020-03-29/el-gobierno-ordena-la-hibernacion-de-la-economia-para-evitar-el-colapso-del-sistema-sanitario.html> [Consultado 27-06-2022].
- Cherchye, L., Moesen, W., Rogge, N., y Puyenbroeck, T. V. (2007). An Introduction to Benefit of the Doubt Composite Indicators. *Social Indicators Research*, 82(1):111–145.
- Cherchye, L., Moesen, W., Rogge, N., Van Puyenbroeck, T., Saisana, M., Saltelli, A., Liska, R., y Tarantola, S. (2008). Creating composite indicators with DEA and robustness analysis: the case of the Technology Achievement Index. *Journal of the Operational Research Society*, 59(2):239–251.
- Chumacero, J. A. (2015). Detección de la multicolinealidad y heterocedasticidad. *Universidad Nacional Mayor de San Marcos*. Disponible en: <https://www.studocu.com/pe/document/universidad-nacional-mayor-de-san-marcos/econometria-i/deteccion-de-la-multicolinealidad-y-heteroscedasticidad/9763391> [Consultado 01-06-2022].
- Cortina, J. M. (1993). What is coefficient alpha? An examination of theory and applications. *Journal of Applied Psychology*, 78(1):98–104.
- CRITERIA (2022). *Observatorio Criteria: un análisis económico y social*. Disponible en: <https://www.comillas.edu/noticias/63-comillas-icade/icade-cee/cee-investigacion/2351-observatorio-criteria-analisis-de-la-situacion-en-espana> [Consultado 01-11-2021].
- Cronbach, L. J. (1951). Coefficient alpha and the internal structure of tests. *Psychometrika*, 16(3):297–334.
- Cuadras, C. M. (2007). *Nuevos métodos de análisis multivariante*. CMC Editions, Barcelona.
- Cámara de Comercio de España (2021). El año 2020 concluye con un mercado laboral muy debilitado y alejado de su plena recuperación en 2021 | Cámara de España. Disponible en: <http://www.camara.es/encuesta-poblacion-activa-cuarto-trimestre-2020> [Consultado 27-06-2022].
- Davino, C. y Romano, R. (2014). Assessment of Composite Indicators Using the ANOVA Model Combined with Multivariate Methods. *Social Indicators Research*, 119(2):627–646.
- Deloitte (2021). *360 Smart Vision | Deloitte España*. Disponible en: <https://www2.deloitte.com/es/es/pages/strategy/articles/360-smart-vision.html#dashboard> [Consultado 01-10-2021].

- Dialga, I. y Thi Hang Giang, L. (2017). Highlighting Methodological Limitations in the Steps of Composite Indicators Construction. *Social Indicators Research*, 131(2):441–465.
- Donders, A. R. T., van der Heijden, G. J. M. G., Stijnen, T., y Moons, K. G. M. (2006). Review: A gentle introduction to imputation of missing values. *Journal of Clinical Epidemiology*, 59(10):1087–1091.
- Dziechciarz, J., Walesiak, M., y Bk, A. (1999). An application of conjoint analysis for preference measurement. *Argumenta Oeconomica*, 1(7):169–178.
- EC (2022). Step 5: Normalisation: Knowledge for policy. Disponible en: https://knowledge4policy.ec.europa.eu/composite-indicators/10-step-guide/step-5-normalisation_en [Consultado 02-04-2022].
- Emerson, Jay, Esty, Daniel C., Kim, Christine, Srebotnjak, Tanja, Levy, Marc A., Mara, Valentina, De Sherbinin, Alex, y Jaiteh, Malanding (2008). 2010 Environmental Performance Index. *Choice Reviews Online*, 45(05).
- Eurostat (2017). Código de buenas prácticas de las estadísticas europeas. Technical report, Sistema Estadístico Europeo, Luxemburgo.
- Expansión (2022). ¿Qué significa mercado de trabajo? Disponible en: <https://www.expansion.com/economia-para-todos/economia/que-significa-mercado-de-trabajo.html> [Consultado 10-12-2021].
- Floridi, M., Pagni, S., Falorni, S., y Luzzati, T. (2011). An exercise in composite indicators construction: Assessing the sustainability of Italian regions. *Ecological Economics*, 70(8):1440–1447.
- French, A., Macedo, M., Poulsen, J., Waterson, T., y Yu, A. (2008). Multivariate Analysis of Variance (MANOVA). Disponible en: <https://docplayer.net/20883992-Multivariate-analysis-of-variance-manova.html> [Consultado 01-06-2022].
- Freudenberg, M. (2003). Composite Indicators of Country Performance: A Critical Assessment. OECD Science, Technology and Industry Working Papers 2003/16.
- Gan, X., Fernandez, I. C., Guo, J., Wilson, M., Zhao, Y., Zhou, B., y Wu, J. (2017). When to use what: Methods for weighting and aggregating sustainability indicators. *Ecological Indicators*, 81:491–502.
- Gold, M. S. y Bentler, P. M. (2000). Treatments of Missing Data: A Monte Carlo Comparison of RBHDI, Iterative Stochastic Regression Imputation, and Expectation-Maximization. *Structural Equation Modeling: A Multidisciplinary Journal*, 7(3):319–355.
- Greco, S., Ishizaka, A., Tasiou, M., y Torrisi, G. (2019). On the Methodological Framework of Composite Indices: A Review of the Issues of Weighting, Aggregation, and Robustness. *Social Indicators Research*, 141(1):61–94.
- Groeneveld, R. A. y Meeden, G. (1984). Measuring Skewness and Kurtosis. *Journal of the Royal Statistical Society. Series D (The Statistician)*, 33(4):391–399. Publisher: [Royal Statistical Society, Wiley].
- Grupp, H. y Mogee, M. E. (2004). Indicators for national science and technology policy: how robust are composite indicators? *Research Policy*, 33(9):1373–1384.

- Gutiérrez, H. (2021). El turismo nacional se desboca y supera en julio los récords previos a la pandemia. *El País*. Disponible en: <https://elpais.com/economia/2021-08-24/el-turismo-nacional-se-desboca-y-supera-en-julio-los-niveles-prepandemia.html> [Consultado 10-06-2021].
- Guttman, L. (1954). A new approach to factor analysis: the Radex. En *Mathematical thinking in the social sciences*, pp. 258–348. Free Press, New York, US.
- Guz, T., Sengun, G., e Incekara, A. (2017). Measuring the technology achievement index: comparison and ranking of countries. *Pressacademia*, 4(2):164–174.
- Gómez, . P., Pacheco, J. A., y Herrero, M. J. L. (2017). *Comparativa de análisis de imputación de datos faltantes con análisis de casos completos en pruebas diagnósticas*. Tesis doctoral, Universidad Complutense.
- Gómez Bengoechea, G. (2020). Tasa de Paro y ERTes: una relación incomprendida. *360 Smart Vision*. Disponible en: <https://360smartvision.comillas.edu/tasa-de-paro-y-ertes-una-relacion-incomprendida/> [Consultado 15-06-2022].
- Halpern, B. S., Longo, C., Hardy, D., McLeod, K. L., Samhour, J. F., Katona, S. K., Kleisner, K., Lester, S. E., OLeary, J., Ranelletti, M., Rosenberg, A. A., Scarborough, C., Selig, E. R., Best, B. D., Brumbaugh, D. R., Chapin, F. S., Crowder, L. B., Daly, K. L., Doney, S. C., ..., y Zeller, D. (2012). An index to assess the health and benefits of the global ocean. *Nature*, 488(7413):615–620.
- Haq, M. U. (1995). El paradigma del desarrollo humano. Disponible en: <https://isfcolombia.uniandes.edu.co/images/documentos/paradigma%20de%20desarrollo%20humano%201.pdf> [Consultado 16-04-2022].
- Hartigan, J. A. (1975). *Clustering Algorithms*. John Wiley Sons.
- Heo, M., Kim, N., y Faith, M. S. (2015). Statistical power as a function of Cronbach alpha of instrument questionnaire items. *BMC Medical Research Methodology*, 15(1):86.
- Hernández Díaz, G. A. y Marín Jaramillo, M. (2016). Pronóstico del Consumo Privado: Usando datos de alta frecuencia para el pronóstico de variables de baja frecuencia. Disponible en: <https://ideas.repec.org/p/col/000118/014828.html> [Consultado 27-05-2022].
- Horton, N. J. y Kleinman, K. P. (2007). Much Ado About Nothing: A Comparison of Missing Data Methods and Software to Fit Incomplete Data Regression Models. *The American Statistician*, 61(1):79–90.
- Hsu, A. y Zomer, A. (2016). Environmental Performance Index. *Wiley StatsRef: Statistics Reference Online*, pp. 1–5.
- Hudrliková, L. (2013). Composite Indicators as a Useful Tool for International Comparison: The Europe 2020 Example. *Prague Economic Papers*, 22(4):459–473.
- IBM (2021). IBM SPSS Missing Values 28. Disponible en: https://www.ibm.com/docs/en/SSLVMB_28.0.0/pdf/IBM_SPSS_Missing_Values.pdf [Consultado 13-02-2022].
- Illahi, U. y Mir, M. S. (2021). Sustainable Transportation Attainment Index: multivariate analysis of indicators with an application to selected states and National Capital Territory (NCT) of India. *Environment, Development and Sustainability*, 23(3):3578–3622.

- Illes, A. y Lombardi, M. (2013). Interest rate pass-through since the financial crisis. p. 10.
- Instituto de Salud Carlos III (2022). COVID-19. Disponible en: <https://cnecovid.isciii.es/covid19/> [Consultado 13-05-2022].
- J, K. y W, M. C. (1978). Factor Analysis - Statistical Methods and Practical Issues | Office of Justice Programs.
- Jackson, D. A. (1993). Stopping Rules in Principal Components Analysis: A Comparison of Heuristical and Statistical Approaches. *Ecology*, 74(8).
- Jain, A., Nandakumar, K., y Ross, A. (2005). Score normalization in multimodal biometric systems. *Pattern Recognition*, 38(12):2270–2285.
- Jolliffe, I. (2002). *Principal Components Analysis, Second Edition*. Springer, Nueva York. Disponible en: [http://cda.psych.uiuc.edu/statistical_learning_course/Jolliffe%20I.%20Principal%20Component%20Analysis%20\(2ed.,%20Springer,%202002\)\(518s\)_MVsa_.pdf](http://cda.psych.uiuc.edu/statistical_learning_course/Jolliffe%20I.%20Principal%20Component%20Analysis%20(2ed.,%20Springer,%202002)(518s)_MVsa_.pdf) [Consultado 22-05-2022].
- José Antonio, P.-G., Vicent, A.-L., y Ramón, F.-P. (2022). A composite indicator index as a proxy for measuring the quality of water supply as perceived by users for urban water services. *Technological Forecasting and Social Change*, 174:121300.
- Korhonen, P., Tainio, R., y Wallenius, J. (2001). Value efficiency analysis of academic research. *European Journal of Operational Research*, 130(1):121–132.
- Lebart, L., Warwick, K. M., y Morineau, A. (1984). *Multivariate descriptive statistical analysis: correspondence analysis and related techniques for large matrices*.
- Lee, G. K. L. y Chan, E. H. W. (2008). The Analytic Hierarchy Process (AHP) Approach for Assessment of Urban Renewal Proposals. *Social Indicators Research*, 89(1):155–168.
- Li, C. (2013). Little's Test of Missing Completely at Random. *The Stata Journal*, 13(4):795–809.
- Little, R. J. A. (1988). A Test of Missing Completely at Random for Multivariate Data with Missing Values. *Journal of the American Statistical Association*, 83(404):1198–1202.
- Lorenzo-Seva, U. y Ferrando, P. J. (2021). Not Positive Definite Correlation Matrices in Exploratory Item Factor Analysis: Causes, Consequences and a Proposed Solution. *Structural Equation Modeling: A Multidisciplinary Journal*, 28(1):138–147.
- Mahlberg, B. y Obersteiner, M. (2001). Remeasuring the HDI by Data Envelopment Analysis. *SSRN Electronic Journal*.
- Malan, L., Smuts, C. M., Baumgartner, J., y Ricci, C. (2020). Missing data imputation via the expectation-maximization algorithm can improve principal component analysis aimed at deriving biomarker profiles and dietary patterns. *Nutrition Research*, 75:67–76.
- Marimon, R. (1997). Una reflexión sobre el desempleo en España. *Els Opuscles del CREI*. Disponible en: https://crei.cat/wp-content/uploads/opuscles/090429182629_ESP_op1cas.pdf [Consultado 21-06-2022].

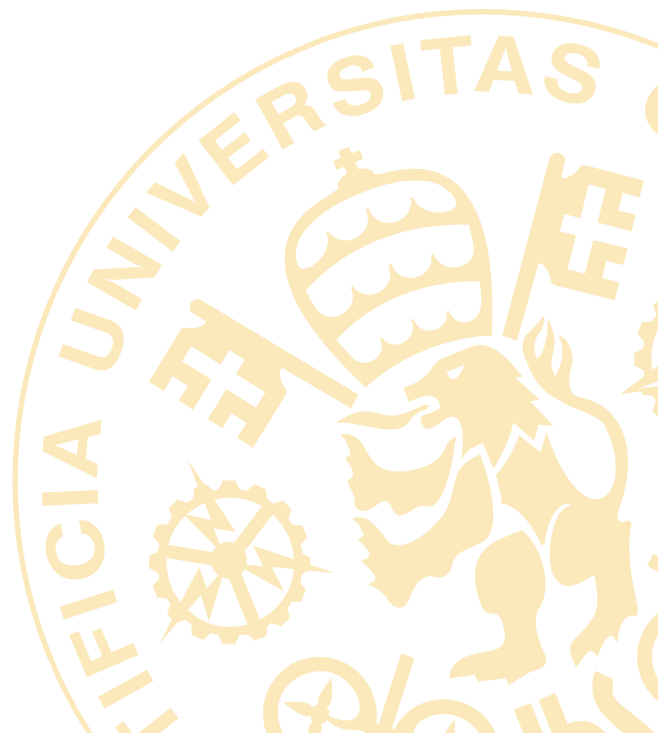
- Marina Gómez García, Laura Hospido Quintana, y Carlos Sanz Alonso (2021). El impacto diferencial por sexos de la crisis del Covid-19 en el mercado de trabajo español. *Informe trimestral de la economía española Banco de España*, pp. 39–42.
- Matthiesen, R., editor (2010). *Bioinformatics Methods in Clinical Research*, volumen 593 de *Methods in Molecular Biology*. Humana Press, Totowa, NJ.
- Mazziotta, M. y Pareto, A. (2017). Synthesis of Indicators: The Composite Indicators Approach. En *Complexity in Society: From Indicators Construction to their Synthesis*, pp. 159–191.
- Medel Esquivel, R., Gomez-Vargas, I., Vazquez, J., y García-Salcedo, R. (2019). *An introduction to Markov Chain Monte Carlo*. Ciudad de México.
- Mertens, B. J. A. (2017). Transformation, Normalization, and Batch Effect in the Analysis of Mass Spectrometry Data for Omics Studies. En Datta, S. y Mertens, B. J. A., editores, *Statistical Analysis of Proteomics, Metabolomics, and Lipidomics Data Using Mass Spectrometry*, Frontiers in Probability and the Statistical Sciences, pp. 1–21. Springer International Publishing, Cham.
- Mtessigwa, W. M. (2013). *Labor Economics*. Disponible en: https://www.academia.edu/6861728/LABOR_ECONOMICS [Consultado 20-02-2022].
- Munda, G. (2012). Choosing Aggregation Rules for Composite Indicators. *Social Indicators Research*, 109(3):337–354.
- Munda, G. y Nardo, M. (2003). On the Methodological Foundations of Composite Indicators Used for Ranking Countries. p. 19.
- Murias, P., de Miguel, J. C., y Rodríguez, D. (2008). A Composite Indicator for University Quality Assesment: The Case of Spanish Higher Education System. *Social Indicators Research*, 89(1):129–146.
- Musil, C. M., Warner, C. B., Yobas, P. K., y Jones, S. L. (2002). A Comparison of Imputation Techniques for Handling Missing Data. *Western Journal of Nursing Research*, 24(7):815–829. Publisher: SAGE Publications Inc.
- Nardo, M., Saisana, M., Saltelli, A., y Tarantola, S. (2005). Tools for Composite Indicators Building. *Joint Research Centre*.
- Newman, D. A. (2014). Missing Data: Five Practical Guidelines. *Organizational Research Methods*, 17(4):372–411. Publisher: SAGE Publications Inc.
- Nicoletti, G., Scarpetta, S., y Boylaud, O. (2000). Summary Indicators of Product Market Regulation with an Extension to Employment Protection Legislation. Technical report, OECD, Paris.
- Núñez, J. A. (2019). Análisis de la ola de calor de junio de 2019 en un contexto de crisis climática. Technical report, AEMET, Madrid. Disponible en: <http://www.aemet.es/documentos/es/noticias/2019/Oladecalorjun2019.pdf> [Consultado 17-05-2022].
- OCDE (2008). *Handbook on Constructing Composite Indicators: Methodology and User Guide*. Disponible en: <https://www.oecd.org/sdd/42495745.pdf> [Consultado 02-10-2021].

- Orchard, T. y Woodbury, M. A. (1972). A missing information principle: Theory and applications. En Le Cam, L. M., Neyman, J., y Scott, E. L., editores, *Theory of Statistics*, pp. 697–716. University of California Press.
- Paruolo, P., Saisana, M., y Saltelli, A. (2013). Ratings and rankings: voodoo or science?: *Ratings and Rankings. Journal of the Royal Statistical Society: Series A (Statistics in Society)*, 176(3):609–634.
- Peres-Neto, P. R., Jackson, D. A., y Somers, K. M. (2005). How many principal components? stopping rules for determining the number of non-trivial axes revisited. *Computational Statistics & Data Analysis*, 49(4):974–997.
- PNUD (1990). *Desarrollo Humano Informe 1990*. Disponible en: <http://desarrollohumano.org.gt/wp-content/uploads/2016/04/HDR-1990.pdf> [Consultado 15-12-2021].
- Profit, J., Typpo, K. V., Hysong, S. J., Woodard, L. D., Kallen, M. A., y Petersen, L. A. (2010). Improving benchmarking by using an explicit framework for the development of composite indicators: an example using pediatric quality of care. *Implementation Science*, 5(1):13.
- Razali, N. M. y Wah, Y. B. (2011). Power comparisons of Shapiro-Wilk, Kolmogorov-Smirnov, Lilliefors and Anderson-Darling tests. p. 14.
- Reisi, M., Aye, L., Rajabifard, A., y Ngo, T. (2014). Transport sustainability index: Melbourne case study. *Ecological Indicators*, 43:288–296.
- Riedler, B., Pernkopf, L., Strasser, T., Lang, S., y Smith, G. (2015). A composite indicator for assessing habitat quality of riparian forests derived from Earth observation data. *International Journal of Applied Earth Observation and Geoinformation*, 37:114–123.
- Roeser, S. y Pesch, U. (2016). An Emotional Deliberation Approach to Risk. *Science, Technology, & Human Values*, 41(2):274–297. Publisher: SAGE Publications Inc.
- Rosas, L. A. A. y Jiménez-Bandala, C. A. (2018). El desempleo y la probabilidad de caer en trampas de pobreza - Unemployment and the Probability of Falling into Poverty Traps: consideraciones para países en vías de desarrollo. *Reis: Revista Española de Investigaciones Sociológicas*, (164):3–20. Publisher: Centro de Investigaciones Sociológicas.
- RTVE (2022). *Dos años del primer caso de COVID-19 en España*. Disponible en: <https://www.rtve.es/noticias/20220131/dos-anos-primer-caso-covid-19-espana/2275660.shtml> [Consultado 15-12-2021].
- Rua Vieites, A. (2021). Modelos cuantitativos: T5. modelización. [Material del aula].
- Sainani, K. L. (2015). Dealing With Missing Data. *PM&R*, 7(9):990–994. Disponible en: <https://onlinelibrary.wiley.com/doi/abs/10.1016/j.pmrj.2015.07.011> [Consultado 15-03-2022].
- Saisana, M. y Saltelli, A. (2011). Rankings and Ratings: Instructions for Use. *Hague Journal on the Rule of Law*, 3(02):247–268.
- Saisana, M., Saltelli, A., y Tarantola, S. (2005). Uncertainty and sensitivity analysis techniques as tools for the quality assessment of composite indicators. *Journal of the Royal Statistical Society: Series A (Statistics in Society)*, 168(2):307–323.

- Saisana, M. y Tarantola, S. (2002). State-of-the-Art Report on Current Methodologies and Practices for Composite Indicator Development. p. 80.
- Saltelli, A. (2007). Composite Indicators between Analysis and Advocacy. *Social Indicators Research*, 81(1):65–77.
- Santeramo, F. G. (2015). On the Composite Indicators for Food Security: Decisions Matter! *Food Reviews International*, 31(1):63–73.
- Santos, V. A. (1996). Métodos multivariantes en bioestadística. Disponible en: https://www.popularlibros.com/libro/metodos-multivariantes-en-bioestadistica_11307 [Consultado 20-02-2022].
- Seguridad Social (2022). *Seguridad Social: Estadísticas*. Disponible en: <https://www.seg-social.es/wps/portal/wss/internet/EstadisticasPresupuestosEstudios/Estadisticas/EST8/EST10/EST305> [Consultado 20-02-2022].
- Sen, A. (1998). *Bienestar, justicia y mercado*. Ediciones Paidós Ibérica, Barcelona.
- Shcherbakov, M. V., Brebels, A., Shcherbakova, N. L., Tyukov, A. P., Janovsky, T. A., y Kamaev, V. A. (2013). A Survey of Forecast Error Measures.
- Soto de la Rosa, H. y Schuschny, A. R. (2009). *Guía metodológica: diseño de indicadores compuestos de desarrollo sostenible*. CEPAL, Santiago de Chile.
- Talukder, B., W. Hipel, K., y W. vanLoon, G. (2017). Developing Composite Indicators for Agricultural Sustainability Assessment: Effect of Normalization and Aggregation Techniques. *Resources*, 6(4):66.
- Van de Kerk, G. y Manuel, A. R. (2008). A comprehensive index for a sustainable society: The SSI – the Sustainable Society Index. *Ecological Economics*, 66(2-3):228–242.
- Van Der Sluijs, J. (2005). Uncertainty as a monster in the science policy interface: four coping strategies. *Water Science and Technology*, 52(6):87–92.
- Van Der Sluijs, J., Craye, M., Funtowicz, S., Klopogge, P., Ravetz, J., y Risbey, J. (2005). Combining Quantitative and Qualitative Measures of Uncertainty in ModelBased Environmental Assessment: The NUSAP System. *Risk Analysis*, 25(2).
- Van Der Sluijs, J. P., Potting, J., Risbey, J., van Vuuren, D., de Vries, B., Beusen, A., Heuberger, P., Corral Quintana, S., Funtowicz, S., Klopogge, P., Nuijten, D., Petersen, A., y Ravetz, J. (2002). Uncertainty assessment of the image/timer b1 co2 emissions scenario, using the nusap method. *Dutch National Research Programme on Global Air Pollution and Climate Change*, 25. Disponible en: <http://www.nusap.net/workshop/report/finalrep.pdf> [Consultado 05-03-2022].
- Willmott, C., Ackleson, S., Davis, R., Feddema, J., Klink, K., Legates, D., O'Donnell, J., y Rowe, C. (1985). Statistics for the Evaluation and Comparison of Models. *Journal of Geophysical Research*.
- Wolf, M. J., Emerson, J. W., Esty, D. C., de Sherbinin, A., y Wendling, Z. A. (2022). *Environmental Performance Index*. Yale Center for Environmental Law & Policy, New Haven.
- Zhou, P., Ang, B., y Poh, K. (2007). A mathematical programming approach to constructing composite indicators. *Ecological Economics*, 62(2):291–297.

PARTE II

ANEXOS



Anexo A

Alineación con los ODS

Los Objetivos de Desarrollo Sostenible (ODS) permiten a la ONU establecer en 2030 un horizonte de mejora de las condiciones de vida a escala global. Hay un total de 17 objetivos, de los cuales se aplican al proyecto los siguientes:

- *Trabajo decente y crecimiento económico (ODS 8)*: El indicador sintético pone el foco sobre el mercado laboral, clave para el crecimiento económico y el bienestar social. La visibilización de los problemas transitorios y estructurales que afrontan el mercado laboral y el tejido productivo español permite mejorar las soluciones propuestas por los diferentes agentes económicos, lo cual sigue la línea establecida por los ODS 8.2 y 8.3: mejora de la productividad y la innovación.
- *Igualdad de género (ODS 5)*: Dada la reducción de las categorías realizada en la sección 3.6, la categoría de *Habilidades y capacitación* atañe principalmente al colectivo femenino, proporcionándole por tanto una importancia superlativa en el sintético final. Las mujeres, dado que se han visto especialmente damnificadas por la pandemia de coronavirus (Marina Gómez García *et al.*, 2021), deben verse especialmente protegidas y representadas por el indicador. Por esta razón, el sintético construido consta de cierta perspectiva de género y persigue empoderar a las mujeres en un mercado laboral que, tradicionalmente, no las ha tratado con justicia. Su sobrerrepresentación en el indicador contribuye por tanto a reducir las desigualdades en el mercado laboral, en tanto que se pone el foco y se visibiliza su situación. Así, se facilita el análisis de iniciativas que aspiren a garantizar la igualdad de las mujeres en el mercado laboral, lo cual está alineado con los objetivos ODS 5.5 y 5.c.

