

Esplugues Conca, I. (2026). Análisis comparativo de técnicas de representación latente para entidades territoriales: caso de estudio en la Comunidad de Madrid. *GeoFocus, Revista Internacional de Ciencia y Tecnología de la Información Geográfica* (Artículos), 37, 63-85. <https://dx.doi.org/10.21138/GF.946>

ANÁLISIS COMPARATIVO DE TÉCNICAS DE REPRESENTACIÓN LATENTE PARA ENTIDADES TERRITORIALES: CASO DE ESTUDIO EN LA COMUNIDAD DE MADRID

Ignacio Esplugues Conca 
ignacioespluguesconca@gmail.com

RESUMEN

El presente trabajo propone un enfoque innovador para la representación vectorial de municipios mediante técnicas de *graph embeddings*, aplicado al caso de la Comunidad de Madrid. La investigación parte del contexto de desigualdad territorial y reconfiguración socioeconómica que caracteriza al área metropolitana madrileña, y plantea como hipótesis que las relaciones entre municipios pueden aprenderse en un espacio latente capaz de preservar su estructura urbana, económica y temporal. Para ello, se construye un grafo urbano en el que cada municipio constituye un nodo y las aristas representan interacciones de diversa naturaleza (distancias, colindancias, flujos migratorios o conectividad). A partir de un conjunto heterogéneo de fuentes públicas se desarrollan distintos modelos basados en redes neuronales que integran información nodal, internodal y temporal. La validación de los *embeddings* combina métricas intrínsecas de coherencia y tareas predictivas externas, contrastadas mediante tests estadísticos (Friedman, Nemenyi y Wilcoxon-Holm) y *bootstrap* paramétrico. Los resultados evidencian diferencias significativas entre modelos, destacando E2A-SAGE-MAE como el de mejor desempeño y mayor robustez, alcanzando el mejor rendimiento en torno al 80 % de las configuraciones analizadas. Asimismo, se confirma que la ciudad de Madrid actúa como *outlier*, afectando de forma notable la estructura latente del grafo. En conjunto, los resultados demuestran que es posible generar representaciones vectoriales de unidades urbanas que capturen relaciones complejas, ofreciendo una herramienta útil para el análisis territorial.

Palabras clave: *embeddings*; grafo urbano; analítica urbana

COMPARATIVE ANALYSIS OF LATENT REPRESENTATION TECHNIQUES FOR TERRITORIAL ENTITIES: CASE STUDY IN THE COMMUNITY OF MADRID

ABSTRACT

This paper proposes an innovative approach to the vector representation of municipalities using *graph embedding* techniques, applied to the case of the Community of Madrid. The research is based on the context of territorial inequality and socio-economic reconfiguration that

Recepción: 28/10/2025

Editor al cargo: Dr. Ismael Vallejo

Attribution-NonCommercial-NoDerivatives 4.0 International (CC BY-NC-ND 4.0)

Aceptación definitiva: 27/07/2026

www.geofocus.org

characterises the Madrid metropolitan area, and hypothesises that the relationships between municipalities can be learned in a latent space capable of preserving their urban, economic and temporal structure. To this end, an urban graph is constructed in which each municipality constitutes a node and the edges represent interactions of various kinds (distances, boundaries, migratory flows or connectivity). Based on a heterogeneous set of public sources, different neural networks models are developed which integrate nodal, internodal and temporal information. The validation of the *embeddings* combines intrinsic consistency metrics and external predictive tasks, verified by statistical tests (Friedman, Nemenyi and Wilcoxon-Holm) and parametric bootstrap. The results show significant differences between models, with E2A-SAGE-MAE standing out as the best performing and most robust, achieving the best performance in around 80 % of the configurations analysed. Furthermore, it is confirmed that the city of Madrid acts as an outlier, significantly affecting the latent structure of the graph. Overall, the results demonstrate that it is possible to generate vector representations of urban units that capture complex relationships, providing a useful tool for territorial analysis.

Keywords: *embeddings*; urban graph; urban analytics

1. Introducción

1.1. Contexto

La Comunidad de Madrid ha experimentado una transformación urbana a lo largo de las últimas décadas, marcada por procesos de expansión territorial, reconfiguración socioeconómica y tensiones en el acceso a la vivienda. Desde la implementación del Plan de Ensanche en el siglo XIX (Castro 1862) hasta la formación de la actual área metropolitana, la estructura urbana madrileña refleja una constante negociación entre crecimiento demográfico, modernización y desigualdad.