Al igual que los Objetivos de Desarrollo Sostenible aquí mencionados, los objetivos referentes a la pobreza (ODS 1), la reducción de las desigualdades entre países (ODS 10) y el fomento de la industria, la innovación y las infraestructuras (ODS 9) se ven representados de forma tangencial. Al fin y al cabo, el mercado laboral español está también relacionado con estas tres metas, ya que su buen hacer repercute sobre un sistema productivo más justo, digno y competitivamente desarrollado frente a otros países del entorno europeo. De hecho, a través del indicador construido se visibiliza la magnitud real de toda problemática asociada al mercado laboral, de forma que los mandatarios, los empresarios y cualquier otro interesado puedan evaluar el contexto social con la mayor precisión posible. De esta manera, se pueden tomar medidas más acertadas para transformar la sociedad y alcanzar dichos Objetivos de Desarrollo Sostenible.

Anexo B

Tablas de datos

En este anexo aparecen las tablas completas que son mencionadas a lo largo de la memoria del trabajo:

Recolección de datos

Subgrupo	Indicadores	Entrega
Desempleo	ERTEs	Mensual
Desempleo	ERTEs completos	Mensual
Desempleo	ERTEs parciales	Mensual
Desempleo	Impacto en el mercado laboral (ICADE)	Mensual
Desempleo	Paro Registrado	Mensual
Desempleo	Paro Registrado (<25 años)	Mensual
Desempleo	Paro Registrado Agricultura y Pesca	Mensual
Desempleo	Paro Registrado Construcción	Mensual
Desempleo	Paro Registrado CVE	Mensual
Desempleo	Paro Registrado Industria	Mensual
Desempleo	Paro Registrado Mujeres	Mensual
Desempleo	Paro Registrado Servicios	Mensual
Desempleo	Paro Registrado Sin Empleo Previo	Mensual
Desempleo	Paro Registrado Varones	Mensual
Desempleo	Solicitudes Relativas / Vacantes por Mes Finanzas	Mensual
Desempleo	Solicitudes Relativas / Vacantes por Mes Hospitality	Mensual
Desempleo	Solicitudes Relativas / Vacantes por Mes Industria	Mensual
Desempleo	Solicitudes Relativas / Vacantes por Mes Logística	Mensual
Desempleo	Solicitudes Relativas / Vacantes por Mes Retail	Mensual
Desempleo	Solicitudes Relativas / Vacantes por Mes Telemarketing	Mensual
Desempleo	Solicitudes Relativas / Vacantes por Mes Total General	Mensual
Desempleo	Solicitudes Relativas / Vacantes por Mes Transporte	Mensual
Desempleo	Tasa de paro ajustada Eurostat	Mensual
Desempleo	Tasa de Paro Trimestral	Mensual

Empleo	Altas Afiliados	Diario
Empleo	Bajas Afiliados	Diario
Empleo	Contratos Registrados	Mensual
Empleo	Contratos Registrados a Tiempo Parcial	Mensual
Empleo	Contratos Registrados CVE	Mensual
Empleo	Contratos Registrados Indefinidos	Mensual
Empleo	Contratos Registrados Interinidades	Mensual
Empleo	Contratos Registrados Temporales	Mensual
Empleo	Ocupacion crisis 1976	Mensual
Empleo	Ocupacion crisis 1991	Mensual
Empleo	Ocupacion crisis 2007	Mensual
Empleo	Ocupacion crisis 2020	Mensual
Empleo	Productividad por hora efectivamente trabajada	Mensual
Empleo	Productividad por puesto de trabajo a tiempo completo	Mensual
Empleo	Total Afiliados	Diario
Habilidades y Capacitación	Tasa de Paro Analfabetos Ambos sexos	Mensual
Habilidades y Capacitación	Tasa de Paro Analfabetos Hombres	Mensual
Habilidades y Capacitación	Tasa de Paro Analfabetos Mujeres	Mensual
Habilidades y Capacitación	Tasa de Paro Educación Superior Ambos sexos	Mensual
Habilidades y Capacitación	Tasa de Paro Educación Superior Hombres	Mensual
Habilidades y Capacitación	Tasa de Paro Educación Superior Mujeres	Mensual
Habilidades y Capacitación	Tasa de Paro Estudios Primarios Ambos sexos	Mensual
Habilidades y Capacitación	Tasa de Paro Estudios Primarios Hombres	Mensual
Habilidades y Capacitación	Tasa de Paro Estudios Primarios Incompletos Ambos sexos	Mensual
Habilidades y Capacitación	Tasa de Paro Estudios Primarios Incompletos Hombres	Mensual
Habilidades y Capacitación	Tasa de Paro Estudios Primarios Incompletos Mujeres	Mensual
Habilidades y Capacitación	Tasa de Paro Estudios Primarios Mujeres	Mensual
Habilidades y Capacitación	Tasa de Paro FP Ambos sexos	Mensual

Habilidades y Capacitación	Tasa de Paro FP Hombres	Mensual
Habilidades y Capacitación	Tasa de Paro FP Mujeres	Mensual
Habilidades y Capacitación	Tasa de Paro Primera Etapa de Educación Secundaria Ambos sexos	Mensual
Habilidades y Capacitación	Tasa de Paro Primera Etapa de Educación Secundaria Hombres	Mensual
Habilidades y Capacitación	Tasa de Paro Primera Etapa de Educación Secundaria Mujeres	Mensual
Habilidades y Capacitación	Tasa de Paro Segunda Etapa de Educación Secundaria Ambos sexos	Mensual
Habilidades y Capacitación	Tasa de Paro Segunda Etapa de Educación Secundaria Hombres	Mensual
Habilidades y Capacitación	Tasa de Paro Segunda Etapa de Educación Secundaria Mujeres	Mensual
Habilidades y Capacitación	Tasa de Paro Total Ambos sexos	Mensual
Habilidades y Capacitación	Tasa de Paro Total Hombres	Mensual
Habilidades y Capacitación	Tasa de Paro Total Mujeres	Mensual
Protección a los desempleados	Beneficiarios Asistenciales	Mensual
Protección a los desempleados	Prestaciones Contributivas Desempleo Parcial	Mensual
Protección a los desempleados	Prestaciones Contributivas Desempleo Total	Mensual
Salarios y costes laborales	Incremento Salarial Pactado en Negociación Colectiva	Mensual
Salarios y costes laborales	Variación Interanual de los Rendimientos del Trabajo Agricultura y Ganadería	Mensual
Salarios y costes laborales	Variación Interanual de los Rendimientos del Trabajo Comercio	Mensual
Salarios y costes laborales	Variación Interanual de los Rendimientos del Trabajo Construcción	Mensual
Salarios y costes laborales	Variación Interanual de los Rendimientos del Trabajo Hostelería	Mensual
Salarios y costes laborales	Variación Interanual de los Rendimientos del Trabajo Industria	Mensual
Salarios y costes laborales	Variación Interanual de los Rendimientos del Trabajo Servicios	Mensual

Salarios y costes laborales	Variación Interanual de los Rendimientos del Trabajo Total	Mensual
-----------------------------	--	---------

Tabla 24. Lista completa de los indicadores componentes del trabajo junto a su granularidad**Análisis multivariante**

Matriz de componente rotados					
Indicadores	Componentes				
	1	2	3	4	5
<i>ImpactoenelmercadolaboralICADE</i>	0.068	0.143	0.843	-0.167	-0.325
<i>ParoRegistrado</i>	0.985	-0.031	-0.009	0.107	0.117
<i>ParoRegistrado125años</i>	0.915	-0.129	-0.033	0.291	0.054
<i>ParoRegistradoAgriculturayPesca</i>	0.739	-0.534	-0.071	0.033	0.059
<i>ParoRegistradoConstrucción</i>	0.794	0.515	0.042	-0.021	0.100
<i>ParoRegistradoCVE</i>	0.946	-0.158	0.076	-0.072	-0.119
<i>ParoRegistradoIndustria</i>	0.894	0.251	0.277	0.018	0.096
<i>ParoRegistradoMujeres</i>	0.967	-0.176	-0.051	0.103	0.102
<i>ParoRegistradoServicios</i>	0.982	-0.004	0.071	0.071	0.130
<i>ParoRegistradoSinEmpleoPrevio</i>	0.710	-0.487	-0.351	0.227	-0.005
<i>ParoRegistradoVarones</i>	0.972	0.121	0.035	0.107	0.129
<i>SolicitudesRelativasVacantesporMesFinanzas</i>	-0.010	0.933	0.213	0.055	-0.051
<i>SolicitudesRelativasVacantesporMesHospitality</i>	-0.313	0.774	0.120	-0.382	0.230
<i>SolicitudesRelativasVacantesporMesIndustria</i>	0.330	-0.122	0.142	0.036	0.830
<i>SolicitudesRelativasVacantesporMesLogística</i>	0.028	0.153	0.800	0.085	0.216
<i>SolicitudesRelativasVacantesporMesRetail</i>	-0.073	-0.065	0.145	-0.677	-0.373
<i>SolicitudesRelativasVacantesporMesTelemarketing</i>	-0.021	0.657	0.551	-0.281	-0.200
<i>SolicitudesRelativasVacantesporMesTotalGeneral</i>	0.054	0.394	0.804	0.015	0.365
<i>SolicitudesRelativasVacantesporMesTransporte</i>	0.146	0.903	0.213	-0.081	-0.114
<i>TasadeparoajustadaEurostat</i>	0.078	-0.105	0.089	0.890	-0.125
<i>TasadeParoTrimestral</i>	0.252	-0.512	-0.102	0.747	-0.024

Tabla 25. Matriz de componentes rotados obtenidos por PCA para la categoría *Desempleo*

A continuación se presentan las tablas correspondientes al análisis PCA de la sección 3.4:

Varianza total explicada									
Comp.	Autovalores iniciales			Extracción original			Rotación		
	Total	% var.	% acu.	Total	% var.	% acu.	Total	% var.	% acu.

1	8.857	42.179	42.179	8.857	42.179	42.179	8.340	39.714	39.714
2	5.472	26.059	68.237	5.472	26.059	68.237	4.147	19.747	59.461
3	1.837	8.748	76.985	1.837	8.748	76.985	2.687	12.796	72.257
4	1.425	6.788	83.773	1.425	6.788	83.773	2.261	10.768	83.025
5	1.178	5.608	89.381	1.178	5.608	89.381	1.335	6.356	89.381

Tabla 26. Matriz de varianza total explicada por los componentes obtenidos por PCA para *Desempleo*

Matriz de componente rotados				
Indicadores	Componentes			
	1	2	3	4
<i>TasadeParoAnalfabetos</i>	0.155	-0.602	0.191	0.657
<i>TasadeParoAnalfabetosHombres</i>	-0.088	-0.958	-0.136	0.187
<i>TasadeParoAnalfabetosMujeres</i>	0.375	0.223	0.539	0.711
<i>TasadeParoEducaciónSuperior</i>	0.779	-0.072	0.403	0.464
<i>TasadeParoEducaciónSuperiorHombres</i>	0.684	-0.402	0.187	0.518
<i>TasadeParoEducaciónSuperiorMujeres</i>	0.770	0.097	0.484	0.402
<i>TasadeParoEstudiosPrimarios</i>	0.456	-0.063	0.823	0.305
<i>TasadeParoEstudiosPrimariosHombres</i>	0.655	0.223	0.682	0.179
<i>TasadeParoEstudiosPrimariosIncompletos</i>	0.382	-0.121	-0.095	0.910
<i>TasadeParoEstudiosPrimariosIncompletosHombres</i>	0.009	0.399	-0.793	0.291
<i>TasadeParoEstudiosPrimariosIncompletosMujeres</i>	0.376	-0.377	0.372	0.724
<i>TasadeParoEstudiosPrimariosMujeres</i>	0.247	-0.278	0.847	0.360
<i>TasadeParoFP</i>	0.134	0.897	-0.343	-0.209
<i>TasadeParoFPHombres</i>	0.584	0.657	-0.374	-0.226
<i>TasadeParoFPMujeres</i>	-0.195	0.926	-0.278	-0.158
<i>TasadeParoPrimeraEtapadeSecundaria</i>	0.751	0.583	0.132	0.142
<i>TasadeParoPrimeraEtapadeSecundariaHombres</i>	0.909	0.016	0.217	0.104
<i>TasadeParoPrimeraEtapadeSecundariaMujeres</i>	0.474	0.845	-0.011	0.195
<i>TasadeParoSegundaEtapadeSecundaria</i>	0.874	0.396	0.126	0.191
<i>TasadeParoSegundaEtapadeSecundariaHombres</i>	0.933	0.012	-0.025	0.140
<i>TasadeParoSegundaEtapadeSecundariaMujeres</i>	0.731	0.598	0.226	0.239
<i>TasadeParoTotal</i>	0.824	0.276	0.332	0.364
<i>TasadeParoTotalHombres</i>	0.935	-0.109	0.174	0.289
<i>TasadeParoTotalMujeres</i>	0.654	0.503	0.389	0.399

Tabla 27. Matriz de componentes rotados de PCA para la categoría *Habilidades y capacitación*

Varianza total explicada									
Comp.	Autovalores iniciales			Extracción original			Rotación		
	Total	% var.	% acu.	Total	% var.	% acu.	Total	% var.	% acu.
1	12.959	53.996	53.996	12.959	53.996	53.996	8.987	37.447	37.447
2	6.829	28.452	82.448	6.829	28.452	82.448	5.957	24.819	62.267
3	1.993	8.305	90.754	1.993	8.305	90.754	4.172	17.383	79.650
4	1.320	5.500	96.254	1.320	5.500	96.254	3.985	16.604	96.254

Tabla 28. Matriz de varianza explicada por los componentes de PCA para *Habilidades y capacitación*

Matriz de componente rotados		
Indicadores	Componentes	
	1	2
<i>Incremento Salarial Pactado en Negociación Colectiva</i>	-0.366	0.829
<i>Variación Interanual de los Rendimientos del Trabajo Agricultura y Ganadería</i>	-0.060	0.939
<i>Variación Interanual de los Rendimientos del Trabajo Comercio</i>	0.955	-0.233
<i>Variación Interanual de los Rendimientos del Trabajo Construcción</i>	0.972	-0.210
<i>Variación Interanual de los Rendimientos del Trabajo Hostelería</i>	0.936	-0.171
<i>Variación Interanual de los Rendimientos del Trabajo Industria</i>	0.949	-0.220
<i>Variación Interanual de los Rendimientos del Trabajo Servicios</i>	0.978	-0.176
<i>Variación Interanual de los Rendimientos del Trabajo Total</i>	0.980	-0.188

Tabla 29. Matriz de componentes rotados de PCA para la categoría *Salarios y costes laborales*

Varianza total explicada									
Comp.	Autovalores iniciales			Extracción original			Rotación		
	Total	% var.	% acu.	Total	% var.	% acu.	Total	% var.	% acu.
1	6.208	77.602	77.602	6.208	77.602	77.602	5.689	71.112	71.112
2	1.292	16.144	93.747	1.292	16.144	93.747	1.811	22.635	93.747

Tabla 30. Matriz de varianza explicada por los componentes de PCA para *Salarios y costes laborales*

Estadísticas de total de elemento		
	Correlación	Alfa de Cronbach
<i>Impacto en el mercado laboral ICADE</i>	0.047	0.856
<i>Paro Registrado</i>	0.990	0.819
<i>Paro Registrado < 25 años</i>	0.928	0.842
<i>Paro Registrado Agricultura y Pesca</i>	0.766	0.852

<i>ParoRegistradoConstrucción</i>	0.745	0.850
<i>ParoRegistradoCVE</i>	0.908	0.832
<i>ParoRegistradoIndustria</i>	0.875	0.850
<i>ParoRegistradoMujeres</i>	0.982	0.819
<i>ParoRegistradoServicios</i>	0.983	0.813
<i>ParoRegistradoSinEmpleoPrevio</i>	0.732	0.847
<i>ParoRegistradoVarones</i>	0.962	0.821
<i>SolicitudesRelativasVacantesporMesFinanzas</i>	-0.074	0.856
<i>SolicitudesRelativasVacantesporMesHospitality</i>	-0.376	0.856
<i>SolicitudesRelativasVacantesporMesIndustria</i>	0.386	0.856
<i>SolicitudesRelativasVacantesporMesLogística</i>	0.049	0.856
<i>SolicitudesRelativasVacantesporMesRetail</i>	-0.142	0.856
<i>SolicitudesRelativasVacantesporMesTelemarketing</i>	-0.084	0.856
<i>SolicitudesRelativasVacantesporMesTotalGeneral</i>	0.058	0.856
<i>SolicitudesRelativasVacantesporMesTransporte</i>	0.074	0.856
<i>TasadeparoajustadaEurostat</i>	0.131	0.856
<i>TasadeParoTrimestral</i>	0.336	0.856

Tabla 31. Correlación y alfa de Cronbach tras excluir variables de *Desempleo*

Estadísticas de total de elemento		
	Correlación	Alfa de Cronbach
<i>AltasAfiliados</i>	0.330	0.845
<i>BajasAfiliados</i>	0.299	0.847
<i>ContratosRegistrados</i>	0.920	0.793
<i>ContratosRegistradosaTiempoParcial</i>	0.918	0.812
<i>ContratosRegistradosCVE</i>	0.866	0.800
<i>ContratosRegistradosIndefinidos</i>	0.777	0.843
<i>ContratosRegistradosInterinidades</i>	0.874	0.849
<i>ContratosRegistradosTemporales</i>	0.918	0.793
<i>Ocupacioncrisis2020</i>	0.361	0.853
<i>Productividadporhoraefectivamentetrabajada</i>	-0.672	0.853
<i>Productividadporpuestodetrabajoatiempocompleto</i>	0.554	0.853
<i>TotalAfiliados</i>	0.716	0.822

Tabla 32. Correlación y alfa de Cronbach tras excluir variables de *Empleo*

Estadísticas de total de elemento		
	Correlación	Alfa de Cronbach
<i>TasadeParoAnalfabetos</i>	0.634	0.805
<i>TasadeParoAnalfabetosHombres</i>	-0.046	0.866
<i>TasadeParoAnalfabetosMujeres</i>	0.822	0.791
<i>TasadeParoEducaciónSuperior</i>	0.935	0.816
<i>TasadeParoEducaciónSuperiorHombres</i>	0.865	0.820
<i>TasadeParoEducaciónSuperiorMujeres</i>	0.898	0.814
<i>TasadeParoEstudiosPrimarios</i>	0.727	0.811
<i>TasadeParoEstudiosPrimariosHombres</i>	0.687	0.815
<i>TasadeParoEstudiosPrimariosIncompletos</i>	0.831	0.801
<i>TasadeParoEstudiosPrimariosIncompletosHombres</i>	-0.202	0.842
<i>TasadeParoEstudiosPrimariosIncompletosMujeres</i>	0.861	0.799
<i>TasadeParoEstudiosPrimariosMujeres</i>	0.647	0.809
<i>TasadeParoFP</i>	-0.293	0.833
<i>TasadeParoFPHombres</i>	-0.003	0.829
<i>TasadeParoFPMujeres</i>	-0.460	0.839
<i>TasadeParoPrimeraEtapadeEducaciónSecundaria</i>	0.542	0.821
<i>TasadeParoPrimeraEtapadeEducaciónSecundariaHombres</i>	0.722	0.821
<i>TasadeParoPrimeraEtapadeEducaciónSecundariaMujeres</i>	0.294	0.824
<i>TasadeParoSegundaEtapadeEducaciónSecundaria</i>	0.632	0.817
<i>TasadeParoSegundaEtapadeEducaciónSecundariaHombres</i>	0.620	0.819
<i>TasadeParoSegundaEtapadeEducaciónSecundariaMujeres</i>	0.588	0.816
<i>TasadeParoTotal</i>	0.837	0.819
<i>TasadeParoTotalHombres</i>	0.845	0.821
<i>TasadeParoTotalMujeres</i>	0.746	0.817

Tabla 33. Correlación y alfa de Cronbach tras excluir variables de *Habilidades y capacitación*

Estadísticas de total de elemento		
	Correlación	Alfa de Cronbach
<i>Beneficiarios Asistenciales</i>	0.151	0.426
<i>Prestaciones Contributivas Desempleo Parcial</i>	0.861	0.152
<i>Prestaciones Contributivas Desempleo Total</i>	0.485	0.262

Tabla 34. Correlación y alfa de Cronbach tras excluir variables de *Protección a los desempleados*

Estadísticas de total de elemento		
	Correlación	Alfa de Cronbach
<i>Incremento Salarial Pactado en Negociación Colectiva</i>	-0.483	0.726
<i>Variación Interanual de los Rendimientos del Trabajo Agroganadería</i>	-0.246	0.733
<i>Variación Interanual de los Rendimientos del Trabajo Comercio</i>	0.940	0.649
<i>Variación Interanual de los Rendimientos del Trabajo Construcción</i>	0.974	0.627
<i>Variación Interanual de los Rendimientos del Trabajo Hostelería</i>	0.934	0.903
<i>Variación Interanual de los Rendimientos del Trabajo Industria</i>	0.921	0.679
<i>Variación Interanual de los Rendimientos del Trabajo Servicios</i>	0.983	0.639
<i>Variación Interanual de los Rendimientos del Trabajo Total</i>	0.982	0.648

Tabla 35. Correlación y alfa de Cronbach tras excluir variables de *Salarios y costes laborales*

Normalización

En lo referente a la normalización, la transformación logarítmica puede considerarse viable si el coeficiente de curtosis supera 3.5 y la asimetría es mayor que 1. Para averiguarlo, se efectúa para cada categoría un análisis univariado que permita excluir este procedimiento de la fase de normalización:

Categorías	N	%
<i>Desempleo</i>	4	19.0
<i>Empleo</i>	0	0.0
<i>Habilidades y capacitación</i>	2	8.3
<i>Protección a los desempleados</i>	0	0.0
<i>Salarios y costes laborales</i>	0	0.0

Tabla 36. Número de indicadores que cumplen los criterios para la transformación logarítmica

Ponderación y agregación

Una vez realizada la reducción de las categorías en la sección 3.6, la etapa de ponderación a través de PCA da comienzo. Así, se obtienen las siguientes matrices de componentes rotadas:

Matriz de componente rotada			
Indicadores	Componente		
	1	2	3

<i>Impacto en el mercado laboral (ICADE)</i>	0.044	0.758	-0.061
<i>Paro Registrado (<25 años)</i>	0.897	-0.040	0.376
<i>Paro Registrado Agricultura y Pesca</i>	0.766	-0.380	0.223
<i>Paro Registrado Construcción</i>	0.788	0.361	-0.207
<i>Paro Registrado CVE</i>	0.949	-0.027	0.020
<i>Paro Registrado Industria</i>	0.893	0.408	-0.033
<i>Paro Registrado Mujeres</i>	0.968	-0.113	0.193
<i>Paro Registrado Servicios</i>	0.979	0.096	0.116
<i>Paro Registrado Sin Empleo Previo</i>	0.702	-0.545	0.412
<i>Solicitudes Relativas / Vacantes por Mes Total General</i>	0.067	0.902	-0.017
<i>Tasa de paro ajustada Eurostat</i>	0.015	0.104	0.936
<i>Tasa de Paro Trimestral</i>	0.234	-0.308	0.850

Tabla 37. Matriz de componentes rotada con rotación *varimax* de la categoría *Desempleo*

<i>Matriz de componente rotada</i>		
Indicadores	Componente	
	1	2
<i>Tasa de Paro Analfabetos Mujeres</i>	0.903	0.199
<i>Tasa de Paro Educación Superior Mujeres</i>	0.957	0.239
<i>Tasa de Paro Estudios Primarios Mujeres</i>	0.613	0.703
<i>Tasa de Paro Segunda Etapa de Educación Secundaria Mujeres</i>	0.922	-0.344
<i>Tasa de Paro FP Mujeres</i>	0.005	-0.973
<i>Tasa de Paro Educación Superior Ambos sexos</i>	0.899	0.368

Tabla 38. Matriz de componentes rotada por *varimax* de la categoría *Habilidades y capacitación*

<i>Matriz de componente rotada</i>		
Indicadores	Componente	
	1	2
<i>Incremento Salarial Pactado en Negociación Colectiva</i>	-0.366	0.829
<i>Variación Interanual de los Rendimientos del Trabajo Agroganadería</i>	-0.060	0.939
<i>Variación Interanual de los Rendimientos del Trabajo Comercio</i>	0.955	-0.233
<i>Variación Interanual de los Rendimientos del Trabajo Construcción</i>	0.972	-0.210
<i>Variación Interanual de los Rendimientos del Trabajo Hostelería</i>	0.936	-0.171
<i>Variación Interanual de los Rendimientos del Trabajo Industria</i>	0.949	-0.220
<i>Variación Interanual de los Rendimientos del Trabajo Servicios</i>	0.978	-0.176

<i>Variación Interanual de los Rendimientos del Trabajo Total</i>	0.980	-0.188
---	--------------	--------

Tabla 39. Matriz de componentes rotada por *varimax* de la categoría *Salarios y costes laborales*

La categoría restante, *Protección a los desempleados*, no es adecuada para experimentar la técnica PCA, por lo que sus tablas no se consideran pertinentes ni relevantes. A continuación, se presentan los pesos (α_i) de los indicadores de estas tres categorías:

Indicadores	α_i
<i>Impacto en el mercado laboral (ICADE)</i>	0.0554
<i>Paro Registrado (<25 años)</i>	0.0905
<i>Paro Registrado Agricultura y Pesca</i>	0.0745
<i>Paro Registrado Construcción</i>	0.0758
<i>Paro Registrado CVE</i>	0.0861
<i>Paro Registrado Industria</i>	0.0920
<i>Paro Registrado Mujeres</i>	0.0942
<i>Paro Registrado Servicios</i>	0.0937
<i>Paro Registrado Sin Empleo Previo</i>	0.0917
<i>Solicitudes Relativas / Vacantes por Mes Total General</i>	0.0781
<i>Tasa de paro ajustada Eurostat</i>	0.0847
<i>Tasa de Paro Trimestral</i>	0.0833

Tabla 40. Pesos de los diferentes indicadores para la subagregación aritmética en *Desempleo*

Indicadores	α_i
<i>Tasa de Paro Analfabetos Mujeres</i>	0.1540
<i>Tasa de Paro Educación Superior Ambos sexos</i>	0.1699
<i>Tasa de Paro Educación Superior Mujeres</i>	0.1749
<i>Tasa de Paro Estudios Primarios Mujeres</i>	0.1564
<i>Tasa de Paro FP Mujeres</i>	0.1705
<i>Tasa de Paro Segunda Etapa de Educación Secundaria Mujeres</i>	0.1743

Tabla 41. Pesos de los indicadores para la subagregación aritmética en *Habilidades y Capacitación*

Indicadores	α_i
<i>Incremento Salarial Pactado en Negociación Colectiva</i>	0.1094
<i>Variación Interanual de los Rendimientos del Trabajo Agricultura y Ganadería</i>	0.1181
<i>Variación Interanual de los Rendimientos del Trabajo Comercio</i>	0.1288
<i>Variación Interanual de los Rendimientos del Trabajo Construcción</i>	0.1318

<i>Variación Interanual de los Rendimientos del Trabajo Hostelería</i>	0.1206
<i>Variación Interanual de los Rendimientos del Trabajo Industria</i>	0.1266
<i>Variación Interanual de los Rendimientos del Trabajo Servicios</i>	0.1318
<i>Variación Interanual de los Rendimientos del Trabajo Total</i>	0.1329

Tabla 42. Pesos de los indicadores para la subagregación aritmética en *Salarios y costes laborales*

De esta manera, a partir de la agregación aritmética, se subagregan los indicadores en torno a sus categorías. El proceso de construcción del sintético final llega a término en el documento principal de la memoria del proyecto.

Figuras

Recolección de datos

Figura 23. Instantánea de la tabla de estructuración y selección de datos en el subgrupo *Desempleo*

Figura 24. Estructuración y selección de datos en el subgrupo *Habilidades y capacitación*

Fecha	Beneficiarios Asistenciales	Prestaciones Contributivas Desempleo Parcial	Prestaciones Contributivas Desempleo Total
Jul-21	985574	412	899773
Jun-21	1024856	464	840466
May-21	1069636	488	908902
Apr-21	1078546	454	998086
Mar-21	1083193	477	1052787
Feb-21	1103832	476	1161888
Jan-21	1133776	582	1148021
Dec-20	1093782	486	1109271
Nov-20	1044347	5	1238390
Oct-20	1020903	545	1433338
Sep-20	997281	137235	1310006
Aug-20	1020261	183125	1640863
Jul-20	1043767	246098	1854408
Jun-20	1087033	315966	2481334
May-20	1112379	346071	3393321
Apr-20	1106757	208251	3236640
Mar-20	1038462	325	973726
Feb-20	1011915	234	892065
Jan-20	1008548	223	939443

Figura 25. Estructuración y selección de datos en el subgrupo *Protección a los desempleados*

Fecha	Incremento Salarial Pactado en Negociación Colectiva	Variación Interanual de los Rendimientos del Trabajo Agricultura y Ganadería	Variación Interanual de los Rendimientos del Trabajo Comercio	Variación Interanual de los Rendimientos del Trabajo Construcción	Variación Interanual de los Rendimientos del Trabajo Hostelería	Variación Interanual de los Rendimientos del Trabajo Industria	Variación Interanual de los Rendimientos del Trabajo Servicios	Variación Interanual de los Rendimientos del Trabajo Total
Jul-21	1.5	0.885	4.171	4.723	19.106	3.324	5.696	5.175
Jun-21	1.6	1.823	8.598	7.711	51.926	3.898	8.609	7.149
May-21	1.6	0.790	16.564	15.494	83.868	7.691	15.476	13.240
Apr-21	1.6	-1.892	20.515	22.972	98.329	11.531	15.649	15.138
Mar-21	1.6	-0.183	2.657	-0.328	-34.175	1.233	-0.969	-0.468
Feb-21	1.5	-0.989	-0.345	-3.932	-45.000	-0.980	-3.743	-3.226
Jan-21	1.4	-0.374	0.962	-1.748	-42.763	0.059	-3.317	-2.259
Dec-20	1.9	0.019	1.781	-3.332	-46.140	0.521	-3.356	-2.475
Nov-20	1.9	4.221	0.549	-5.248	-54.643	0.477	-3.959	-3.079
Oct-20	1.9	3.525	0.902	-4.684	-55.012	0.285	-4.677	-3.563
Sep-20	1.9	5.065	0.230	-5.290	-47.593	-0.717	-5.260	-4.184
Aug-20	1.9	3.305	-0.113	-5.812	-44.644	-1.639	-5.364	-4.494
Jul-20	1.9	1.225	-1.728	-7.160	-50.361	-2.421	-7.040	-5.997
Jun-20	2	4.216	-4.941	-9.121	-62.840	-3.340	-10.862	-8.677
May-20	1.96	-1.912	-12.572	-16.097	-72.259	-8.502	-16.311	-14.298
Apr-20	1.96	6.517	-14.909	-20.899	-75.761	-11.571	-15.797	-14.934
Mar-20	1.96	6.415	-0.974	-0.713	-24.203	-0.356	-2.068	-1.489
Feb-20	1.97	6.659	3.330	4.083	2.145	1.877	3.232	2.807
Jan-20	1.98	5.647	3.636	3.460	3.319	1.867	3.327	2.882

Figura 26. Estructuración y selección de datos en el subgrupo *Salarios y costes laborales*

Imputación de información ausente

Las dos primeras figuras representan la técnica basada en NRMSE de elección del método de imputación:

II. ANEXOS C. FIGURAS

Fecha	Altas Afiliados	Bajas Afiliados	Contratos Registrados	Contratos Registrados a Tiempo Parcial	Contratos Registrados CVE	Contratos Registrados Indefinidos	Contratos Registrados Interinidades	Contratos Registrados Temporales	Ocupacion crisis 2020	Productividad por hora efectivamente trabajada	Productividad por puesto de trabajo a tiempo completo	Total Afiliados
Jun-21	109591.2	109613.6	1798047.000	657589.000	1558935.000	172866.000	68264.000	1625181.000	98.522	-6.986	0.740	19500277.4
May-21	99487.5	87156.4	1545308.000	518699.000	1513358.000	156148.000	63912.000	1389160.000	97.745	-4.876	-0.308	19267221.0
Apr-21	80052.5	75257.7	1332886.529	406034.119	1425082.122	140121.569	59966.602	1192765.000	96.824	-3.140	-1.817	19016930.3
Mar-21	78293.5	81870.9	1404107.000	453042.000	1475289.574	207191.000	62590.000	1196916.000	96.193	-0.655	-2.404	18920901.9
Feb-21	74704.6	76258.2	1212284.000	339408.000	1459780.957	132431.000	63754.000	1079853.000	96.423	-1.403	-2.873	18850111.7
Jan-21	93691.4	97517.7	1302429.000	358990.000	1451485.209	124191.000	66557.000	1178238.000	96.652	-2.151	-3.343	18829480.2
Dec-20	87074.5	88322.7	1355147.000	401285.000	1381692.350	111822.000	66636.000	1243325.000	96.882	-2.899	-3.812	19048433.3
Nov-20	76003.9	84163.5	1449810.000	419228.000	1512946.466	128189.000	71882.000	1321621.000	96.374	-2.816	-3.676	19022001.6
Oct-20	95941.8	89648.6	1541866.028	517602.123	1381229.532	142828.052	66109.193	1399038.000	95.270	-1.336	-3.051	18977067.9
Sep-20	143821.0	120663.1	1632484.000	542413.000	1408332.000	163209.000	74131.000	1469275.000	95.358	-2.652	-3.403	18862712.8
Aug-20	71455.7	76186.9	1118663.000	353195.000	1419365.000	96275.000	53965.000	1022388.000	94.636	-0.576	-3.732	18792376.1
Jul-20	87149.2	79818.9	1536122.000	531513.000	1312696.000	141105.000	67421.000	1395017.000	93.913	1.500	-4.062	18785554.3
Jun-20	74456.0	80719.9	1159602.000	362772.000	1003079.000	114393.000	47439.000	1045209.000	93.190	3.576	-4.392	18624336.7
May-20	55230.2	45519.9	850617.000	195131.000	832856.000	76692.000	32902.000	773925.000	94.983	2.438	-4.105	18556128.9
Apr-20	47802.3	49994.4	671249.000	134515.000	717819.000	59042.000	38898.000	614107.000	96.776	1.300	-3.819	18458666.8
Mar-20	80012.3	118714.1	1256510.000	397997.000	1327930.000	145393.000	64815.000	1111117.000	98.570	0.162	-3.533	19006759.6

Figura 27. Instantánea de la categoría *Empleo* tras la imputación EM en IBM SPSS

Fecha	Altas Afiliados	Bajas Afiliados	Contratos Registrados	Contratos Registrados a Tiempo Parcial	Contratos Registrados CVE	Contratos Registrados Indefinidos	Contratos Registrados Interinidades	Contratos Registrados Temporales	Ocupacion crisis 2020	Productividad por hora efectivamente trabajada	Productividad por puesto de trabajo a tiempo completo	Total Afiliados
Jun-21	109591.2	109613.6	1798047.000	657589.000	1558935.000	172866.000	68264.000	1625181.000	98.522	-6.986	0.740	19500277.4
May-21	99487.5	87156.4	1545308.000	518699.000	1513358.000	156148.000	63912.000	1389160.000	97.745	-4.876	-0.308	19267221.0
Apr-21	83946.0	75257.7	1421595.93	432314.663	1369981.210	133755.324	58621.709	1192765.000	97.483	-2.946	-1.627	18939678.4
Mar-21	78293.5	81870.9	1404107.000	453042.000	1475289.574	207191.000	62590.000	1196916.000	96.193	-0.655	-2.404	18920901.9
Feb-21	74704.6	76258.2	1212284.000	339408.000	1459780.957	132431.000	63754.000	1079853.000	96.423	-1.403	-2.873	18850111.7
Jan-21	93691.4	97517.7	1302429.000	358990.000	1451485.209	124191.000	66557.000	1178238.000	96.652	-2.151	-3.343	18829480.2
Dec-20	87074.5	88322.7	1355147.000	401285.000	1381692.350	111822.000	66636.000	1243325.000	96.882	-2.899	-3.812	19048433.3
Nov-20	76003.9	84163.5	1449810.000	419228.000	1512946.466	128189.000	71882.000	1321621.000	96.374	-2.816	-3.676	19022001.6
Oct-20	87077.3	89648.6	1103138.69	465697.302	1644368.184	165432.152	64167.958	1399038.000	95.437	1.422	-3.775	19015198.6
Sep-20	143821.0	120663.1	1632484.000	542413.000	1408332.000	163209.000	74131.000	1469275.000	95.358	-2.652	-3.403	18862712.8
Aug-20	71455.7	76186.9	1118663.000	353195.000	1419365.000	96275.000	53965.000	1022388.000	94.636	-0.576	-3.732	18792376.1
Jul-20	87149.2	79818.9	1536122.000	531513.000	1312696.000	141105.000	67421.000	1395017.000	93.913	1.500	-4.062	18785554.3
Jun-20	74456.0	80719.9	1159602.000	362772.000	1003079.000	114393.000	47439.000	1045209.000	93.190	3.576	-4.392	18624336.7
May-20	55230.2	45519.9	850617.000	195131.000	832856.000	76692.000	32902.000	773925.000	94.983	2.438	-4.105	18556128.9
Apr-20	47802.3	49994.4	673149.000	134515.000	717819.000	59042.000	38898.000	614107.000	96.776	1.300	-3.819	18458666.8
Mar-20	80012.3	118714.1	1256510.000	397997.000	1327930.000	145393.000	64815.000	1111117.000	98.570	0.162	-3.533	19006759.6

Figura 28. Instantánea de la categoría *Empleo* tras la imputación por regresión

A continuación se muestran las matrices de datos completos ya imputados:

Indicadores presentados en el mismo orden que en tablas anteriores																					
Fecha	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	21
Aug-21	0.195	3333915	245291	161678	271609	3414779	271855	1972216	2345613	283160	1361699	9.743	6.584	0.200	2.713	1.686	2.745	3.685	2.191	0.152	16.118
Jul-21	0.247	3416498	262411	175177	270470	3563982	272981	2017719	2391837	306033	1398779	11.695	6.383	2.955	2.673	11.083	6.214	2.864	0.076	0.155	15.408
Jun-21	0.192	3614339	299337	184057	280624	3755738	286139	2122610	2525495	338024	1491729	8.123	8.742	1.093	2.989	49.000	1.730	2.887	0.250	0.151	15.100
May-21	0.197	3781250	322894	182175	292387	3858342	298837	2201471	2656712	351139	1579779	13.864	8.919	5.444	2.913	12.455	1.054	3.015	0.305	0.154	15.393
Apr-21	0.195	3910828	358884	191330	300536	3887771	308240	2263125	2750039	360483	1647503	18.664	2.565	4.587	4.168	24.000	1.444	2.979	0.342	0.156	15.887
Mar-21	0.196	3949640	357793	193952	304483	3835302	313367	2278099	2782231	355607	1671541	37.239	6.025	5.020	5.558	2.400	2.101	4.229	1.461	0.154	15.980
Feb-21	0.198	4008789	366403	191584	312168	3866579	317042	2304779	2835917	352078	1704001	13.091	9.642	12.338	6.509	1.056	1.921	5.656	0.549	0.157	16.030
Jan-21	0.201	3964353	357123	185410	317284	3846356	316413	2273375	2799040	346206	1609078	9.654	1.844	10.085	4.370	0.000	2.095	6.032	0.123	0.161	16.080
Dec-20	0.204	3888137	362997	182138	318155	3863896	315290	2225121	2720951	351603	1663016	15.894	1.621	0.593	2.899	0.000	1.464	2.645	1.875	0.162	16.130
Nov-20	0.207	3851312	365719	183449	299659	3787509	305405	2222254	2712817	349982	1629058	22.255	0.896	1.429	3.726	21.000	1.835	3.289	0.000	0.164	16.173
Apr-19	0.205	3826043	361986	188073	298177	3785751	305707	2203285	2687858	346228	1622758	29.476	4.388	3.614	5.055	9.000	2.251	4.437	0.545	0.163	16.217
May-19	0.219	3776485	348925	177839	298542	3831765	304921	2181794	2657234	337949	1594691	41.667	2.616	2.034	6.071	5.431	1.444	4.653	3.462	0.164	16.260
Jun-19	0.224	3802814	321968	187342	306224	3898283	313016	2197913	2670601	325631	1604901	30.615	3.354	2.090	10.513	5.556	1.811	5.370	0.309	0.159	15.940
Jul-19	0.239	3773034	321364	200595	298241	3932654	310035	2177586	2650385	317788	1595448	9.605	3.069	3.789	6.784	1.771	2.852	5.076	1.308	0.154	15.620
Aug-19	0.275	3802814	335763	189487	304797	3951592	319479	2197913	2734948	314172	1604901	11.566	7.105	3.703	8.646	36.913	5.533	6.699	1.167	0.153	15.300
Sep-19	0.293	3857776	326574	164145	320724	3936418	327249	2191678	2762267	283391	1666098	35.186	6.988	3.356	6.219	18.722	11.108	5.824	8.000	0.152	15.000
Oct-19	0.263	3831200	318820	163440	344440	3813957	327510	2151800	2721480	274330	1679400	137.667	23.469	3.146	7.602	17.538	13.899	7.409	50.000	0.155	14.700
Nov-19	0.160	3548312	287560	159420	319386	3443787	300679	2019370	2502355	266472	1538942	49.483	14.000	2.221	4.385	14.889	2.021	4.139	3.285	0.150	14.400
Dec-19	0.246	3246407	261448	152900	259835	3132750	275485	1896072	2226339	261488	1349975	47.841	11.907	4.585	6.586	22.652	5.313	6.202	2.299	0.160	15.417
Jan-20	0.208	3253583	254240	150045	264654	3151793	277744	1896873	2305824	255586	1356980	41.208	13.108	2.383	8.514	13.247	3.956	5.242	2.367	0.159	15.137

Figura 29. Subgrupo *Desempleo* tras la aplicación de la imputación regresiva

Indicadores presentados en el mismo orden que en tablas anteriores																								
Fecha	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	21	22	23	24
Jun-21	30.790	16.340	52.090	9.520	8.300	10.600	26.350	24.310	28.010	29.340	25.780	29.670	17.280	14.450	20.420	20.330	16.600	26.060	17.360	14.460	20.490	15.260	13.390	17.360
May-21	32.773	17.773	51.993	9.737	8.400	10.923	26.383	24.387	28.490	28.427	28.703	29.573	17.357	14.780	20.227	20.733	17.097	26.377	17.543	14.507	20.807	15.500	13.617	17.61

El resto de categorías ya se encontraban completas, por lo que se recomienda visualizar las figuras 25 y 26.

Análisis multivariante

Las matrices de correlación copan la primera parte del análisis multivariante:

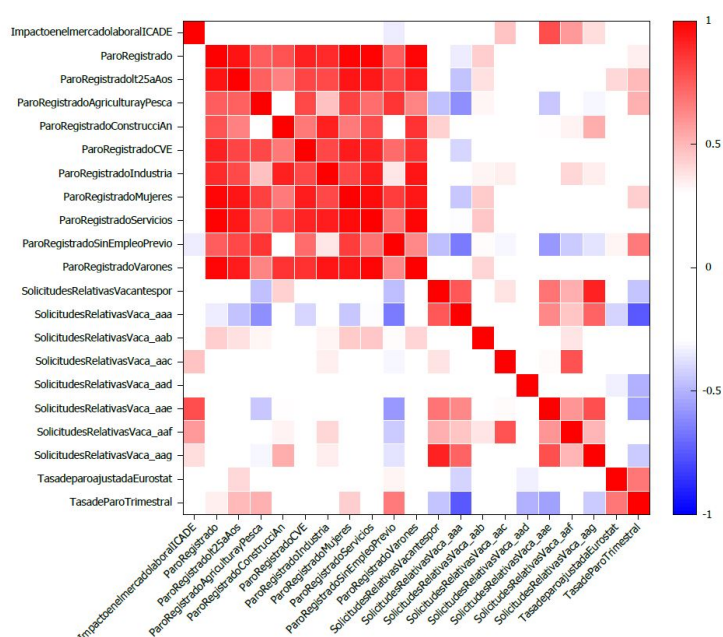


Figura 31. Matriz de correlaciones de los indicadores de la categoría *Desempleo*

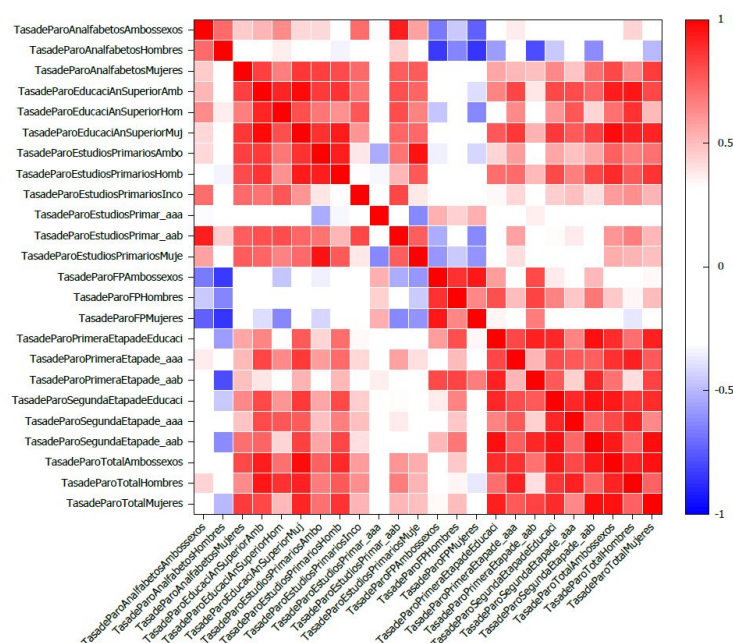


Figura 32. Matriz de correlaciones de los indicadores de *Habilidades y capacitación*

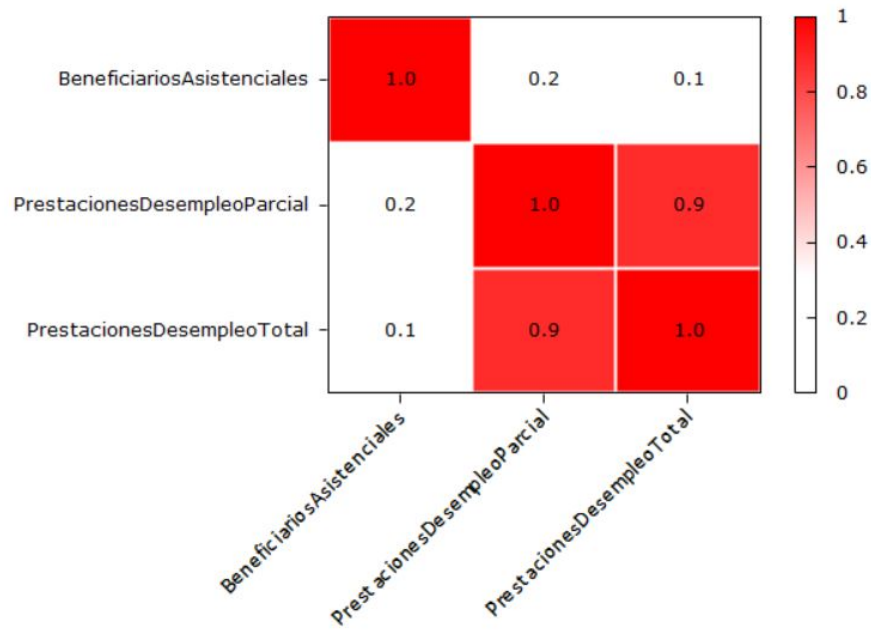


Figura 33. Matriz de correlaciones de los indicadores de *Protección a los desempleados*