En las últimas dos décadas el mercado inmobiliario de Madrid ha sido escenario de una escalada de precios sin precedentes, con un incremento acumulado superior al 100 % desde 2015, pasando de 2691 €/m² en enero de 2015 a 5104 €/m² en enero de 2025 (Idealista 2025a), situando el precio medio de la vivienda en torno a 5500 €/m² en junio de 2025 (Idealista 2025a). Esta evolución ha llevado a que se reaviven los temores de una nueva burbuja inmobiliaria, comparable en intensidad a la del año 2007. A este fenómeno se le suma una progresiva tensión en el mercado de alquiler, donde la subida interanual ha acumulado un 40 % en los últimos cinco ejercicios, pasando de 15.8 €/m² en agosto de 2020 a 22 €/m² en junio de 2025 (Idealista 2025b), generando un entorno residencial inaccesible para algunas capas de la población.

La escasez de vivienda asequible y la presión de agentes especulativos han favorecido procesos de gentrificación en barrios céntricos (Lees *et al.* 2013) y desplazamientos hacia la periferia, dando lugar a una horda de migraciones internas que principalmente afectan a jóvenes, migrantes y colectivos vulnerables. Los datos muestran cómo las migraciones entre municipios de la comunidad, además de estar en claro aumento, representan en media el 63 % anual de migraciones con destino la Comunidad de Madrid (Figura 1).

A nivel global, Madrid no es un caso aislado: fenómenos similares se observan en otras capitales europeas, donde el mercado de la vivienda se enfrenta a un déficit estructural y a un desequilibrio entre oferta y demanda, en parte permitido por la turistificación (Wachsmuth & Weiler 2018) y la financiarización del suelo urbano (Aalbers 2016, Fernández & Aalbers 2016). Estas dinámicas ponen de manifiesto la importancia de contar con herramientas analíticas capaces de identificar patrones comunes y divergentes en la evolución urbana, tanto en un contexto internacional como dentro de la propia Comunidad de Madrid.

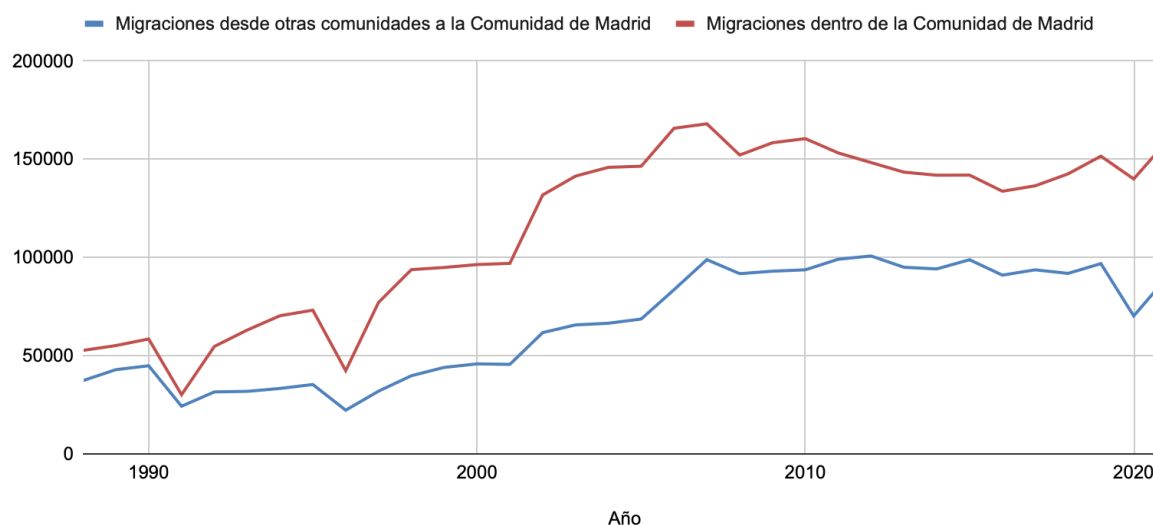


Figura 1. Migraciones desde otras comunidades con destino la Comunidad de Madrid vs. Migraciones dentro de la Comunidad de Madrid

Fuente: Instituto de Estadística de la Comunidad de Madrid (IECM)

De esta manera, el modelo se concibe como una herramienta para fortalecer la planificación regional y la articulación de políticas públicas en un contexto caracterizado por procesos de fragmentación, suburbanización y creciente interdependencia funcional entre municipios. El análisis de estas dinámicas, desde una perspectiva crítica y basada en datos, sitúa a la Comunidad de Madrid como un laboratorio idóneo para explorar cómo las nuevas formas de representación urbana pueden apoyar estrategias de gobernanza más inclusivas, resilientes y adaptadas a las territorialidades emergentes.