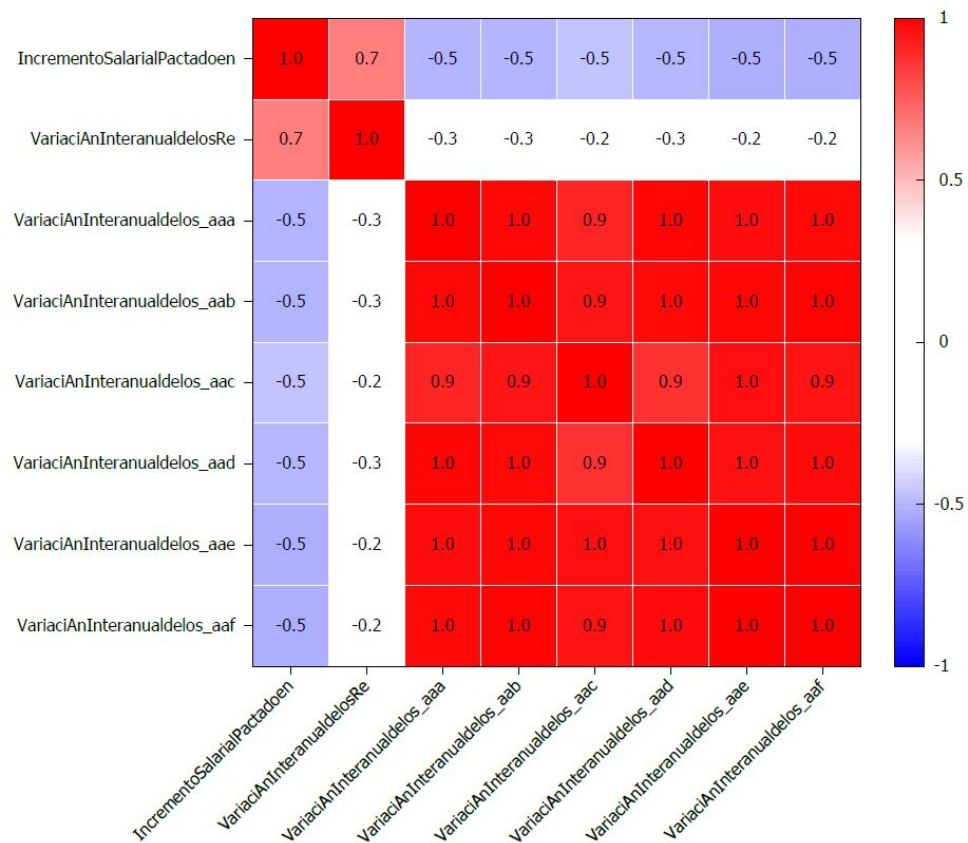


Figura 34. Matriz de correlaciones de los indicadores de *Salarios y costes laborales*

En segundo lugar, se presentan los análisis *cluster*:

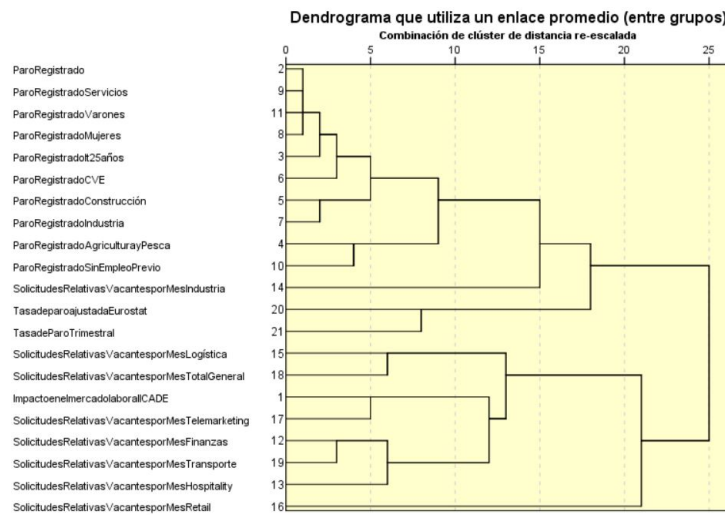


Figura 35. Dendrograma del análisis *cluster* de la categoría *Desempleo*

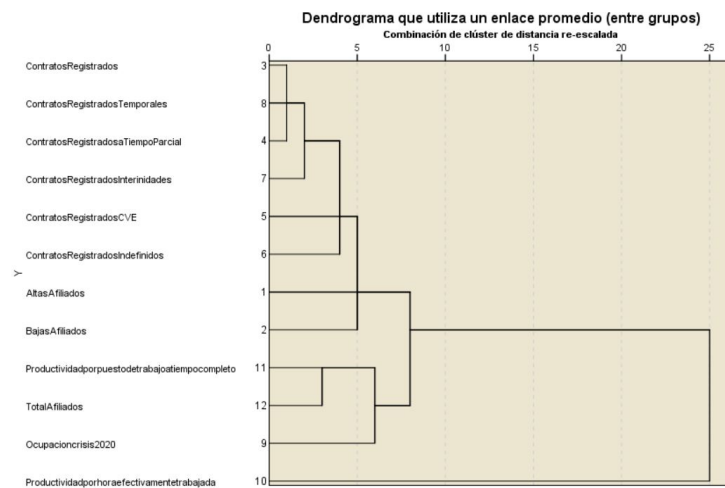


Figura 36. Dendrograma del análisis *cluster* de la categoría *Empleo*

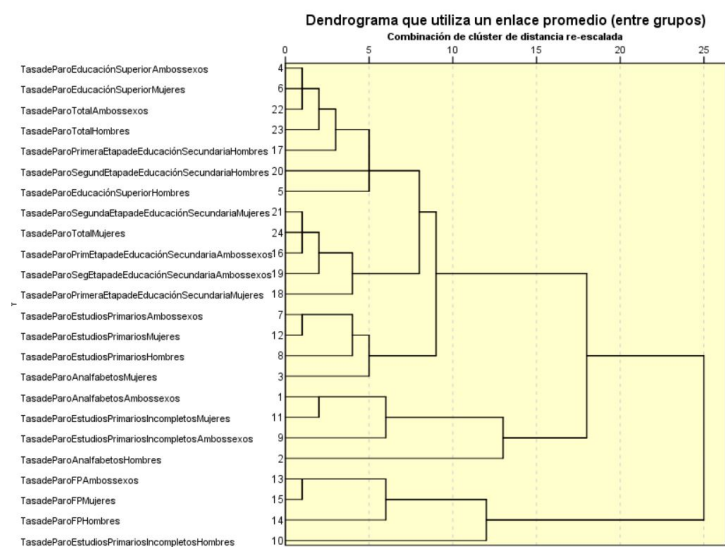


Figura 37. Dendrograma del análisis *cluster* de la categoría *Habilidades y capacitación*



Figura 38. Dendrograma del análisis *cluster* de la categoría *Protección a los desempleados*



Figura 39. Dendrograma del análisis *cluster* de la categoría *Salarios y costes laborales*

Normalización

Llegados a este punto, tiene lugar la normalización Mín-Máx:

Indicadores presentados en el mismo orden que en tablas anteriores																					
Fecha	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	21
Aug-21	0.741	0.885	1.000	0.770	0.861	0.656	1.000	0.814	0.909	0.737	0.967	0.987	0.748	1.000	0.995	0.966	0.868	0.782	0.956	0.824	0.077
Jul-21	0.345	0.777	0.859	0.503	0.874	0.473	0.980	0.702	0.823	0.519	0.862	0.972	0.757	0.773	1.000	0.774	0.598	0.954	0.998	0.641	0.458
Jun-21	0.764	0.517	0.554	0.327	0.754	0.239	0.743	0.446	0.575	0.214	0.600	1.000	0.652	0.926	0.960	0.000	0.947	0.949	0.995	0.929	0.624
May-21	0.725	0.298	0.359	0.364	0.615	0.114	0.515	0.253	0.332	0.089	0.351	0.956	0.645	0.568	0.969	0.746	1.000	0.923	0.994	0.714	0.466
Apr-21	0.738	0.129	0.087	0.183	0.519	0.078	0.346	0.102	0.159	0.000	0.160	0.919	0.926	0.639	0.809	0.510	0.970	0.930	0.993	0.571	0.308
Mar-21	0.734	0.078	0.071	0.131	0.472	0.142	0.254	0.065	0.099	0.046	0.092	0.775	0.773	0.603	0.632	0.951	0.918	0.667	0.971	0.714	0.151
Feb-21	0.714	0.000	0.000	0.178	0.381	0.104	0.188	0.000	0.000	0.080	0.000	0.962	0.613	0.000	0.511	0.978	0.932	0.368	0.989	0.500	0.124
Jan-21	0.693	0.058	0.077	0.300	0.321	0.129	0.199	0.077	0.068	0.136	0.037	0.988	0.958	0.186	0.784	1.000	0.919	0.289	0.998	0.214	0.097
Dec-20	0.669	0.158	0.028	0.365	0.311	0.107	0.220	0.195	0.213	0.085	0.116	0.940	0.968	0.968	0.971	1.000	0.968	1.000	0.963	0.143	0.070
Nov-20	0.649	0.206	0.006	0.339	0.529	0.200	0.397	0.202	0.228	0.100	0.212	0.891	1.000	0.899	0.866	0.571	0.939	0.865	1.000	0.000	0.047
Apr-19	0.663	0.240	0.036	0.248	0.547	0.203	0.392	0.248	0.274	0.136	0.230	0.835	0.845	0.719	0.696	0.816	0.907	0.624	0.989	0.071	0.023
May-19	0.560	0.305	0.170	0.450	0.542	0.146	0.406	0.301	0.331	0.215	0.309	0.741	0.924	0.849	0.567	0.889	0.970	0.579	0.931	0.000	0.000
Jun-19	0.522	0.270	0.307	0.262	0.452	0.065	0.260	0.261	0.306	0.332	0.280	0.826	0.891	0.844	0.000	0.887	0.941	0.428	0.994	0.357	0.172
Jul-19	0.404	0.309	0.372	0.000	0.546	0.023	0.314	0.311	0.344	0.445	0.307	0.989	0.904	0.704	0.476	0.964	0.860	0.490	0.974	0.714	0.344
Aug-19	0.136	0.270	0.253	0.220	0.469	0.000	0.144	0.261	0.187	0.441	0.280	0.973	0.725	0.711	0.238	0.206	0.651	0.149	0.977	0.786	0.516
Sep-19	0.000	0.198	0.329	0.721	0.280	0.019	0.005	0.277	0.136	0.735	0.107	0.791	0.730	0.740	0.548	0.618	0.217	0.333	0.840	0.857	0.677
Oct-19	0.223	0.233	0.393	0.735	0.000	0.168	0.000	0.374	0.212	0.821	0.070	0.000	0.000	0.757	0.371	0.642	0.000	0.000	0.643	0.839	
Nov-19	1.000	0.604	0.651	0.815	0.296	0.620	0.482	0.698	0.618	0.896	0.494	0.681	0.419	0.833	0.782	0.696	0.925	0.687	0.934	1.000	1.000
Dec-19	0.353	1.000	0.867	0.944	1.000	1.000	0.935	1.000	1.000	0.944	1.000	0.693	0.512	0.639	0.501	0.538	0.668	0.253	0.954	0.255	0.453
Jan-20	0.636	0.990	0.926	1.000	0.943	0.977	0.894	0.998	0.982	1.000	0.980	0.745	0.459	0.820	0.255	0.730	0.774	0.455	0.953	0.381	0.604

Figura 40. Instantánea del subgrupo *Desempleo* tras la normalización y reorientación de los datos

II. ANEXOS ▯ C. FIGURAS

Indicadores presentados en el mismo orden que en tablas anteriores																								
Fecha	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	21	22	23	24
Jun-21	1.000	1.000	0.645	0.750	0.824	0.716	1.000	0.638	0.889	0.000	1.000	0.944	0.125	0.384	0.000	0.494	0.856	0.257	0.208	0.269	0.247	0.541	0.625	0.479
May-21	0.842	0.885	0.654	0.650	0.765	0.590	0.990	0.616	0.817	0.132	0.865	0.963	0.083	0.256	0.071	0.329	0.571	0.171	0.152	0.254	0.165	0.411	0.483	0.360
Apr-21	0.684	0.770	0.663	0.549	0.706	0.464	0.980	0.593	0.746	0.263	0.731	0.981	0.042	0.128	0.142	0.165	0.285	0.086	0.096	0.238	0.082	0.281	0.342	0.240
Mar-21	0.526	0.654	0.672	0.449	0.647	0.339	0.970	0.571	0.674	0.395	0.596	1.000	0.000	0.000	0.212	0.000	0.000	0.000	0.040	0.223	0.000	0.151	0.200	0.121
Feb-21	0.543	0.682	0.547	0.366	0.580	0.246	0.647	0.380	0.760	0.597	0.563	0.667	0.243	0.177	0.376	0.076	0.025	0.117	0.027	0.177	0.003	0.124	0.179	0.090
Jan-21	0.561	0.711	0.422	0.282	0.514	0.154	0.323	0.190	0.847	0.798	0.529	0.333	0.486	0.354	0.540	0.152	0.050	0.234	0.013	0.131	0.005	0.097	0.158	0.059
Dec-20	0.578	0.739	0.297	0.199	0.447	0.062	0.000	0.000	0.934	1.000	0.495	0.000	0.728	0.531	0.703	0.229	0.075	0.351	0.000	0.085	0.008	0.070	0.138	0.028
Nov-20	0.386	0.550	0.198	0.133	0.298	0.042	0.091	0.106	0.623	0.781	0.330	0.048	0.804	0.572	0.763	0.252	0.113	0.351	0.013	0.083	0.032	0.047	0.092	0.019
Oct-20	0.193	0.361	0.099	0.066	0.149	0.021	0.182	0.212	0.311	0.563	0.165	0.096	0.880	0.614	0.823	0.275	0.151	0.351	0.027	0.081	0.056	0.023	0.046	0.009
Sep-20	0.000	0.172	0.000	0.000	0.000	0.000	0.273	0.318	0.000	0.344	0.000	0.144	0.957	0.655	0.883	0.298	0.190	0.351	0.040	0.079	0.081	0.000	0.000	0.000
Aug-20	0.160	0.115	0.333	0.171	0.131	0.193	0.513	0.462	0.277	0.395	0.217	0.377	0.933	0.585	0.922	0.390	0.199	0.505	0.115	0.052	0.237	0.168	0.054	0.259
Jul-20	0.321	0.057	0.667	0.343	0.263	0.387	0.754	0.606	0.554	0.445	0.435	0.611	0.909	0.516	0.961	0.483	0.209	0.659	0.191	0.026	0.393	0.335	0.108	0.518
Jun-20	0.481	0.000	1.000	0.514	0.394	0.580	0.994	0.750	0.831	0.496	0.652	0.845	0.886	0.446	1.000	0.576	0.218	0.814	0.266	0.000	0.549	0.503	0.163	0.777
May-20	0.462	0.058	0.947	0.676	0.596	0.720	0.991	0.833	0.887	0.596	0.657	0.766	0.924	0.630	0.878	0.717	0.479	0.876	0.511	0.333	0.700	0.668	0.442	0.851
Apr-20	0.444	0.117	0.894	0.838	0.798	0.860	0.988	0.917	0.944	0.695	0.663	0.687	0.962	0.815	0.756	0.859	0.739	0.938	0.755	0.667	0.850	0.834	0.721	0.926
Mar-20	0.425	0.175	0.840	1.000	1.000	1.000	0.985	1.000	1.000	0.795	0.668	0.608	1.000	1.000	0.634	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
Feb-20	0.438	0.324	0.576	0.234	0.375	0.156	0.194	0.149	0.889	0.800	0.470	0.143	0.888	0.541	0.941	0.314	0.200	0.604	0.148	0.195	0.229	0.186	0.103	0.297
Jan-20	0.723	0.419	0.585	0.271	0.172	0.331	0.495	0.359	0.741	0.479	0.692	0.463	0.837	0.531	0.750	0.540	0.410	0.672	0.053	0.017	0.285	0.286	0.193	0.478

Figura 41. *Habilidades y capacitación* tras la normalización y reorientación de los datos

Fecha	Beneficiarios Asistenciales	Prestaciones Contributivas Desempleo Parcial	Prestaciones Contributivas Desempleo Total
Jul-21	0.00000	0.00118	0.02323
Jun-21	0.26506	0.00133	0.00000
May-21	0.56721	0.00140	0.02681
Apr-21	0.62733	0.00130	0.06174
Mar-21	0.65869	0.00136	0.08317
Feb-21	0.79795	0.00136	0.12591
Jan-21	1.00000	0.00167	0.12047
Dec-20	0.73014	0.00139	0.10530
Nov-20	0.39657	0.00000	0.15587
Oct-20	0.23838	0.00156	0.23224
Sep-20	0.07899	0.39654	0.18393
Aug-20	0.23405	0.52915	0.31353
Jul-20	0.39266	0.71112	0.39718
Jun-20	0.68460	0.91301	0.64276
May-20	0.85562	1.00000	1.00000
Apr-20	0.81769	0.60175	0.93863
Mar-20	0.35686	0.00092	0.05220
Feb-20	0.17774	0.00066	0.02021
Jan-20	0.15502	0.00063	0.03877

Figura 42. *Protección a los desempleados* tras la normalización y reorientación de los datos

Fecha	Incremento Salarial Pactado en Negociación Colectiva	Variación Interanual de los Rendimientos del Trabajo Agricultura y Ganadería	Variación Interanual de los Rendimientos del Trabajo Comercio	Variación Interanual de los Rendimientos del Trabajo Construcción	Variación Interanual de los Rendimientos del Trabajo Hostelería	Variación Interanual de los Rendimientos del Trabajo Industria	Variación Interanual de los Rendimientos del Trabajo Servicios	Variación Interanual de los Rendimientos del Trabajo Total
Jul-21	0.1667	0.3263	0.5386	0.5840	0.5449	0.6448	0.6886	0.6687
Jun-21	0.3333	0.4357	0.6636	0.6521	0.7335	0.6696	0.7797	0.7343
May-21	0.3333	0.3153	0.8884	0.8295	0.9169	0.8338	0.9946	0.9369
Apr-21	0.3333	0.0023	1.0000	1.0000	1.0000	1.0000	1.0000	1.0000
Mar-21	0.3333	0.2018	0.4959	0.4689	0.2389	0.5542	0.4801	0.4810
Feb-21	0.1667	0.1077	0.4111	0.3868	0.1767	0.4584	0.3932	0.3893
Jan-21	0.0000	0.1794	0.4480	0.4365	0.1895	0.5034	0.4066	0.4215
Dec-20	0.8333	0.2252	0.4711	0.4004	0.1701	0.5234	0.4054	0.4143
Nov-20	0.8333	0.7155	0.4364	0.3568	0.1213	0.5215	0.3865	0.3942
Oct-20	0.8333	0.6344	0.4463	0.3696	0.1192	0.5132	0.3640	0.3781
Sep-20	0.8333	0.8140	0.4274	0.3558	0.1618	0.4698	0.3458	0.3575
Aug-20	0.8333	0.6086	0.4177	0.3439	0.1787	0.4299	0.3425	0.3472
Jul-20	0.8333	0.3660	0.3721	0.3132	0.1459	0.3961	0.2901	0.2972
Jun-20	1.0000	0.7150	0.2814	0.2685	0.0742	0.3563	0.1705	0.2081
May-20	0.9333	0.0000	0.0660	0.1095	0.0201	0.1328	0.0000	0.0212
Apr-20	0.9333	0.9834	0.0000	0.0000	0.0000	0.0000	0.0161	0.0000
Mar-20	0.9333	0.9715	0.3934	0.4601	0.2962	0.4854	0.4457	0.4471
Feb-20	0.9500	1.0000	0.5149	0.5694	0.4475	0.5821	0.6115	0.5900
Jan-20	0.9667	0.8819	0.5235	0.5552	0.4542	0.5817	0.6145	0.5924

Figura 43. *Instantánea de Salarios y costes laborales* tras la normalización y reorientación de los datos

Ponderación y agregación

A continuación se muestran las gráficas individualizadas de cada una de las categorías:

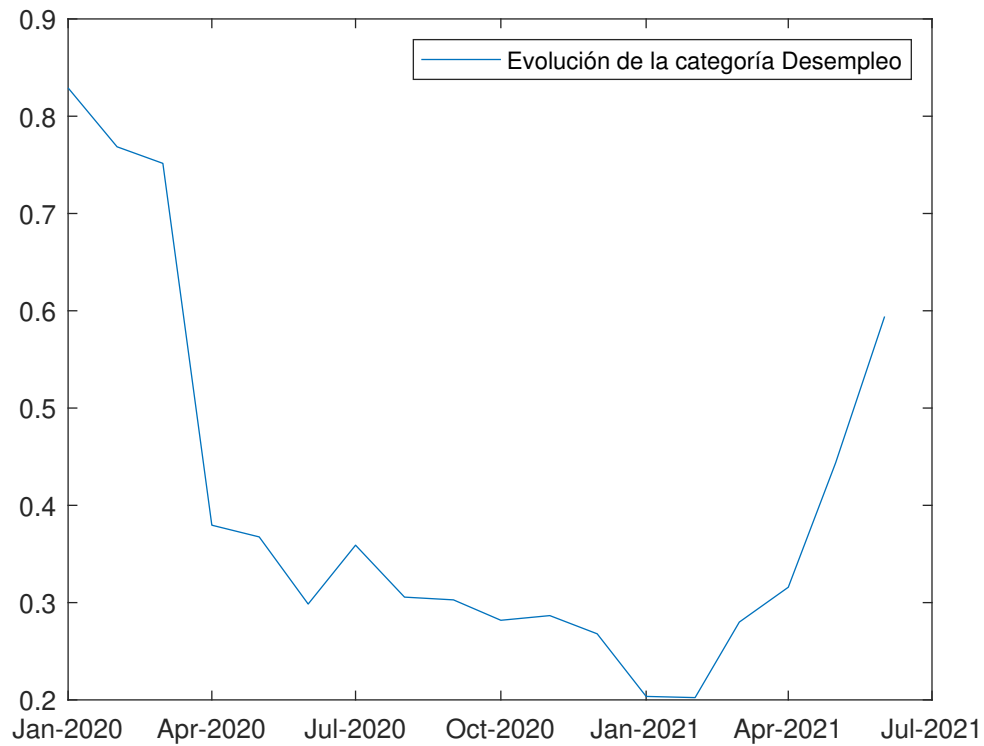


Figura 44. Evolución individualizada de la categoría *Desempleo*

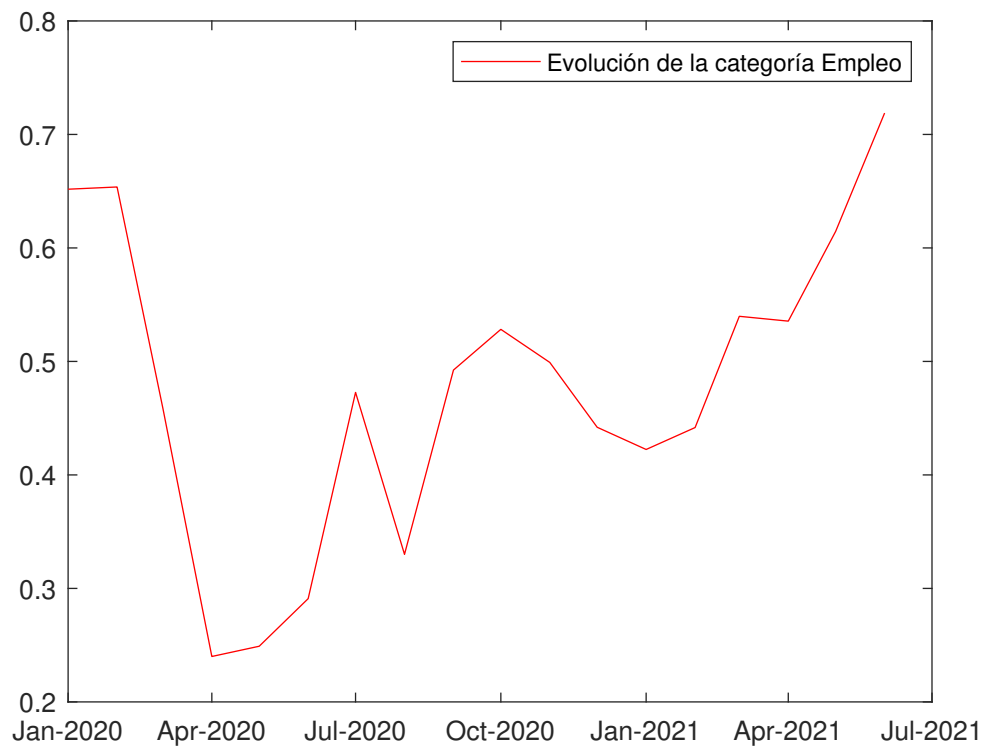


Figura 45. Evolución individualizada de la categoría *Empleo*

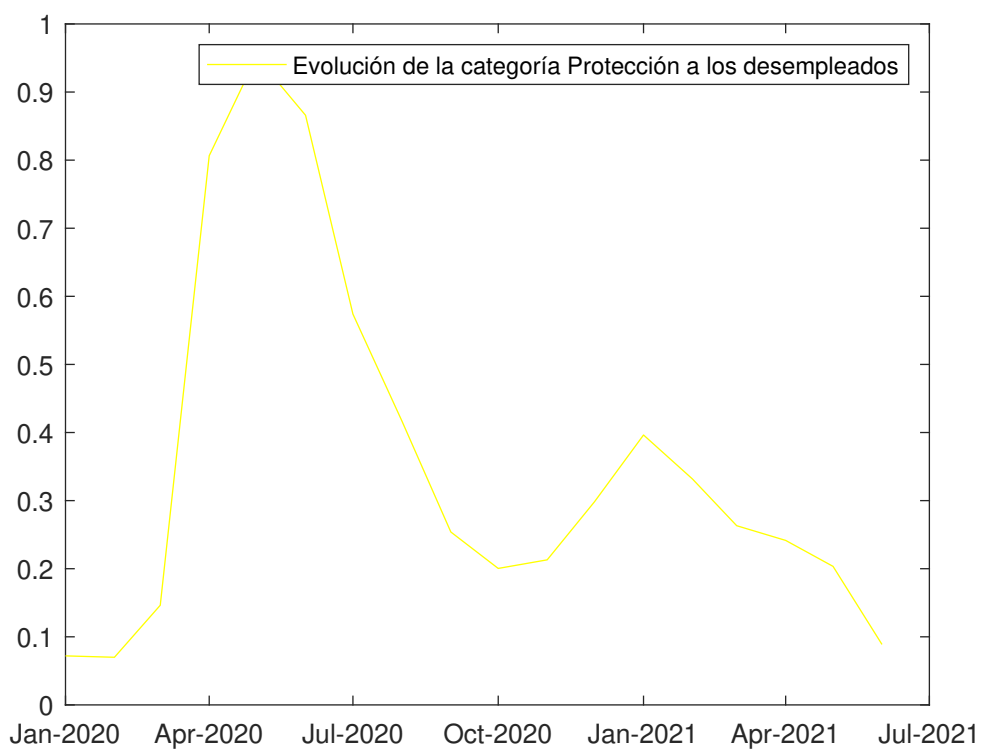


Figura 46. Evolución individualizada de la categoría *Protección a los desempleados*

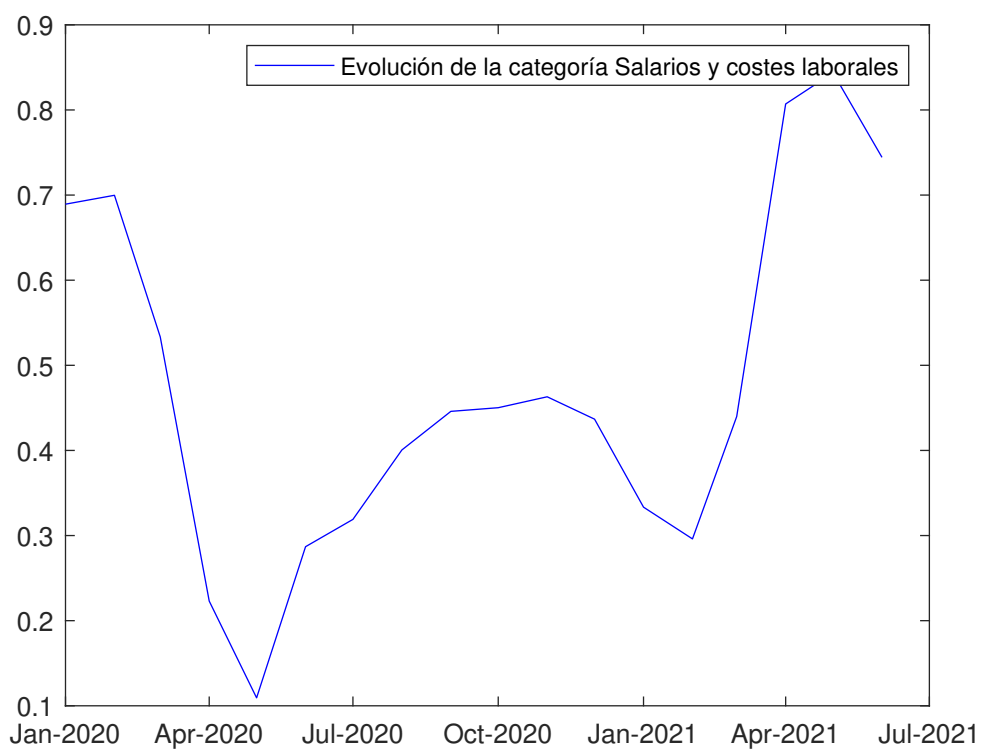


Figura 47. Evolución individualizada de la categoría *Salarios y costes laborales*

Análisis de robustez

Las próximas gráficas forman parte del análisis de robustez, aunque son presentadas en la memoria del proyecto en forma de matriz de gráficas:

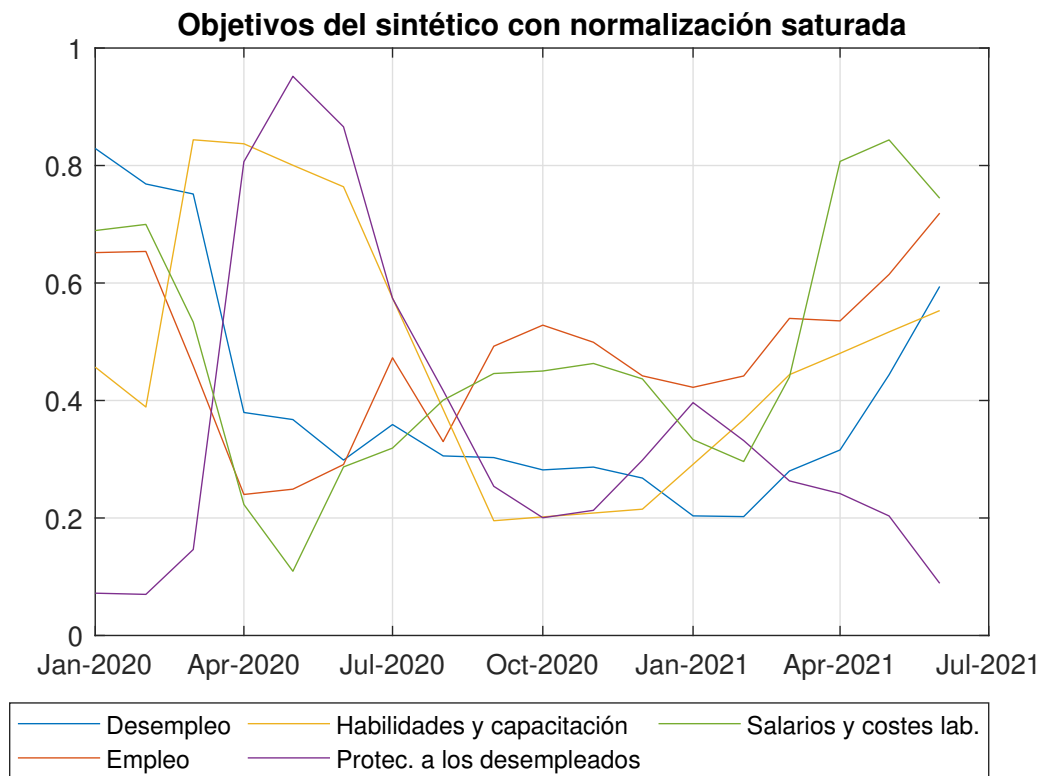


Figura 48. Evolución de los objetivos del sintético con saturación en la normalización

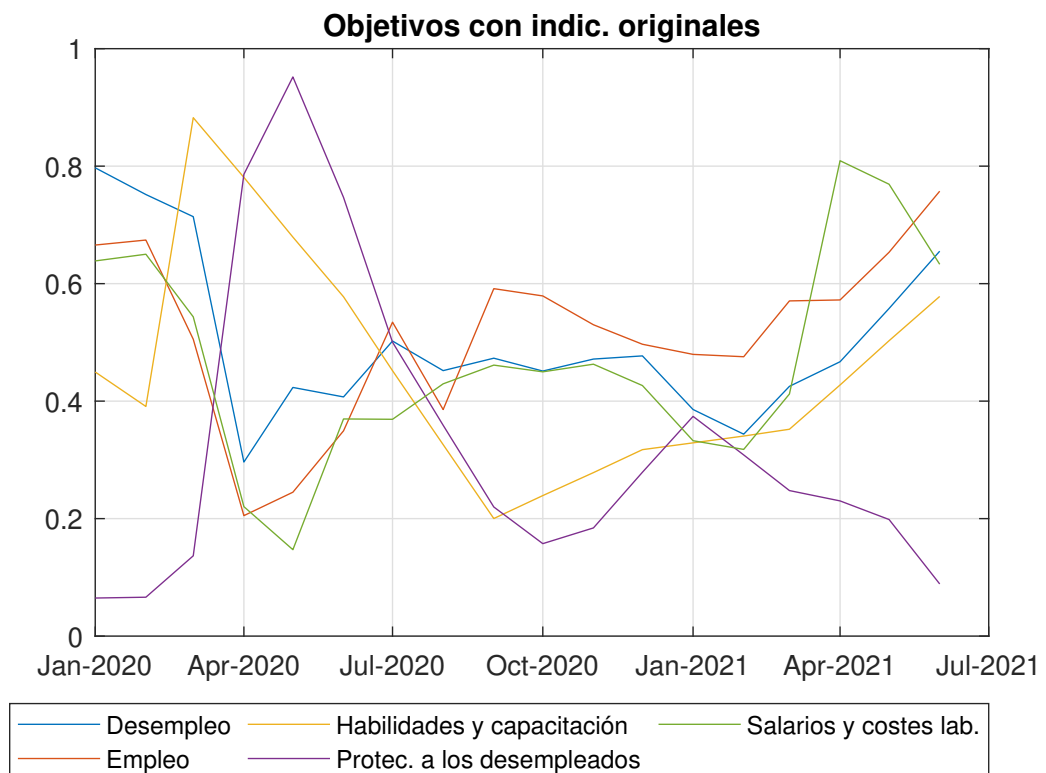


Figura 49. Evolución de los objetivos del sintético a partir de los indicadores originales

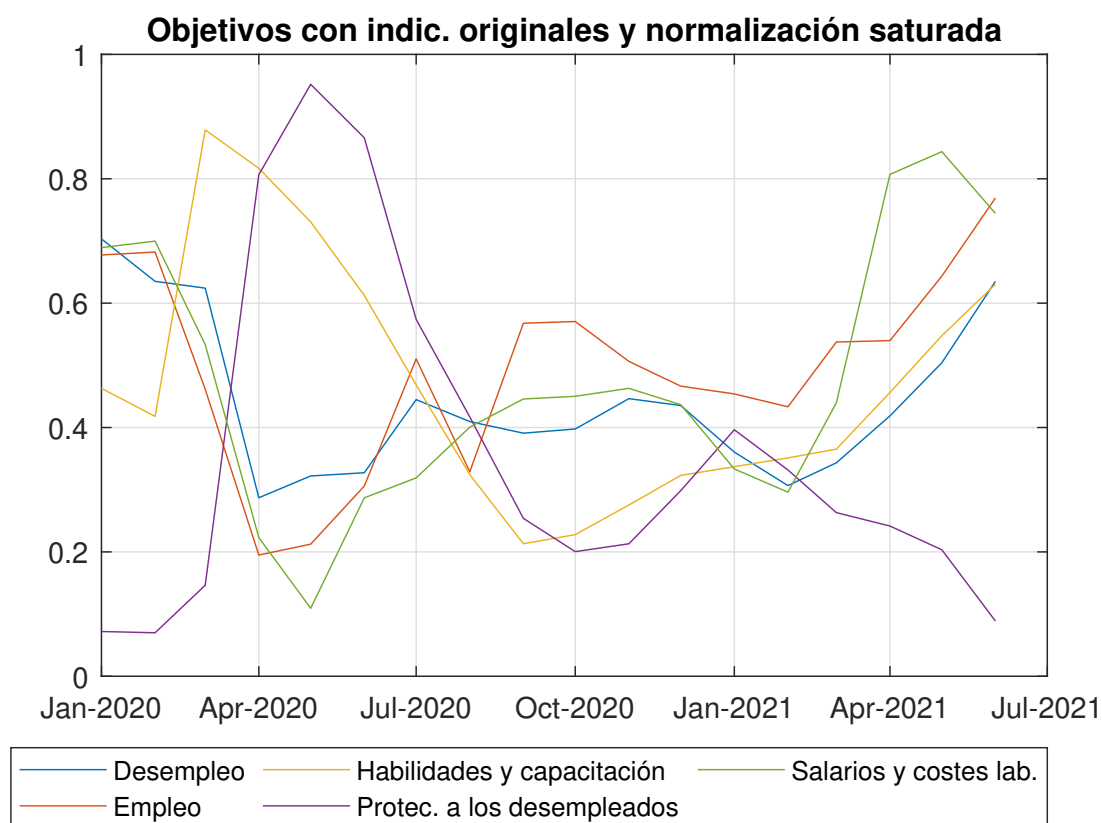


Figura 50. Evolución de los objetivos del sintético original con saturación en la normalización

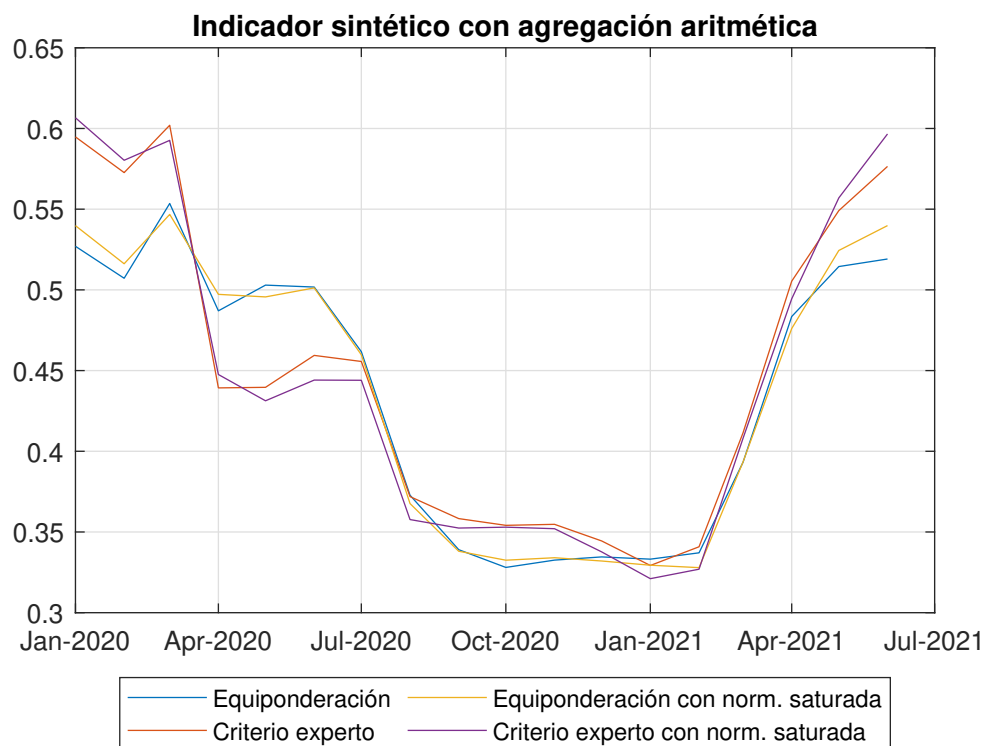


Figura 51. Indicador sintético con todos los indicadores originales tras la agregación final aritmética

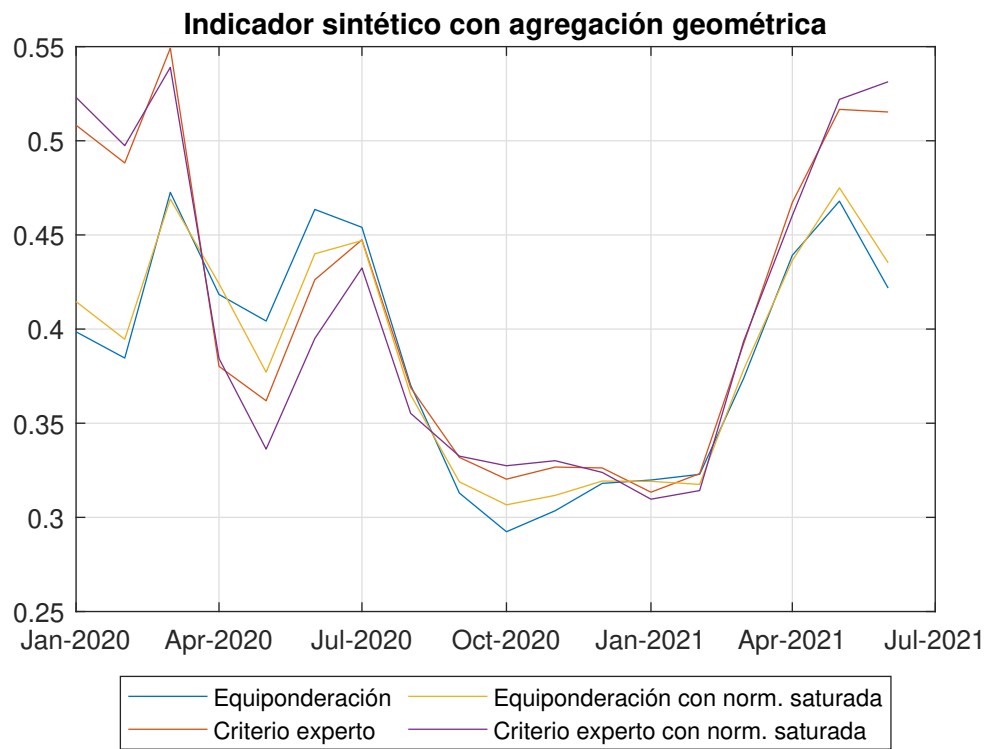


Figura 52. Evolución del indicador sintético tras la agregación geométrica

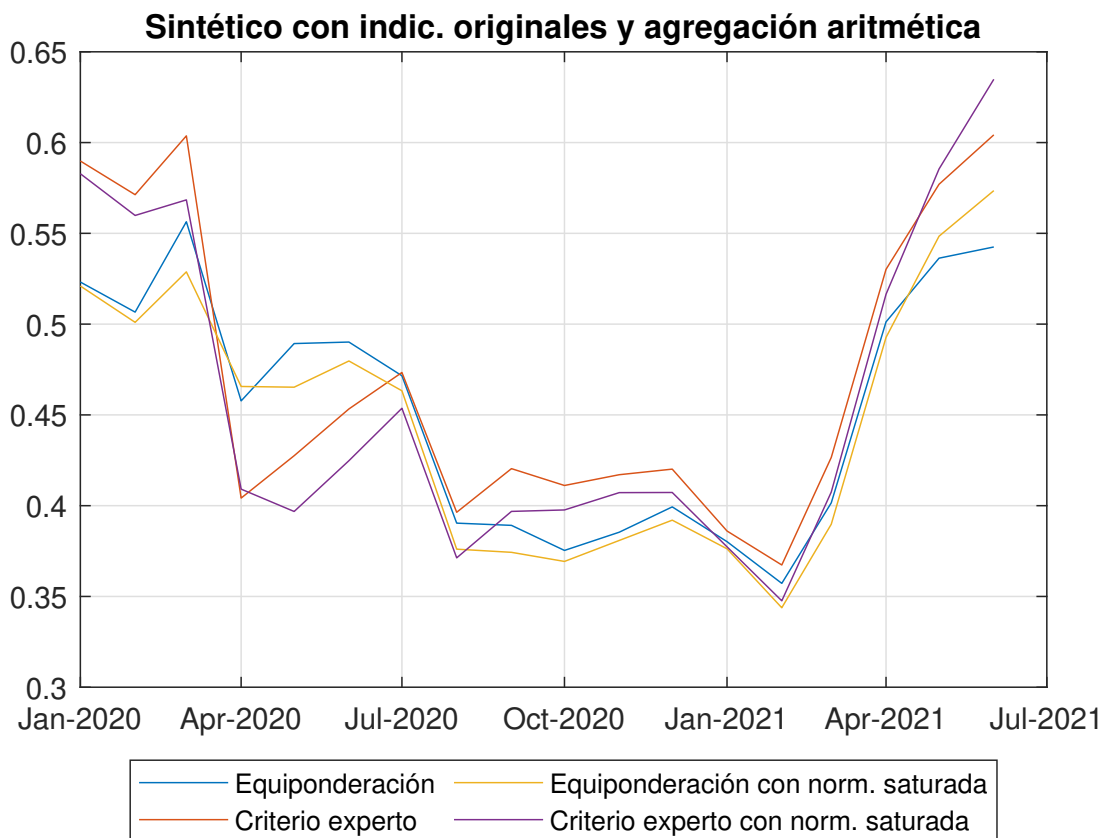


Figura 53. Evolución del sintético con todos los indicadores originales tras la agregación aritmética

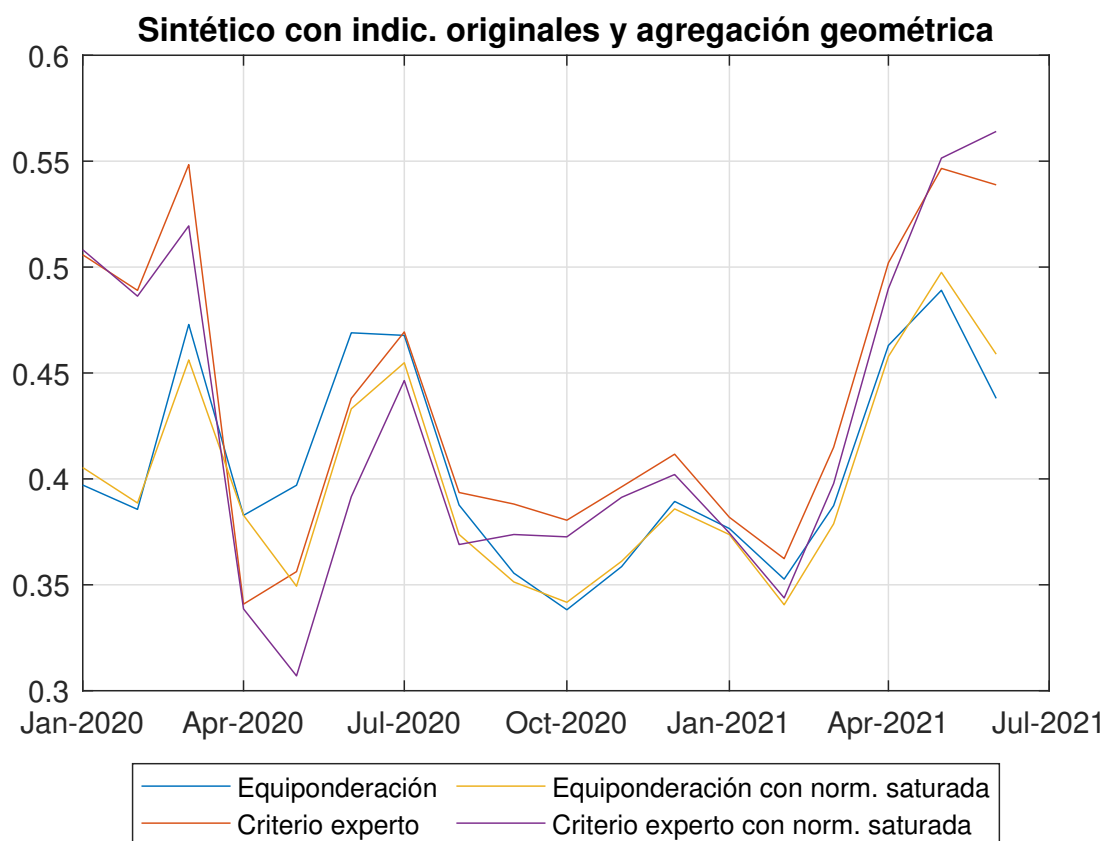


Figura 54. Evolución del sintético con los indicadores originales tras la agregación geométrica

Anexo D

Código fuente

EL siguiente código de MATLAB presenta la selección de datos y la creación de archivos para la imputación de información ausente propias de la elaboración del indicador "Salud del Mercado de Trabajo":

```
clear

%% Cargar archivos
trabajo='Mercado de Trabajo.xlsx';
[num_trabajo,txt_trabajo,lista_trabajo] = xlsread(trabajo);
longitud_trabajo=size(lista_trabajo,1);
clear num_trabajo
clear txt_trabajo

%% Adaptar formato de las fechas y eliminar filas excluidas
j=2;
descarte_1= 'Solicitudes / Usuarios Activos';
descarte_2= 'Vacantes Disponibles / Usuarios Activos';
descarte_3= 'Jornada Media Pactada';
for i=2:longitud_trabajo
    lista_trabajo{j,6}=datetime(lista_trabajo{j,6},'InputFormat','M/d/yyyy'); %Depende del
        idioma en que esté el PC
    if contains(lista_trabajo{j,3}, descarte_1) || contains(lista_trabajo{j,3}, descarte_2)
        || contains(lista_trabajo{j,3}, descarte_3)
        lista_trabajo(j,:)=[];

    else if isequaln('España', lista_trabajo{j,4}) && isequaln('Total', lista_trabajo{j,5})

        j=j+1;
        else
            lista_trabajo(j,:)=[];
        end
    end
end

longitud_trabajo=size(lista_trabajo,1);

%% Tabla con granularidad
j=1;
for i=1:(longitud_trabajo-1)
    if isequaln(lista_trabajo{i+1,3},lista_trabajo{i,3})
    else
        resumen_trabajo{j,1}=lista_trabajo{i+1,2};
        resumen_trabajo{j,2}=lista_trabajo{i+1,3};
        j=j+1;
    end
end

j=1;
z=0;
granularidad=duration(datetime('16-04-2000','InputFormat','dd-MM-yyyy') -
    datetime('15-01-2000','InputFormat','dd-MM-yyyy'));
for i=2:(longitud_trabajo-1)
    if isequaln(lista_trabajo{i+1,3},lista_trabajo{i,3}) && i~=(longitud_trabajo-1)
        diferencia=duration(lista_trabajo{i,6}-lista_trabajo{i+1,6});
```

```

if diferencia<duration((datetime('16-04-2000','InputFormat','dd-MM-yyyy') -
    datetime('10-04-2000','InputFormat','dd-MM-yyyy')))
    z=z+1;
end
if z==0 && diferencia<granularidad
    granularidad=diferencia;
else if z>0 && diferencia<granularidad
    granularidad=diferencia;
end
end
else if granularidad>duration(datetime('16-04-2000','InputFormat','dd-MM-yyyy')-
    datetime('25-03-2000','InputFormat','dd-MM-yyyy'))
    resumen_trabajo{j,3}='Mensual';
    j=j+1;
    z=0;
else if granularidad>duration(datetime('16-04-2000','InputFormat','dd-MM-yyyy') -
    datetime('06-04-2000','InputFormat','dd-MM-yyyy'))
    resumen_trabajo{j,3}='Bisemanal';
    j=j+1;
    z=0;
else if granularidad>duration(datetime('16-04-2000','InputFormat','dd-MM-yyyy')
    - datetime('11-04-2000','InputFormat','dd-MM-yyyy'))
    resumen_trabajo{j,3}='Semanal';
    j=j+1;
    z=0;
else
    resumen_trabajo{j,3}='Diario';
    j=j+1;
    z=0;
end
end
end
end
if isequaln(lista_trabajo{i+1,3},lista_trabajo{i,3})
else
    granularidad= duration(datetime('16-04-2000','InputFormat','dd-MM-yyyy') -
        datetime('15-03-2000','InputFormat','dd-MM-yyyy'));
end
end

clear granularidad
clear diferencia
clear descarte_1 descarte_2 descarte_3

%% Vector con las posiciones de inicio y fin de las categorías
j=1;
k=1;
posicion_cat(1,j)=2;
posicion_kpi(1,k)=2;
for i=2:(longitud_trabajo-1)
    if isequaln(lista_trabajo{i+1,2},lista_trabajo{i,2})
    else
        posicion_cat(2,j)=i;
        j=j+1;
        posicion_cat(1,j)=i+1;
    end
    if isequaln(lista_trabajo{i+1,3},lista_trabajo{i,3})
    else
        posicion_kpi(2,k)=i;
        k=k+1;
        posicion_kpi(1,k)=i+1;
    end
end
posicion_cat(2,j)=i+1;
posicion_kpi(2,k)=i+1;

%% Señalar outliers
j=0;
z=1;
flag=true;
flag_covid=true;
for i=1:longitud_trabajo-1
    if (isequaln(lista_trabajo{i,3},lista_trabajo{i+1,3}) || i==1)
        j=j+1;
        vector(j,1)=lista_trabajo{i+1,7};
    end
end

```

```

end
if isequaln(lista_trabajo{i,3},lista_trabajo{i+1,3}) && i<(longitud_trabajo-1)
else if j>1
    quartiles=quantile(vector,[0.25 0.75]);
    rango_iqr=iqr(vector);
    umbral_down=quartiles(1)-3*rango_iqr;
    umbral_up=quartiles(2)+3*rango_iqr;
    for k=1:size(vector,1)
        if (vector(k,1)<umbral_down || vector(k,1)>umbral_up) && flag_covid==true
            resumen_trabajo{z,4}='ATÍPICO';
            flag=false;
            fecha=lista_trabajo{posicion_kpi(1,z)+k-1,6};
            resumen_trabajo{z,5}=fecha;
            if isbetween(fecha, datetime('14-03-2020','InputFormat','dd-MM-yyyy'),
                datetime('15-06-2021','InputFormat','dd-MM-yyyy'))
                resumen_trabajo{z,4}='COVID Atípico';
                flag_covid=false;
            end
            resumen_trabajo{z,5}=char(fecha);
        else if flag==true
            resumen_trabajo{z,4}='Correcto';
        end
    end
    end
    flag=true;
    flag_covid=true;
    clear vector
    j=1;
    vector(j,1)=lista_trabajo{i+1,7};
end
end
if isequaln(lista_trabajo{i,3},lista_trabajo{i+1,3}) || i==1
else
    z=z+1;
end
end

clear flag
clear flag_covid
clear rango_iqr
clear umbral_down
clear umbral_up
clear quartiles
clear vector
clear fecha

%% Pasar lista trabajo a mensual
for i=2:longitud_trabajo
    lista_trabajo{i,6}.Format='MMM-yyyy';
end

%% Leer número de categorías
j=1;
for i=2:(longitud_trabajo-1)
    if isequaln(lista_trabajo{i+1,2},lista_trabajo{i,2})
    else
        j=j+1;
    end
    i=i+1;
end
n_cat=j;

%% Gran bucle selección de datos e imputación

% Fecha de inicio del COVID en España (31 de enero, primer caso)
inicio_covid_date=datetime('Jan-2020','InputFormat','MMM-yyyy');
inicio_covid_date.Format='MMM-yyyy';

for j=1:n_cat
    v_tiempo_trabajo(1,j)=lista_trabajo{posicion_cat(1,j),6};
    v_tiempo_trabajo(2,j)=lista_trabajo{posicion_cat(1,j),6};

    for i=posicion_cat(1,j):posicion_cat(2,j)

```

```

        if v_tiempo_trabajo(1,j)>lista_trabajo{i,6}
            v_tiempo_trabajo(1,j)=lista_trabajo{i,6};
            if v_tiempo_trabajo(1,j)<inicio_covid_date
                v_tiempo_trabajo(1,j)=inicio_covid_date;
            end
        end
        if v_tiempo_trabajo(2,j)<lista_trabajo{i,6}
            v_tiempo_trabajo(2,j)=lista_trabajo{i,6};
        end
    end

    vt{1,j}=transpose(v_tiempo_trabajo(2,j):calmonths(-1):v_tiempo_trabajo(1,j));
    vt{1,j}.Format='MMM-yyyy';

    dimensiones_trabajo(1,j)=size(vt{1,j},1)+1;%Incluye título 'Fecha'

    n=0;
    for i=posicion_cat(1,j):posicion_cat(2,j)
        if isequaln(lista_trabajo{i,3},lista_trabajo{i-1,3})
            else
                n=n+1;
            end
        end
    end
    dimensiones_trabajo(2,j)=n+1;%+1 por los títulos

    tablas_trabajo{1,j}=cell(dimensiones_trabajo(1,j),dimensiones_trabajo(2,j));

    for i=1:dimensiones_trabajo(1,j)
        if i==1
            tablas_trabajo{1,j}(1,1)=cellstr('Fecha');
        else
            tablas_trabajo{1,j}(i,1)=cellstr(char(vt{1,j}(i-1,1)));
        end
    end

    k=1;
    for i=posicion_cat(1,j):posicion_cat(2,j)
        if isequaln(lista_trabajo{i,3},lista_trabajo{i-1,3})
            else
                k=k+1;
                tablas_trabajo{1,j}{1,k}=lista_trabajo{i,3};
            end
        end
    end

    for i=2:longitud_trabajo
        lista_trabajo(i,6)=cellstr(char(lista_trabajo{i,6}));
    end

    for i=posicion_cat(1,j):posicion_cat(2,j)
        if isequaln(lista_trabajo{i,3},lista_trabajo{i-1,3})
            else
                for z=1:size(resumen_trabajo,1)
                    if isequaln(resumen_trabajo{z,2},lista_trabajo{i,3}) &&
                        isequaln(resumen_trabajo{z,3},'Diario')
                        a=0;
                        n=0;
                        for t=i:posicion_cat(2,j)
                            if isequaln(lista_trabajo{t,3},lista_trabajo{t-1,3}) || t==i
                                else
                                    break
                                end
                            if isequaln(lista_trabajo{t,6},lista_trabajo{t+1,6})
                                a=a+lista_trabajo{t,7};
                                n=n+1;
                            else
                                n=n+1;
                                a=a+lista_trabajo{t,7};
                                media=a/n;

                                for x=(t-(n-1)):t
                                    lista_trabajo{x,7}=media;
                                end

                                a=0;
                                n=0;
                            end
                        end
                    end
                end
            end
        end
    end

```

```

        end
    end
end
end

for k=2:dimensiones_trabajo(2,j)
    if isequaln(tablas_trabajo{1,j}{1,k},lista_trabajo{i,3})
        for z=2:dimensiones_trabajo(1,j)
            if isequaln(tablas_trabajo{1,j}{z,1},lista_trabajo{i,6})
                tablas_trabajo{1,j}{z,k}=lista_trabajo{i,7};
            end
        end
    end
end
end

for k=2:dimensiones_trabajo(2,j)
    for z=2:dimensiones_trabajo(1,j)
        if isequaln(tablas_trabajo{1,j}{z,k},[])
            tablas_trabajo{1,j}{z,k}=missing;
            tablas_trabajo{1,j}{z,k}=str2double(tablas_trabajo{1,j}{z,k});
        end
    end
end

% Matrices para IBM SPSS
r=1;
for k=1:dimensiones_trabajo(2,j)
    x=0;
    y=0;
    w=0;
    a=0;
    for z=1:dimensiones_trabajo(1,j)
        if isequaln(tablas_trabajo{1,j}{z,k},NaN)
            y=y+1;
            w=w+1;
        else
            w=0;
        end
        if z==dimensiones_trabajo(1,j)
            else
                if isequaln(tablas_trabajo{1,j}{z+1,k},NaN)
                    else if w==2
                        a=a+1;
                    end
                end
            end
        end
        x=x+1;
    end
    if (y/x)<0.3 || a>4
        for z=1:dimensiones_trabajo(1,j)
            matrices_trabajo{1,j}{z,r}=tablas_trabajo{1,j}{z,k};
        end
        r=r+1;
    end
end

dimensiones_matrices_trabajo(1,j)=size(matrices_trabajo{1,j},1);
dimensiones_matrices_trabajo(2,j)=size(matrices_trabajo{1,j},2);

%Interpolación datos trimestrales
for k=1:dimensiones_matrices_trabajo(2,j)
    w=0;
    a=0;
    for z=1:dimensiones_matrices_trabajo(1,j)
        if isequaln(matrices_trabajo{1,j}{z,k},NaN)
            w=w+1;
        else
            w=0;
        end
    end
end

```

```

        if z==dimensiones_matrices_trabajo(1,j)
        else
            if isequaln(matrices_trabajo{1,j}{z+1,k},NaN)
            else if w==2
                a=a+1;
            end
            end
        end
    end
end
w=0;
if a>4
    for z=1:dimensiones_trabajo(1,j)
        if isequaln(matrices_trabajo{1,j}{z,k},NaN)
            w=w+1;
        else
            w=0;
        end
        if z==dimensiones_matrices_trabajo(1,j) || z<4
        else
            if isequaln(matrices_trabajo{1,j}{z+1,k},NaN)
            else if w==2
                matrices_trabajo{1,j}{z-1,k}= ((matrices_trabajo{1,j}{z+1,k} -
                    matrices_trabajo{1,j}{z-2,k})/3)+
                    matrices_trabajo{1,j}{z-2,k};
                matrices_trabajo{1,j}{z,k}= ((matrices_trabajo{1,j}{z+1,k} -
                    matrices_trabajo{1,j}{z-2,k})/3)*2+
                    matrices_trabajo{1,j}{z-2,k};
            end
        end
    end
end
end
end
end

for k=1:dimensiones_matrices_trabajo(2,j)
    for z=1:dimensiones_matrices_trabajo(1,j)
        if isequaln(matrices_trabajo{1,j}{z,k},NaN)
            ausentes(z,k)=1;
        else
            ausentes(z,k)=0;
        end
    end
end

% Matrices con listwise deletion para hacer el RMSE
suma=sum(ausentes,2);
for z=1:dimensiones_trabajo(1,j)
    if suma(z)>0
        ausentes(z,:)=1;
    end
end

for k=1:dimensiones_matrices_trabajo(2,j)
    u=1;
    for z=1:dimensiones_matrices_trabajo(1,j)
        if ausentes(z,k)==0
            listwise_trabajo{1,j}{u,k}=matrices_trabajo{1,j}{z,k};
            RMSE_trabajo{1,j}{u,k}=matrices_trabajo{1,j}{z,k}; %de esta se eliminarán
                posteriormente dos valores
            u=u+1;
        end
    end
end

clear ausentes

dimensiones_listwise_trabajo(1,j)=size(listwise_trabajo{1,j},1);
dimensiones_listwise_trabajo(2,j)=size(listwise_trabajo{1,j},2);

%Randomly delete 2 values & select 2 reference variables
pos_1=round(rand*(dimensiones_listwise_trabajo(1,j)));
pos_2=round(rand*(dimensiones_listwise_trabajo(1,j)));
variable_ref1=round(rand*(dimensiones_listwise_trabajo(2,j)));
variable_ref2=round(rand*(dimensiones_listwise_trabajo(2,j)));

```

```

while pos_1<2 || pos_2<2 || pos_1==pos_2
pos_1=round(rand*(dimensiones_listwise_trabajo(1,j)));
pos_2=round(rand*(dimensiones_listwise_trabajo(1,j)));
end

while (variable_ref1<2 || variable_ref2<2) || (variable_ref1==variable_ref2 &&
dimensiones_listwise_trabajo(2,j)>2)
variable_ref1=round(rand*(dimensiones_listwise_trabajo(2,j)));
variable_ref2=round(rand*(dimensiones_listwise_trabajo(2,j)));
end

for k=2:dimensiones_listwise_trabajo(2,j)
if k==variable_ref1 || k== variable_ref2
else
RMSE_trabajo{1,j}{pos_1,k}=str2double(missing);
RMSE_trabajo{1,j}{pos_2,k}=str2double(missing);
end
end

end

%% Creación de los archivos EXCEL y presentación tras la selección de datos

for i=1:n_cat
if isequaln(matrices_trabajo{1,i},[])
else
xlswrite([strcat('Tabla', num2str(i)), ' Mercado de Trabajo.xlsx'],
matrices_trabajo{1,i});
end

if isequaln(RMSE_trabajo{1,i},[])
else
xlswrite([strcat('Imputacion/RMSE', num2str(i)), ' Mercado de Trabajo.xlsx'],
RMSE_trabajo{1,i});
end
end

clear w a r x y j n v_tiempo_trabajo trabajo variable_ref1 variable_ref2 u suma pos_1 pos_2
inicio_covid_date

```

Código 1. Selección de datos y creación de archivos para la imputación

Tras la selección y la creación de los archivos para la imputación múltiple en IBM SPSS, el código 2 refleja la conclusión de la etapa de imputación, en la que se comparan los valores de NRMSE obtenidos a partir de los tres métodos de imputación estudiados:

```

%% Contrastar métodos de imputación

for i=1:n_cat

IM_trabajo{1,i}=xlsread([strcat('Imputacion/IM', num2str(i)), '.xlsx']);
EM_trabajo{1,i}=xlsread([strcat('Imputacion/EM', num2str(i)), '.xlsx']);
Reg_trabajo{1,i}=xlsread([strcat('Imputacion/Reg', num2str(i)), '.xlsx']);

end

RMSE_IM=cell(1,n_cat);
RMSE_EM=cell(1,n_cat);
RMSE_Reg=cell(1,n_cat);

for i=1:n_cat

longitud=(size(IM_trabajo{1,i},1)); %Comparten todos la misma longitud
ancho=size(IM_trabajo{1,i},2); %Comparten todos el mismo ancho

for k=1:ancho

```

```

prom{1,i}(1,k)=0;
RMSE_IM{1,i}(1,k)=0;
RMSE_EM{1,i}(1,k)=0;
RMSE_Reg{1,i}(1,k)=0;

for z=1:longitud
    RMSE_IM{1,i}(1,k)=RMSE_IM{1,i}(1,k)+(IM_trabajo{1,i}(z,k)-
        listwise_trabajo{1,i}{z+1,k+1})^2;
    RMSE_EM{1,i}(1,k)=RMSE_EM{1,i}(1,k)+(EM_trabajo{1,i}(z,k)-
        listwise_trabajo{1,i}{z+1,k+1})^2;
    RMSE_Reg{1,i}(1,k)=RMSE_Reg{1,i}(1,k)+(Reg_trabajo{1,i}(z,k)-
        listwise_trabajo{1,i}{z+1,k+1})^2;
    prom{1,i}(1,k)=prom{1,i}(1,k)+listwise_trabajo{1,i}{z+1,k+1};
end

prom{1,i}(1,k)=prom{1,i}(1,k)/longitud;
RMSE_IM{1,i}(1,k)=(RMSE_IM{1,i}(1,k)/(longitud))^(1/2);
RMSE_EM{1,i}(1,k)=(RMSE_EM{1,i}(1,k)/(longitud))^(1/2);
RMSE_Reg{1,i}(1,k)=(RMSE_Reg{1,i}(1,k)/(longitud))^(1/2);

%Normalized Root Mean Square Error
NRMSE_IM{1,i}(1,k)=(RMSE_IM{1,i}(1,k)/prom{1,i}(1,k));
NRMSE_EM{1,i}(1,k)=(RMSE_EM{1,i}(1,k)/prom{1,i}(1,k));
NRMSE_Reg{1,i}(1,k)=(RMSE_Reg{1,i}(1,k)/prom{1,i}(1,k));

end

NRMSE_IM{1,i}=abs(sum(NRMSE_IM{1,i}(1,:)));
NRMSE_EM{1,i}=abs(sum(NRMSE_EM{1,i}(1,:)));
NRMSE_Reg{1,i}=abs(sum(NRMSE_Reg{1,i}(1,:)));

if NRMSE_IM{1,i}<NRMSE_EM{1,i} && NRMSE_IM{1,i}<NRMSE_Reg{1,i}
    IMP{1,i}='IM';
else if NRMSE_EM{1,i}<NRMSE_IM{1,i} && NRMSE_EM{1,i}<NRMSE_Reg{1,i}
    IMP{1,i}='EM';
else if NRMSE_EM{1,i}==NRMSE_IM{1,i} && NRMSE_EM{1,i} == NRMSE_Reg{1,i}
    IMP{1,i}=[];
else
    IMP{1,i}='Reg';
end

end

end

clear z i k prom ancho longitud

```

Código 2. Comparación de métodos de imputación según su error NRMSE

Tras la selección de la técnica de imputación y su implementación en IBM SPSS, el siguiente código 3 representa la etapa de normalización del indicador (junto a la “saturación” que se estudia en la sección 3.8), de forma que se aplica el método Mín-Máx y se invierten los valores de algunas variables para emplear un sentido común de interpretación:

```

%% Datos imputados

for i=1:n_cat
    imputacion{1,i}=xlsread([strcat('Imputacion/Tabla', num2str(i)), '_Imp Mercado de
        Trabajo.xlsx']);
end

imputacion_sat=imputacion; %Matriz donde saturar los outliers

%% Normalización min-max
for i=1:n_cat

```

```

for k=1:size(imputacion{1,i},2)
    outliers=isoutlier(imputacion{1,i}(:,k), 'quartiles');
    for z=1:size(outliers,1)
        if outliers(z,1)==1
            imputacion_sat{1,i}(z,k)=NaN;
        end
    end
    minimo=min(imputacion{1,i}(:,k));
    maximo=max(imputacion{1,i}(:,k));
    minimo_sat=min(imputacion_sat{1,i}(:,k));
    maximo_sat=max(imputacion_sat{1,i}(:,k));
    media=mean(imputacion{1,i}(:,k));
    for z=1:size(imputacion{1,i},1)
        normalizacion{1,i}(z,k)=(imputacion{1,i}(z,k)-minimo)/(maximo-minimo);
        if outliers(z,1)==1
            if imputacion{1,i}(z,k)>media
                imputacion_sat{1,i}(z,k)=maximo_sat;
            else
                imputacion_sat{1,i}(z,k)=minimo_sat;
            end
        end
        normalizacion_sat{1,i}(z,k) =
            (imputacion_sat{1,i}(z,k)-minimo_sat)/(maximo_sat-minimo_sat);
    end
end
for z=1:size(matrices_trabajo{1,i},1)
    for k=1:size(matrices_trabajo{1,i},2)
        if z==1 || k==1
            trabajo_norm{1,i}{z,k}=matrices_trabajo{1,i}{z,k};
            trabajo_norm_sat{1,i}{z,k}=matrices_trabajo{1,i}{z,k};
        else
            trabajo_norm{1,i}{z,k}=normalizacion{1,i}(z-1,k-1);
            trabajo_norm_sat{1,i}{z,k}=normalizacion_sat{1,i}(z-1,k-1);
        end
    end
end
end
end

clear minimo maximo media maximo_sat minimo_sat outliers

%% Sentido indicadores (+ -) sobre el resumen
for i=1:size(posicion_kpi,2)
    if isequaln(resumen_trabajo{i,1},'Desempleo') ||
        isequaln(resumen_trabajo{i,1},'Habilidades y Capacitación') ||
        isequaln(resumen_trabajo{i,1},'Horas Trabajadas')
        resumen_trabajo{i,6}=0;
    else if isequaln(resumen_trabajo{i,1},'Empleo') ||
        isequaln(resumen_trabajo{i,1},'Salarios y costes laborales') ||
        isequaln(resumen_trabajo{i,1},'Protección a los desempleados')
        resumen_trabajo{i,6}=1;
    end
end

%% Excepciones
if isequaln(resumen_trabajo{i,2},'Bajas Afiliados')
    resumen_trabajo{i,6}=0;
end

end

xlswrite('Mercado de Trabajo Resumen.xlsx',resumen_trabajo);

%% Normalizar unidireccional min-max
i=1;
for j=1:n_cat
    k=2;
    while k~=(dimensiones_matrices_trabajo(2,j)+1)
        if isequaln(resumen_trabajo{i,2},trabajo_norm{1,j}{1,k})
            if resumen_trabajo{i,6}==0
                for z=2:dimensiones_matrices_trabajo(1,j)
                    trabajo_norm{1,j}{z,k}=1-trabajo_norm{1,j}{z,k};
                    trabajo_norm_sat{1,j}{z,k}=1-trabajo_norm_sat{1,j}{z,k};
                end
            end
            k=k+1;
        end
    end
end

```

```

        end
        i=i+1;
    end
end

clear j k z normalizacion normalizacion_sat

% Situación actual: Normalización min-max completada

for i=1:n_cat
    xlswrite([strcat('Normalizacion/Tabla ', num2str(i)), '_Norm Mercado de Trabajo.xlsx'],
        trabajo_norm{1,i});
    xlswrite([strcat('Normalizacion/Tabla ', num2str(i)), '_Norm_Sat Mercado de
        Trabajo.xlsx'], trabajo_norm_sat{1,i});
end

```

Código 3. Proceso de normalización Mín-Máx y la reorientación de escala posterior

El siguiente código 4 arroja información relevante a la sección 3.6, en tanto que permite comprender la verdadera dimensionalidad de los conjuntos de datos estudiados:

```

clear A correlation x C B eigenvalues Param_NPD n_elim matrices_est x y
x=zeros(1, n_cat);
y=zeros(1, n_cat);
for i=1:n_cat
    A=cell2mat(trabajo_norm{1,i}(2:size(trabajo_norm{1,i},1),2:size(trabajo_norm{1,i},2)));
    correlation{1,i}=corrcoef(A);
    eigenvalues{1,i}=eig(correlation{1,i});
    Param_NPD(1,i)=rcond(correlation{1,i});
    n_elim(1,i)=size(correlation{1,i},1)-rank(correlation{1,i});
    if n_elim(1,i)>0
        for k=1:size(correlation{1,i},2)
            B=A;
            B(:,k)=[];
            C=corrcoef(B);
            if rcond(C)>Param_NPD(1,i)
                Param_NPD(1,i)=rcond(C);
                x(1,i)=k;
            end
        end
        end
    matrices_est{1,i}=A;

    if x(1,i)>0
        matrices_est{1,i}(:,x(1,i))=[];
    end

    correlation{1,i}=corrcoef(matrices_est{1,i});
    eigenvalues{1,i}=eig(correlation{1,i});
    Param_NPD(1,i)=rcond(correlation{1,i});

    if n_elim(1,i)>0

        for k=1:(size(matrices_est{1,i},2))
            D=matrices_est{1,i};
            D(:,k)=[];
            E=corrcoef(D);
            if rcond(E)>Param_NPD(1,i)
                Param_NPD(1,i)=rcond(E);
                y(1,i)=k;
            end
        end
        end
    if y(1,i)>0
        matrices_est{1,i}(:,y(1,i))=[];
    end
    correlation{1,i}=corrcoef(matrices_est{1,i});
    eigenvalues{1,i}=eig(correlation{1,i});
    Param_NPD(1,i)=rcond(correlation{1,i});
end

```

```

end
end

```

Código 4. Identificación de sobredimensionalidad en las matrices de correlación de las categorías

El código 5 presenta la implementación en MATLAB de las decisiones tomadas en la sección 3.6:

```

%% Reformulación

descarte_11 = 'Paro Registrado';
descarte_12 = 'Paro Registrado Varones';
descarte_13 = 'Solicitudes Relativas / Vacantes por Mes Finanzas';
descarte_14 = 'Solicitudes Relativas / Vacantes por Mes Hospitality';
descarte_15 = 'Solicitudes Relativas / Vacantes por Mes Industria';
descarte_16 = 'Solicitudes Relativas / Vacantes por Mes Logística';
descarte_17 = 'Solicitudes Relativas / Vacantes por Mes Retail';
descarte_18 = 'Solicitudes Relativas / Vacantes por Mes Telemarketing';
descarte_19 = 'Solicitudes Relativas / Vacantes por Mes Transporte';
descarte_21 = 'Contratos Registrados';
descarte_22 = 'Altas Afiliados';
electo_31 = 'Tasa de Paro FP Mujeres';
electo_32 = 'Tasa de Paro Analfabetos Mujeres';
electo_33 = 'Tasa de Paro Educación Superior Mujeres';
electo_34 = 'Tasa de Paro Estudios Primarios Mujeres';
electo_35 = 'Tasa de Paro Segunda Etapa de Educación Secundaria Mujeres';
electo_36 = 'Tasa de Paro Educación Superior Ambos sexos';
for i=1:n_cat
    b=2;
    for z=1:size(trabajo_norm{1,i},1)
        trabajo_norm_ref{1,i}{z,1}=trabajo_norm{1,i}{z,1};
    end
    for k=2:size(trabajo_norm{1,i},2)
        if isequaln(trabajo_norm{1,i}{1,k}, descarte_11) || isequaln(trabajo_norm{1,i}{1,k},
            descarte_12) || isequaln(trabajo_norm{1,i}{1,k}, descarte_13) ||
            isequaln(trabajo_norm{1,i}{1,k}, descarte_14) || isequaln(trabajo_norm{1,i}{1,k},
            descarte_15) || isequaln(trabajo_norm{1,i}{1,k}, descarte_16) ||
            isequaln(trabajo_norm{1,i}{1,k}, descarte_17) || isequaln(trabajo_norm{1,i}{1,k},
            descarte_18) || isequaln(trabajo_norm{1,i}{1,k}, descarte_19) ||
            isequaln(trabajo_norm{1,i}{1,k}, descarte_21) || isequaln(trabajo_norm{1,i}{1,k},
            descarte_22)
        else
            if (isequaln(trabajo_norm{1,i}{1,k}, electo_31) ||
                isequaln(trabajo_norm{1,i}{1,k}, electo_32) ||
                isequaln(trabajo_norm{1,i}{1,k}, electo_33) ||
                isequaln(trabajo_norm{1,i}{1,k}, electo_34) ||
                isequaln(trabajo_norm{1,i}{1,k}, electo_35) ||
                isequaln(trabajo_norm{1,i}{1,k}, electo_36)) || i~=3
                for z=1:size(trabajo_norm{1,i},1)
                    trabajo_norm_ref{1,i}{z,b}=trabajo_norm{1,i}{z,k};
                end
                b=b+1;
            end
        end
    end
end
end

for i=1:n_cat
    b=2;
    for z=1:size(trabajo_norm_sat{1,i},1)
        trabajo_norm_ref_sat{1,i}{z,1}=trabajo_norm_sat{1,i}{z,1};
    end
    for k=2:size(trabajo_norm_sat{1,i},2)
        if isequaln(trabajo_norm_sat{1,i}{1,k}, descarte_11) ||
            isequaln(trabajo_norm_sat{1,i}{1,k}, descarte_12) ||
            isequaln(trabajo_norm_sat{1,i}{1,k}, descarte_13) ||
            isequaln(trabajo_norm_sat{1,i}{1,k}, descarte_14) ||
            isequaln(trabajo_norm_sat{1,i}{1,k}, descarte_15) ||

```

```

isequaln(trabajo_norm_sat{1,i}{1,k}, descarte_16) ||
isequaln(trabajo_norm_sat{1,i}{1,k}, descarte_17) ||
isequaln(trabajo_norm_sat{1,i}{1,k}, descarte_18) ||
isequaln(trabajo_norm_sat{1,i}{1,k}, descarte_19) ||
isequaln(trabajo_norm_sat{1,i}{1,k}, descarte_21)
||isequaln(trabajo_norm_sat{1,i}{1,k}, descarte_22)
else
    if isequaln(trabajo_norm_sat{1,i}{1,k}, electo_31) ||
        isequaln(trabajo_norm_sat{1,i}{1,k}, electo_32) ||
        isequaln(trabajo_norm_sat{1,i}{1,k}, electo_33) ||
        isequaln(trabajo_norm_sat{1,i}{1,k}, electo_34) ||
        isequaln(trabajo_norm_sat{1,i}{1,k}, electo_35) ||
        isequaln(trabajo_norm_sat{1,i}{1,k}, electo_36) || i~=3
        for z=1:size(trabajo_norm_sat{1,i},1)
            trabajo_norm_ref_sat{1,i}{z,b}=trabajo_norm_sat{1,i}{z,k};
        end
        b=b+1;
    end
end
end
clear trabajo_agreg_ref_directo buf a b c

```

Código 5. Reducción de las categorías según la evaluación intermedia de resultados

Una vez solventados los problemas de correlación, el código 6 implementa el proceso final de creación del sintético:

```

%% Ponderación
for i=1:n_cat
    varianza_comp_ref{1, i}=xlsread([strcat('Ponderacion/Reformulacion/Varianza',
        num2str(i)), '.xlsx']);
    varianza_comp_ref_sat{1, i}=xlsread([strcat('Ponderacion/Reformulacion/Varianza_sat',
        num2str(i)), '.xlsx']);
    factor_load_ref{1, i}=xlsread([strcat('Ponderacion/Reformulacion/Matriz rotada',
        num2str(i)), '.xlsx']);
    factor_load_ref_sat{1, i}=xlsread([strcat('Ponderacion/Reformulacion/Matriz rotada_sat',
        num2str(i)), '.xlsx']);

    A1=factor_load_ref{1, i}.*factor_load_ref{1, i};
    for k=1:size(A1, 2)
        factor_load_ref_sq{1, i}(:, k)=A1(:, k)/sum(A1(isfinite(A1(:, k))), k));
    end

    ancho=size(factor_load_ref{1, i}, 2);
    longitud=size(factor_load_ref{1, i}, 1);
    trabajo_factor_ref_sq{1, i}{1, 1}='Indicadores';
    for k=1:ancho
        trabajo_factor_ref_sq{1, i}{1, k+1}=strcat('Factor_', num2str(k));
        for z=1:longitud
            if k==1
                trabajo_factor_ref_sq{1, i}{z+1, 1}=trabajo_norm_ref{1, i}{1, z+1};
            end
            trabajo_factor_ref_sq{1, i}{z+1, k+1}=factor_load_ref_sq{1, i}(z, k);
        end
    end

    xlswrite([strcat('Ponderacion/Analisis/Reformulacion/Cargas factoriales ', num2str(i)),
        '.xlsx'], trabajo_factor_ref_sq{1, i});

    trabajo_comp_ref{1, i}{1, 1}=trabajo_norm_ref{1, i}{1, 1};

    for z=2:size(trabajo_norm_ref{1, i}, 1)
        trabajo_comp_ref{1, i}{z, 1}=trabajo_norm_ref{1, i}{z, 1};
    end
end

```

```

end

for k=2:(ancho+1)
    trabajo_comp_ref{1, i}{1, k}=strcat('Factor_', num2str(k-1));
    for z=2:size(trabajo_comp_ref{1, i}, 1)
        trabajo_comp_ref{1, i}{z, k}=0;
        for x=1:longitud
            trabajo_comp_ref{1, i}{z, k}=trabajo_comp_ref{1, i}{z, k} +
                factor_load_ref_sq{1, i}(x, k-1)*trabajo_norm_ref{1, i}{z, x+1};
        end
    end
end

%Lo mismo para saturado
ancho_sat=size(factor_load_ref_sat{1, i}, 2);
longitud_sat=size(factor_load_ref_sat{1, i}, 1);
Al_sat=factor_load_ref_sat{1, i}.*factor_load_ref_sat{1, i};
for k=1:size(Al_sat, 2)
    factor_load_ref_sq_sat{1, i}(:, k) = Al_sat(:, k)/sum(Al_sat(isfinite(Al_sat(:, k)),
        k));
end

trabajo_factor_ref_sq_sat{1, i}{1, 1}='Indicadores';
for k=1:ancho_sat
    trabajo_factor_ref_sq_sat{1, i}{1, k+1}=strcat('Factor_', num2str(k));
    for z=1:longitud_sat
        if k==1
            trabajo_factor_ref_sq_sat{1, i}{z+1, 1}= trabajo_norm_ref_sat{1, i}{1, z+1};
        end
        trabajo_factor_ref_sq_sat{1, i}{z+1, k+1}= factor_load_ref_sq_sat{1, i}(z, k);
    end
end

xlswrite([strcat('Ponderacion/Analisis/Reformulacion/Cargas factoriales saturadas',
    num2str(i)), '.xlsx'], trabajo_factor_ref_sq_sat{1, i});

trabajo_comp_ref_sat{1, i}{1, 1}=trabajo_norm_ref_sat{1, i}{1, 1};

for z=2:size(trabajo_norm_ref_sat{1, i}, 1)
    trabajo_comp_ref_sat{1, i}{z, 1}=trabajo_norm_ref_sat{1, i}{z, 1};
end

for k=2:(ancho_sat+1)
    trabajo_comp_ref_sat{1, i}{1, k}=strcat('Factor_', num2str(k-1));
    for z=2:size(trabajo_comp_ref_sat{1, i}, 1)
        trabajo_comp_ref_sat{1, i}{z, k}=0;
        for x=1:longitud_sat
            trabajo_comp_ref_sat{1, i}{z, k}=trabajo_comp_ref_sat{1, i}{z, k}+
                factor_load_ref_sq_sat{1, i}(x, k-1)*trabajo_norm_ref_sat{1, i}{z, x+1};
        end
    end
end

clear ancho_sat ancho longitud longitud_sat

%% Asignar pesos criterio experto agregación final
pesos_cat=[0.25 0.25 0.2 0.1 0.2];

%% Agregación

%Arreglar diferencia periodos temporales por categoría
long=size(trabajo_comp_ref{1, 1}, 1);
periodo_cat=1;

for i=2:n_cat
    if long>size(trabajo_comp_ref{1, i}, 1)
        long=size(trabajo_comp_ref{1, i}, 1);
        periodo_cat=i;
    end
end

```

```

end

long_sat=size(trabajo_comp_ref_sat{1, 1}, 1);
periodo_cat_sat=1;

for i=2:n_cat
    if long_sat>size(trabajo_comp_ref_sat{1, i}, 1)
        long_sat=size(trabajo_comp_ref_sat{1, i}, 1);
        periodo_cat_sat=i;
    end
end

%Emplazar títulos de las categorías
i=1;
for x=2:longitud_trabajo
    if isequaln(lista_trabajo{x, 2}, lista_trabajo{x-1, 2})
    else
        trabajo_agreg_ref{1, i+1}=lista_trabajo{x, 2};
        trabajo_agreg_ref_sat{1, i+1}=lista_trabajo{x, 2};
        i=i+1;
    end
end
trabajo_agreg_ref{1, 1}='Fecha';
trabajo_agreg_ref_sat{1, 1}='Fecha';

trabajo_agreg_ref_directo{1, 1}='Fecha';
i=1;
for x=2:longitud_trabajo
    if isequaln(lista_trabajo{x, 2}, lista_trabajo{x-1, 2})
    else
        trabajo_agreg_ref_directo{1, i+1}=lista_trabajo{x, 2};
        i=i+1;
    end
end

% Bucle subagregación
for i=1:n_cat
    clear A1 b
    A1=factor_load_ref{1, i}.*factor_load_ref{1, i};
    %Método directo
    for k=1:size(A1, 2)
        b(1, k)=sum(A1(:, k));
    end

    for z=1:size(factor_load_ref{1, i}, 1)
        for k=1:size(factor_load_ref{1, i}, 2)
            c{1, i}{z, k}=A1(z, k)/sum(b); %Varianza explicada por el indicador z en el
            factor k
        end
    end
    peso_factorial{1, i}=transpose(sum(c{1, i}, 2));
    longitud=size(trabajo_norm_ref{1, i}, 1);
    ancho=size(trabajo_norm_ref{1, i}, 2);

    x=2;
    for z=2:longitud
        trabajo_agreg_ref_directo{x, 1}=trabajo_norm_ref{1, periodo_cat}{x, 1};
        if trabajo_agreg_ref_directo{x, 1}==trabajo_norm_ref{1, i}{z, 1}
            trabajo_agreg_ref_directo{x, i+1}=0;
            if varianza_comp_ref{1, i}(size(factor_load_ref{1, i}, 2),
                size(varianza_comp_ref{1, i}, 2))>80 % Criterio de exclusión para equiponderar
                for k=2:ancho
                    trabajo_agreg_ref_directo{x, i+1}=trabajo_agreg_ref_directo{x, i+1} +
                        peso_factorial{1, i}(1, k-1)*trabajo_norm_ref{1, i}{z, k};
                end
            else
                for k=2:ancho
                    trabajo_agreg_ref_directo{x, i+1}=trabajo_agreg_ref_directo{x, i+1} +
                        trabajo_norm_ref{1, i}{z, k}/(size(trabajo_norm_ref{1, i}, 2)-1);
                end
            end
            x=x+1;
        end
    end
end

```

```

% Método a partir de los indicadores intermedios
longitud=size(trabajo_comp_ref{1, i}, 1);
ancho=size(trabajo_comp_ref{1, i}, 2);
x=2;
for z=2:longitud
    trabajo_agreg_ref{x, 1}=trabajo_comp_ref{1, periodo_cat}{x, 1};
    if trabajo_agreg_ref{x, 1}==trabajo_comp_ref{1, i}{z, 1}
        trabajo_agreg_ref{x, i+1}=0;
        if varianza_comp_ref{1, i}(ancho-1, size(varianza_comp_ref{1, i}, 2))>80
            for k=2:ancho
                trabajo_agreg_ref{x, i+1}=trabajo_agreg_ref{x, i+1} +
                    varianza_comp_ref{1, i}(k-1, (size(varianza_comp_ref{1, i},
                        2))-1)*trabajo_comp_ref{1, i}{z, k}/varianza_comp_ref{1, i}(ancho-1,
                        size(varianza_comp_ref{1, i}, 2));
            end
        else
            for k=2:size(trabajo_norm_ref{1, i}, 2)
                trabajo_agreg_ref{x, i+1}=trabajo_agreg_ref{x, i+1} + trabajo_norm_ref{1,
                    i}{z, k}/(size(trabajo_norm_ref{1, i}, 2)-1);
            end
        end
        x=x+1;
    end
end

longitud_sat=size(trabajo_comp_ref{1, i}, 1);
ancho_sat=size(trabajo_comp_ref{1, i}, 2);
x=2;
for z=2:longitud_sat
    trabajo_agreg_ref_sat{x, 1}=trabajo_comp_ref_sat{1, periodo_cat}{x, 1};
    if trabajo_agreg_ref_sat{x, 1}==trabajo_comp_ref_sat{1, i}{z, 1}
        trabajo_agreg_ref_sat{x, i+1}=0;
        if varianza_comp_ref_sat{1, i}(ancho_sat-1, size(varianza_comp_ref_sat{1, i},
            2))>80
            for k=2:ancho
                trabajo_agreg_ref_sat{x, i+1}=trabajo_agreg_ref_sat{x, i+1} +
                    varianza_comp_ref_sat{1, i}(k-1, (size(varianza_comp_ref_sat{1, i},
                        2))-1)*trabajo_comp_ref_sat{1, i}{z, k}/varianza_comp_ref_sat{1,
                        i}(ancho_sat-1, size(varianza_comp_ref_sat{1, i}, 2));
            end
        else
            for k=2:size(trabajo_norm_ref_sat{1, i}, 2)
                trabajo_agreg_ref_sat{x, i+1}=trabajo_agreg_ref_sat{x, i+1} +
                    trabajo_norm_ref_sat{1, i}{z, k}/(size(trabajo_norm_ref_sat{1, i},
                        2)-1);
            end
        end
        x=x+1;
    end
end

end

%% Consideración método directo y no directo

%% Ambos métodos proporcionan el mismo output:
%trabajo_agreg_ref_directo=trabajo_agreg_ref

%%Una vez ratificado esto, no hay motivos para duplicar el resto de
%resultados

%% Representación

figure(1);
for k=2:size(trabajo_agreg_ref, 2)
    plot(vt{1, periodo_cat}, (cell2mat(trabajo_agreg_ref(2:size(trabajo_agreg_ref, 1), k))));
    hold on
end
legend('Desempleo', 'Empleo', 'Habilidades y capacitación', 'Protec. a los desempleados',
    'Salarios y costes lab.');
```

```

figure(2);
for k=2:size(trabajo_agreg_ref_sat, 2)
    plot(vt{1, periodo_cat_sat},
        cell2mat(trabajo_agreg_ref_sat(2:size(trabajo_agreg_ref_sat, 1), k)));
    hold on
end
legend('Desempleo', 'Empleo', 'Habilidades y capacitación', 'Protec. a los desempleados',
    'Salarios y costes lab.');
```

```

legend('Location', 'southoutside', 'NumColumns', 3);
title('Objetivos del sintético con normalización saturada');
hold off
grid
print -depsc Graficas/reformulado_areas_saturado

xlswrite(strcat('Agregacion/Reformulacion/Áreas del mercado de trabajo.xlsx'),
    trabajo_agreg_ref);
xlswrite(strcat('Agregacion/Reformulacion/Áreas del mercado de trabajo con normalización
    saturada.xlsx'), trabajo_agreg_ref_sat);

% Bucle final agregación
trabajo_indice_ref_equi_arit{1, 1}='Fecha';
trabajo_indice_ref_exp_arit{1, 1}='Fecha';
trabajo_indice_ref_sat_equi_arit{1, 1}='Fecha';
trabajo_indice_ref_sat_exp_arit{1, 1}='Fecha';
trabajo_indice_ref_equi_arit{1, 2}='Ref_equi_arit';
trabajo_indice_ref_exp_arit{1, 2}='Ref_exp_arit';
trabajo_indice_ref_sat_equi_arit{1, 2}='Ref_sat_equi_arit';
trabajo_indice_ref_sat_exp_arit{1, 2}='Ref_sat_exp_arit';

trabajo_indice_ref_equi_geom{1, 1}='Fecha';
trabajo_indice_ref_exp_geom{1, 1}='Fecha';
trabajo_indice_ref_sat_equi_geom{1, 1}='Fecha';
trabajo_indice_ref_sat_exp_geom{1, 1}='Fecha';
trabajo_indice_ref_equi_geom{1, 2}='Ref_equi_geom';
trabajo_indice_ref_exp_geom{1, 2}='Ref_exp_geom';
trabajo_indice_ref_sat_equi_geom{1, 2}='Ref_sat_equi_geom';
trabajo_indice_ref_sat_exp_geom{1, 2}='Ref_sat_exp_geom';

for z=2:long
    trabajo_indice_ref_equi_arit{z, 2}=0;
    trabajo_indice_ref_exp_arit{z, 2}=0;
    trabajo_indice_ref_equi_arit{z, 1}=trabajo_agreg_ref{z, 1};
    trabajo_indice_ref_exp_arit{z, 1}=trabajo_agreg_ref{z, 1};
    trabajo_indice_ref_equi_geom{z, 2}=1;
    trabajo_indice_ref_exp_geom{z, 2}=1;
    trabajo_indice_ref_equi_geom{z, 1}=trabajo_agreg_ref{z, 1};
    trabajo_indice_ref_exp_geom{z, 1}=trabajo_agreg_ref{z, 1};

    for k=2:size(trabajo_agreg_ref, 2)
        trabajo_indice_ref_equi_arit{z, 2}=trabajo_indice_ref_equi_arit{z, 2} +
            trabajo_agreg_ref{z, k}/n_cat;
        trabajo_indice_ref_exp_arit{z, 2}=trabajo_indice_ref_exp_arit{z, 2} +
            trabajo_agreg_ref{z, k}*pesos_cat(1, k-1);

        trabajo_indice_ref_equi_geom{z, 2}= trabajo_indice_ref_equi_geom{z,
            2}*trabajo_agreg_ref{z, k}^(1/n_cat);
        trabajo_indice_ref_exp_geom{z, 2}= trabajo_indice_ref_exp_geom{z,
            2}*trabajo_agreg_ref{z, k}^pesos_cat(1, k-1);

    end
end

for z=2:long_sat
    trabajo_indice_ref_sat_equi_arit{z, 2}=0;
    trabajo_indice_ref_sat_exp_arit{z, 2}=0;
    trabajo_indice_ref_sat_equi_arit{z, 1}=trabajo_agreg_ref_sat{z, 1};
    trabajo_indice_ref_sat_exp_arit{z, 1}=trabajo_agreg_ref_sat{z, 1};

    trabajo_indice_ref_sat_equi_geom{z, 2}=1;
    trabajo_indice_ref_sat_exp_geom{z, 2}=1;
    trabajo_indice_ref_sat_equi_geom{z, 1}=trabajo_agreg_ref_sat{z, 1};

```

```

trabajo_indice_ref_sat_exp_geom(z, 1)=trabajo_agreg_ref_sat{z, 1};

for k=2:size(trabajo_agreg_ref_sat, 2)
    trabajo_indice_ref_sat_equi_arit{z, 2}=trabajo_indice_ref_sat_equi_arit{z, 2}+
        trabajo_agreg_ref_sat{z, k}/n_cat;
    trabajo_indice_ref_sat_exp_arit{z, 2}=trabajo_indice_ref_sat_exp_arit{z, 2}+
        trabajo_agreg_ref_sat{z, k}*pesos_cat(1, k-1);

    trabajo_indice_ref_sat_equi_geom{z, 2}= trabajo_indice_ref_sat_equi_geom{z,
        2}*trabajo_agreg_ref_sat{z, k}^(1/n_cat);
    trabajo_indice_ref_sat_exp_geom{z, 2}= trabajo_indice_ref_sat_exp_geom{z,
        2}*trabajo_agreg_ref_sat{z, k}^pesos_cat(1, k-1);
end
end

figure(3);
plot(vt{1, periodo_cat},
    cell2mat(trabajo_indice_ref_equi_arit(2:size(trabajo_indice_ref_equi_arit, 1), 2)));
hold on;
plot(vt{1, periodo_cat},
    cell2mat(trabajo_indice_ref_exp_arit(2:size(trabajo_indice_ref_exp_arit, 1), 2)));
hold on;
plot(vt{1, periodo_cat},
    cell2mat(trabajo_indice_ref_sat_equi_arit(2:size(trabajo_indice_ref_sat_equi_arit, 1),
        2)));
hold on;
plot(vt{1, periodo_cat},
    cell2mat(trabajo_indice_ref_sat_exp_arit(2:size(trabajo_indice_ref_sat_exp_arit, 1), 2)));
legend('Equiponderación', 'Criterio experto', 'Equiponderación con norm. saturada',
    'Criterio experto con norm. saturada');
legend('Location', 'southoutside', 'NumColumns', 2);
title('Indicador sintético con agregación aritmética');
hold off
grid
print -depsc Graficas/reformulado_sintético_aritmética

figure(4);
plot(vt{1, periodo_cat},
    cell2mat(trabajo_indice_ref_equi_geom(2:size(trabajo_indice_ref_equi_geom, 1), 2)));
hold on;
plot(vt{1, periodo_cat},
    cell2mat(trabajo_indice_ref_exp_geom(2:size(trabajo_indice_ref_exp_geom, 1), 2)));
hold on;
plot(vt{1, periodo_cat},
    cell2mat(trabajo_indice_ref_sat_equi_geom(2:size(trabajo_indice_ref_sat_equi_geom, 1),
        2)));
hold on;
plot(vt{1, periodo_cat},
    cell2mat(trabajo_indice_ref_sat_exp_geom(2:size(trabajo_indice_ref_sat_exp_geom, 1), 2)));
legend('Equiponderación', 'Criterio experto', 'Equiponderación con norm. saturada',
    'Criterio experto con norm. saturada');
legend('Location', 'southoutside', 'NumColumns', 2);
title('Indicador sintético con agregación geométrica');
hold off
grid
print -depsc Graficas/reformulado_sintético_geométrica

xlswrite(strcat('Agregacion/Reformulacion/Indice equiponderado aritmético.xlsx'),
    trabajo_indice_ref_equi_arit);
xlswrite(strcat('Agregacion/Reformulacion/Indice criterio experto aritmético.xlsx'),
    trabajo_indice_ref_exp_arit);
xlswrite(strcat('Agregacion/Reformulacion/Indice norm_saturada equiponderado
    aritmético.xlsx'), trabajo_indice_ref_sat_equi_arit);
xlswrite(strcat('Agregacion/Reformulacion/Indice norm_saturada criterio experto
    aritmético.xlsx'), trabajo_indice_ref_sat_exp_arit);

xlswrite(strcat('Agregacion/Reformulacion/Indice equiponderado geométrico.xlsx'),
    trabajo_indice_ref_equi_geom);
xlswrite(strcat('Agregacion/Reformulacion/Indice criterio experto geométrico.xlsx'),
    trabajo_indice_ref_exp_geom);
xlswrite(strcat('Agregacion/Reformulacion/Indice norm_saturada equiponderado
    geométrico.xlsx'), trabajo_indice_ref_sat_equi_geom);
xlswrite(strcat('Agregacion/Reformulacion/Indice norm_saturada criterio experto
    geométrico.xlsx'), trabajo_indice_ref_sat_exp_geom);

```

```

%% Ponderación (normal)

for i=1:n_cat
    varianza_comp{1, i}=xlsread([strcat('Ponderacion/Normal/Varianza', num2str(i)),
    '.xlsx']);
    varianza_comp_sat{1, i}=xlsread([strcat('Ponderacion/Normal/Varianza_sat', num2str(i)),
    '.xlsx']);
    factor_load{1, i}=xlsread([strcat('Ponderacion/Normal/Matriz rotada', num2str(i)),
    '.xlsx']);
    factor_load_sat{1, i}=xlsread([strcat('Ponderacion/Normal/Matriz rotada_sat',
    num2str(i)), '.xlsx']);

    ancho=size(factor_load{1, i}, 2);
    longitud=size(factor_load{1, i}, 1);
    A1=factor_load{1, i}.*factor_load{1, i};

    for k=1:size(A1, 2)
        factor_load_sq{1, i}(:, k)=A1(:, k)/sum(A1(isfinite(A1(:, k))), k));
    end

    trabajo_factor_sq{1, i}{1, 1}='Indicadores';
    for k=1:ancho
        trabajo_factor_sq{1, i}{1, k+1}=strcat('Factor_', num2str(k));
        for z=1:longitud
            if k==1
                trabajo_factor_sq{1, i}{z+1, 1}=trabajo_norm{1, i}{1, z+1};
            end
            trabajo_factor_sq{1, i}{z+1, k+1}=factor_load_sq{1, i}(z, k);
        end
    end

    xlswrite([strcat('Ponderacion/Analisis/Normal/Cargas factoriales ', num2str(i)),
    '.xlsx'], trabajo_factor_sq{1, i});

    trabajo_comp{1, i}{1, 1}=trabajo_norm{1, i}{1, 1};

    for z=2:size(trabajo_norm{1, i}, 1)
        trabajo_comp{1, i}{z, 1}=trabajo_norm{1, i}{z, 1};
    end

    for k=2:(ancho+1)
        trabajo_comp{1, i}{1, k}=strcat('Factor_', num2str(k-1));
        for z=2:size(trabajo_comp{1, i}, 1)
            trabajo_comp{1, i}{z, k}=0;
            for x=1:longitud
                trabajo_comp{1, i}{z, k}= trabajo_comp{1, i}{z, k}+factor_load_sq{1, i}(x,
                k-1)*trabajo_norm{1, i}{z, x+1};
            end
        end
    end

    %Lo mismo para saturado
    ancho_sat=size(factor_load_sat{1, i}, 2);
    longitud_sat=size(factor_load_sat{1, i}, 1);
    A1_sat=factor_load_sat{1, i}.*factor_load_sat{1, i};
    for k=1:size(A1_sat, 2)
        factor_load_sq_sat{1, i}(:, k)= A1_sat(:, k)/sum(A1_sat(isfinite(A1_sat(:, k))), k));
    end

    trabajo_factor_sq_sat{1, i}{1, 1}='Indicadores';
    for k=1:ancho_sat
        trabajo_factor_sq_sat{1, i}{1, k+1}=strcat('Factor_', num2str(k));
        for z=1:longitud_sat
            if k==1
                trabajo_factor_sq_sat{1, i}{z+1, 1}=trabajo_norm_sat{1, i}{1, z+1};
            end
            trabajo_factor_sq_sat{1, i}{z+1, k+1}=factor_load_sq_sat{1, i}(z, k);
        end
    end
end

```

```

xlswrite([strcat('Ponderacion/Analisis/Normal/Cargas factoriales saturadas',
    num2str(i)), '.xlsx'], trabajo_factor_sq_sat{1, i});

trabajo_comp_sat{1, i}{1, 1}=trabajo_norm_sat{1, i}{1, 1};

for z=2:size(trabajo_norm_sat{1, i}, 1)
    trabajo_comp_sat{1, i}{z, 1}=trabajo_norm_sat{1, i}{z, 1};
end

for k=2:(ancho_sat+1)
    trabajo_comp{1, i}{1, k}=strcat('Factor_', num2str(k-1));
    for z=2:size(trabajo_comp_sat{1, i}, 1)
        trabajo_comp_sat{1, i}{z, k}=0;
        for x=1:longitud_sat
            trabajo_comp_sat{1, i}{z, k}= trabajo_comp_sat{1, i}{z,
                k)+factor_load_sq_sat{1, i}(x, k-1)*trabajo_norm_sat{1, i}{z, x+1};
        end
    end
end

end

clear ancho_sat ancho longitud longitud_sat

%% Asignar pesos criterio experto agregación final

pesos_cat=[0.25 0.25 0.2 0.1 0.2];

%% Agregación

%Arreglar diferencia periodos temporales por categoría
long=size(trabajo_comp{1, 1}, 1);
periodo_cat=1;

for i=2:n_cat
    if long>size(trabajo_comp{1, i}, 1)
        long=size(trabajo_comp{1, i}, 1);
        periodo_cat=i;
    end
end

long_sat=size(trabajo_comp_sat{1, 1}, 1);
periodo_cat_sat=1;

for i=2:n_cat
    if long_sat>size(trabajo_comp_sat{1, i}, 1)
        long_sat=size(trabajo_comp_sat{1, i}, 1);
        periodo_cat_sat=i;
    end
end

%Emplazar títulos de las categorías
i=1;
for x=2:longitud_trabajo
    if isequaln(lista_trabajo{x, 2}, lista_trabajo{x-1, 2})
    else
        trabajo_agreg{1, i+1}=lista_trabajo{x, 2};
        trabajo_agreg_sat{1, i+1}=lista_trabajo{x, 2};
        i=i+1;
    end
end

trabajo_agreg{1, 1}='Fecha';
trabajo_agreg_sat{1, 1}='Fecha';

% Bucle subagregación
for i=1:n_cat

    longitud=size(trabajo_comp{1, i}, 1);
    ancho=size(trabajo_comp{1, i}, 2);
    x=2;

    for z=2:longitud
        trabajo_agreg{x, 1}=trabajo_comp{1, periodo_cat}{x, 1};
    end
end

```

```

if trabajo_agreg{x, 1}==trabajo_comp{1, i}{z, 1}
    trabajo_agreg{x, i+1}=0;
    if varianza_comp{1, i}(ancho-1, size(varianza_comp{1, i}, 2))>80
        for k=2:ancho
            trabajo_agreg{x, i+1} = trabajo_agreg{x, i+1}+varianza_comp{1, i}(k-1,
                (size(varianza_comp{1, i}, 2))-1)*trabajo_comp{1, i}{z,
                    k}/varianza_comp{1, i}(ancho-1, size(varianza_comp{1, i}, 2));
        end
    else
        for k=2:size(trabajo_norm{1, i}, 2)
            trabajo_agreg{x, i+1} = trabajo_agreg{x, i+1}+trabajo_norm{1, i}{z,
                k}/(size(trabajo_norm{1, i}, 2)-1);
        end
    end
    x=x+1;
end
end

longitud_sat=size(trabajo_comp{1, i}, 1);
ancho_sat=size(trabajo_comp{1, i}, 2);
x=2;
for z=2:longitud_sat
    trabajo_agreg_sat{x, 1}=trabajo_comp_sat{1, periodo_cat}{x, 1};
    if trabajo_agreg_sat{x, 1}==trabajo_comp_sat{1, i}{z, 1}
        trabajo_agreg_sat{x, i+1}=0;
        if varianza_comp_sat{1, i}(ancho_sat-1, size(varianza_comp_sat{1, i}, 2)) > 80
            for k=2:ancho
                trabajo_agreg_sat{x, i+1}= trabajo_agreg_sat{x, i+1}+varianza_comp_sat{1,
                    i}(k-1, (size(varianza_comp_sat{1, i}, 2))-1)*trabajo_comp_sat{1,
                        i}{z, k}/varianza_comp_sat{1, i}(ancho_sat-1,
                            size(varianza_comp_sat{1, i}, 2));
            end
        else
            for k=2:size(trabajo_norm_sat{1, i}, 2)
                trabajo_agreg_sat{x, i+1}= trabajo_agreg_sat{x, i+1}+trabajo_norm_sat{1,
                    i}{z, k}/(size(trabajo_norm_sat{1, i}, 2)-1);
            end
        end
        x=x+1;
    end
end
end

figure(5);
for k=2:size(trabajo_agreg, 2)
    plot(vt{1, periodo_cat}, cell2mat(trabajo_agreg(2:size(trabajo_agreg, 1), k)));
    hold on
end
legend('Desempleo', 'Empleo', 'Habilidades y capacitación', 'Protec. a los desempleados',
    'Salarios y costes lab.');
```

```

legend('Location', 'southoutside', 'NumColumns', 3);
title('Objetivos con indic. originales');
hold off
grid
print -depsc Graficas/original_areas

figure(6);
for k=2:size(trabajo_agreg_sat, 2)
    plot(vt{1, periodo_cat_sat}, cell2mat(trabajo_agreg_sat(2:size(trabajo_agreg_sat, 1),
        k)));
    hold on
end
legend('Desempleo', 'Empleo', 'Habilidades y capacitación', 'Protec. a los desempleados',
    'Salarios y costes lab.');
```

```

legend('Location', 'southoutside', 'NumColumns', 3);
title('Objetivos con indic. originales y normalización saturada');
hold off
grid
print -depsc Graficas/original_areas_saturado

xlswrite(strcat('Agregacion/Normal/Áreas del mercado de trabajo.xlsx'), trabajo_agreg);
xlswrite(strcat('Agregacion/Normal/Áreas del mercado de trabajo con normalización
    saturada.xlsx'), trabajo_agreg_sat);
```

```

% Bucle final agregación
trabajo_indice_equi_arit{1, 1}='Fecha';
trabajo_indice_exp_arit{1, 1}='Fecha';
trabajo_indice_sat_equi_arit{1, 1}='Fecha';
trabajo_indice_sat_exp_arit{1, 1}='Fecha';
trabajo_indice_equi_arit{1, 2}='Orig_equi_arit';
trabajo_indice_exp_arit{1, 2}='Orig_exp_arit';
trabajo_indice_sat_equi_arit{1, 2}='Orig_sat_equi_arit';
trabajo_indice_sat_exp_arit{1, 2}='Orig_sat_exp_arit';

trabajo_indice_equi_geom{1, 1}='Fecha';
trabajo_indice_exp_geom{1, 1}='Fecha';
trabajo_indice_sat_equi_geom{1, 1}='Fecha';
trabajo_indice_sat_exp_geom{1, 1}='Fecha';
trabajo_indice_equi_geom{1, 2}='Orig_equi_geom';
trabajo_indice_exp_geom{1, 2}='Orig_exp_geom';
trabajo_indice_sat_equi_geom{1, 2}='Orig_sat_equi_geom';
trabajo_indice_sat_exp_geom{1, 2}='Orig_sat_exp_geom';

for z=2:long
    trabajo_indice_equi_arit{z, 2}=0;
    trabajo_indice_exp_arit{z, 2}=0;
    trabajo_indice_equi_arit{z, 1}=trabajo_agreg{z, 1};
    trabajo_indice_exp_arit{z, 1}=trabajo_agreg{z, 1};
    trabajo_indice_equi_geom{z, 2}=1;
    trabajo_indice_exp_geom{z, 2}=1;
    trabajo_indice_equi_geom{z, 1}=trabajo_agreg{z, 1};
    trabajo_indice_exp_geom{z, 1}=trabajo_agreg{z, 1};

    for k=2:size(trabajo_agreg, 2)
        trabajo_indice_equi_arit{z, 2}= trabajo_indice_equi_arit{z, 2}+trabajo_agreg{z,
            k}/n_cat;
        trabajo_indice_exp_arit{z, 2}= trabajo_indice_exp_arit{z, 2}+trabajo_agreg{z,
            k}*pesos_cat(1, k-1);

        trabajo_indice_equi_geom{z, 2}= trabajo_indice_equi_geom{z, 2}*trabajo_agreg{z,
            k}^(1/n_cat);
        trabajo_indice_exp_geom{z, 2}= trabajo_indice_exp_geom{z, 2}*trabajo_agreg{z,
            k}*pesos_cat(1, k-1);

    end
end

for z=2:long_sat
    trabajo_indice_sat_equi_arit{z, 2}=0;
    trabajo_indice_sat_exp_arit{z, 2}=0;
    trabajo_indice_sat_equi_arit{z, 1}=trabajo_agreg_sat{z, 1};
    trabajo_indice_sat_exp_arit{z, 1}=trabajo_agreg_sat{z, 1};

    trabajo_indice_sat_equi_geom{z, 2}=1;
    trabajo_indice_sat_exp_geom{z, 2}=1;
    trabajo_indice_sat_equi_geom{z, 1}=trabajo_agreg_sat{z, 1};
    trabajo_indice_sat_exp_geom{z, 1}=trabajo_agreg_sat{z, 1};

    for k=2:size(trabajo_agreg_sat, 2)
        trabajo_indice_sat_equi_arit{z, 2}= trabajo_indice_sat_equi_arit{z,
            2}+trabajo_agreg_sat{z, k}/n_cat;
        trabajo_indice_sat_exp_arit{z, 2}= trabajo_indice_sat_exp_arit{z,
            2}+trabajo_agreg_sat{z, k}*pesos_cat(1, k-1);

        trabajo_indice_sat_equi_geom{z, 2}= trabajo_indice_sat_equi_geom{z,
            2}*trabajo_agreg_sat{z, k}^(1/n_cat);
        trabajo_indice_sat_exp_geom{z, 2}= trabajo_indice_sat_exp_geom{z,
            2}*trabajo_agreg_sat{z, k}*pesos_cat(1, k-1);

    end
end

figure(7);
plot(vt{1, periodo_cat}, cell2mat(trabajo_indice_equi_arit(2:size(trabajo_indice_equi_arit,
    1), 2)));
hold on;

```

```

plot(vt{1, periodo_cat}, cell2mat(trabajo_indice_exp_arit(2:size(trabajo_indice_exp_arit,
1), 2)));
hold on;
plot(vt{1, periodo_cat},
cell2mat(trabajo_indice_sat_equi_arit(2:size(trabajo_indice_sat_equi_arit, 1), 2)));
hold on;
plot(vt{1, periodo_cat},
cell2mat(trabajo_indice_sat_exp_arit(2:size(trabajo_indice_sat_exp_arit, 1), 2)));
legend('Equiponderación', 'Criterio experto', 'Equiponderación con norm. saturada',
'Criterio experto con norm. saturada');
legend('Location', 'southoutside', 'NumColumns', 2);
title('Sintético con indic. originales y agregación aritmética');
hold off
grid
print -depsc Graficas/original_sintético_aritmética

figure(8);
plot(vt{1, periodo_cat}, cell2mat(trabajo_indice_equi_geom(2:size(trabajo_indice_equi_geom,
1), 2)));
hold on;
plot(vt{1, periodo_cat}, cell2mat(trabajo_indice_exp_geom(2:size(trabajo_indice_exp_geom,
1), 2)));
hold on;
plot(vt{1, periodo_cat},
cell2mat(trabajo_indice_sat_equi_geom(2:size(trabajo_indice_sat_equi_geom, 1), 2)));
hold on;
plot(vt{1, periodo_cat},
cell2mat(trabajo_indice_sat_exp_geom(2:size(trabajo_indice_sat_exp_geom, 1), 2)));
legend('Equiponderación', 'Criterio experto', 'Equiponderación con norm. saturada',
'Criterio experto con norm. saturada');
legend('Location', 'southoutside', 'NumColumns', 2);
title('Sintético con indic. originales y agregación geométrica');
hold off
grid
print -depsc Graficas/original_sintético_geométrica

xlswrite(strcat('Agregacion/Normal/Indice equiponderado aritmético.xlsx'),
trabajo_indice_equi_arit);
xlswrite(strcat('Agregacion/Normal/Indice criterio experto aritmético.xlsx'),
trabajo_indice_exp_arit);
xlswrite(strcat('Agregacion/Normal/Indice norm_saturada equiponderado aritmético.xlsx'),
trabajo_indice_sat_equi_arit);
xlswrite(strcat('Agregacion/Normal/Indice norm_saturada criterio experto aritmético.xlsx'),
trabajo_indice_sat_exp_arit);

xlswrite(strcat('Agregacion/Normal/Indice equiponderado geométrico.xlsx'),
trabajo_indice_equi_geom);
xlswrite(strcat('Agregacion/Normal/Indice criterio experto geométrico.xlsx'),
trabajo_indice_exp_geom);
xlswrite(strcat('Agregacion/Normal/Indice norm_saturada equiponderado geométrico.xlsx'),
trabajo_indice_sat_equi_geom);
xlswrite(strcat('Agregacion/Normal/Indice norm_saturada criterio experto geométrico.xlsx'),
trabajo_indice_sat_exp_geom);

%% Análisis de robustez

robustez{1, 2}(:, 1)=[1;1;1;1];
robustez{1, 2}(:, 2)=[1;1;1;0];
robustez{1, 2}(:, 3)=[1;1;0;1];
robustez{1, 2}(:, 4)=[1;1;0;0];
robustez{1, 2}(:, 5)=[1;0;1;1];
robustez{1, 2}(:, 6)=[1;0;1;0];
robustez{1, 2}(:, 7)=[1;0;0;1];
robustez{1, 2}(:, 8)=[1;0;0;0];
robustez{1, 2}(:, 9)=[0;1;1;1];
robustez{1, 2}(:, 10)=[0;1;1;0];
robustez{1, 2}(:, 11)=[0;1;0;1];
robustez{1, 2}(:, 12)=[0;1;0;0];
robustez{1, 2}(:, 13)=[0;0;1;1];
robustez{1, 2}(:, 14)=[0;0;1;0];
robustez{1, 2}(:, 15)=[0;0;0;1];
robustez{1, 2}(:, 16)=[0;0;0;0];

```

```

for z=2:size(trabajo_indice_ref_exp_geom, 1)
    robustez{1, 1}(z-1, 1)=cell2mat(trabajo_indice_ref_exp_geom(z, 2));
    robustez{1, 1}(z-1, 2)=cell2mat(trabajo_indice_ref_exp_arit(z, 2));
    robustez{1, 1}(z-1, 3)=cell2mat(trabajo_indice_ref_equi_geom(z, 2));
    robustez{1, 1}(z-1, 4)=cell2mat(trabajo_indice_ref_equi_arit(z, 2));
    robustez{1, 1}(z-1, 5)=cell2mat(trabajo_indice_ref_sat_exp_geom(z, 2));
    robustez{1, 1}(z-1, 6)=cell2mat(trabajo_indice_ref_sat_exp_arit(z, 2));
    robustez{1, 1}(z-1, 7)=cell2mat(trabajo_indice_ref_sat_equi_geom(z, 2));
    robustez{1, 1}(z-1, 8)=cell2mat(trabajo_indice_ref_sat_equi_arit(z, 2));
    robustez{1, 1}(z-1, 9)=cell2mat(trabajo_indice_exp_geom(z, 2));
    robustez{1, 1}(z-1, 10)=cell2mat(trabajo_indice_exp_arit(z, 2));
    robustez{1, 1}(z-1, 11)=cell2mat(trabajo_indice_equi_geom(z, 2));
    robustez{1, 1}(z-1, 12)=cell2mat(trabajo_indice_equi_arit(z, 2));
    robustez{1, 1}(z-1, 13)=cell2mat(trabajo_indice_sat_exp_geom(z, 2));
    robustez{1, 1}(z-1, 14)=cell2mat(trabajo_indice_sat_exp_arit(z, 2));
    robustez{1, 1}(z-1, 15)=cell2mat(trabajo_indice_sat_equi_geom(z, 2));
    robustez{1, 1}(z-1, 16)=cell2mat(trabajo_indice_sat_equi_arit(z, 2));
    E_y(z-1, 1)=mean(robustez{1, 1}(z-1, :));
    E_y(z-1, 2)=var(robustez{1, 1}(z-1, :));
end

%X_1=1 (Mín-Máx)
%X_1=0 (Mín-Máx saturado)
for z=1:size(robustez{1, 1}, 1)
    a=0;
    b=0;
    c=0;
    d=0;
    for k=1:size(robustez{1, 1}, 2)
        if robustez{1, 2}(1, k)==1
            a=a+robustez{1, 1}(z, k);
            b=b+1;
        else
            c=c+robustez{1, 1}(z, k);
            d=d+1;
        end
    end
    E_y_x1(z, 1)=a/b; %Esperanza de Y con mín-máx
    E_y_x1(z, 2)=c/d; %Esperanza de Y con mín-máx saturado
    E_y_x1(z, 3)=var(E_y_x1(z, :));
    % V_y_x1(z, 1)=0;
    % V_y_x1(z, 2)=0;
    %
    % for k=1:size(robustez{1, 1}, 2)
    %     if robustez{1, 2}(1, k)==1
    %         V_y_x1(z, 1)=V_y_x1(z, 1)+(robustez{1, 1}(z, k)-E_y_x1(z, 1))^2;
    %     else
    %         V_y_x1(z, 2)=V_y_x1(z, 2)+(robustez{1, 1}(z, k)-E_y_x1(z, 2))^2;
    %     end
    % end
    % V_y_x1(z, 1)=V_y_x1(z, 1)/b;
    % V_y_x1(z, 2)=V_y_x1(z, 2)/d;
    % V_y_x1(z, 3)=var(E_y_x1(z, :));
end

s_1=E_y_x1(:, 3)./E_y(:, 2);
for z=2:size(trabajo_agreg_ref, 1)
    for k=2:size(trabajo_agreg_ref, 2)
        comparacion_areas_norm(z-1, k-1)=trabajo_agreg_ref{z, k}-trabajo_agreg_ref_sat{z, k};
    end
    variacion(z-1, 1)=mean(robustez{1, 1}(z-1, :));
    variacion(z-1, 2)=median(robustez{1, 1}(z-1, :));
    variacion(z-1, 3)=quantile(robustez{1, 1}(z-1, :), 0.25);
    variacion(z-1, 4)=quantile(robustez{1, 1}(z-1, :), 0.75);
    variacion(z-1, 5)=max(robustez{1, 1}(z-1, :));
    variacion(z-1, 6)=min(robustez{1, 1}(z-1, :));
    variacion(z-1, 7)=var(robustez{1, 1}(z-1, :));
end
figure(9);
boxchart(transpose(flip(robustez{1, 1}, 1)));
title('Salud del mercado laboral español');
print -depsc Graficas/caja_bigotes
hold off

figure(10);

```

```

boxchart(transpose(flip(robustez{1, 1}, 1)));
title('Salud del mercado laboral español');
hold on
plot([flip(variacion(:, 2), 1), flip(variacion(:, 3), 1), flip(variacion(:, 4), 1)],
      'LineStyle', '--', 'LineWidth', 1);
hold on
plot(flip(cell2mat(trabajo_indice_ref_exp_geom(2:size(trabajo_indice_ref_exp_geom, 1), 2)),
      1), 'k', 'LineWidth', 2);
legend('Distribución', 'Mediana', 'Q_3', 'Q_1', 'Indicador sintético elaborado');
legend('Location', 'southoutside', 'NumColumns', 3);
hold off

```

Código 6. Ponderación y agregación del sintético final junto al análisis de robustez

En el código 6 aparecen también las diferentes variaciones realizadas como parte del análisis de robustez sobre el resultado final.