A partir de estas consideraciones, planteamos la hipótesis de que los municipios de la Comunidad de Madrid pueden representarse en un espacio vectorial latente que preserve sus relaciones socioeconómicas y urbanas. Dicho espacio no es observable directamente en los datos originales, sino que se aprende mediante técnicas de *machine learning*. Su objetivo es condensar en pocas dimensiones la información relevante de cada municipio, de modo que las distancias y posiciones entre vectores reflejen similitudes en términos socioeconómicos, urbanos o demográficos.

Además, esta aproximación no se limita exclusivamente al caso madrileño: podría aplicarse tanto a los municipios de la Comunidad de Madrid, como a otras unidades de población, incluyendo niveles más desagregados, como barrios, así como escalas mayores, como comunidades autónomas o países, siempre que existan contextos socioeconómicos comparables en los que se observe una evolución conjunta de precios de la vivienda, población y otros factores.

1.2. Formulación del problema

El problema central de esta investigación radica en cómo modelar relaciones urbanas complejas de forma que se preserve la información relevante y permita su análisis cuantitativo. Se propone representar la Comunidad de Madrid mediante un grafo urbano, en el que cada nodo equivalga a un municipio y las aristas que los unan a sus respectivas relaciones.

Para ello, resultará necesario establecer un conjunto heterogéneo de características asociadas a los municipios. Dichas variables han de contemplar, por un lado, atributos estáticos descriptivo (como la superficie o la altitud); por otro, dimensiones de naturaleza temporal que recojan la evolución de determinados indicadores a lo largo de los años (por ejemplo, la variación del precio de la vivienda o la densidad de población a lo largo del tiempo); y, adicionalmente, variables relacionales que reflejen las interacciones entre municipios (como la colindancia territorial, las distancias por carretera o en línea recta, o los flujos migratorios intermunicipales).

A partir de este grafo, el objetivo técnico es generar representaciones vectoriales, también llamadas *embeddings* (Zhao *et al.* 2023), que capturen la estructura latente de las relaciones entre municipios. Esto se realizará a partir de modelos de redes neuronales. Estas representaciones facilitarán no solo la visualización y análisis exploratorio de patrones espaciales y temporales, sino también la integración con otros modelos predictivos orientados a distintos casos de uso, como la segmentación de mercados inmobiliarios o la detección temprana de fenómenos de gentrificación. Asimismo, se plantea como objetivo complementario la validación de la calidad de estas representaciones mediante métricas cuantitativas adecuadas, evaluando su capacidad de predicción, tanto de atributos y relaciones presentes en el grafo original, como de variables desconocidas para el mismo.

Por último, el trabajo contempla también objetivos de carácter social y aplicado, proponiendo casos de uso en los que se podrían aplicar las representaciones generadas, como por ejemplo en la toma de decisiones en el ámbito de la política urbana, permitiendo detectar zonas en riesgo de gentrificación o vulnerabilidad.

1.3. Justificación

El enfoque propuesto combina métodos avanzados de análisis de grafos y aprendizaje automático para abordar un problema con futuras implicaciones prácticas. Este trabajo se justifica tanto por su valor científico, al aportar una metodología replicable en otros contextos, como por su relevancia aplicada, al facilitar el análisis de mercados inmobiliarios, la planificación de infraestructuras y el diseño de políticas orientadas a mitigar desigualdades territoriales.

El desarrollo de un aplicativo con estas características permitiría a entidades públicas o privadas disponer de una herramienta para realizar análisis espaciales y temporales robustos, mejorando la capacidad de anticipación en entornos urbanos complejos. Algunos ejemplos de ciudades y proyectos en las que se podrían aplicar este tipo de análisis son:

- **Helsinki Digital Twin** (City of Helsinki, s. f.) se basa en el desarrollo de modelos 3D de la ciudad como una representación virtual del entorno, el funcionamiento y las circunstancias cambiantes de la ciudad. Se basan en una combinación de servicios de tecnología de la información, datos abiertos e información constantemente actualizada.
- **Amsterdam InChange** (Amsterdam Smart City, s. f.) es una iniciativa colaborativa con el objetivo de transformar Ámsterdam en una ciudad más eficiente, sostenible e inclusiva, utilizando la tecnología, los datos y la participación ciudadana.
- **Singapore Virtual** (OECD-OPSI, 2024) supone una representación 3D de toda la nación insular, generando así un gemelo digital que incluye elementos de la superficie, como edificios, carreteras y zonas verdes, así como infraestructuras subterráneas.
- **Smart City Madrid** (Ciudad de Madrid, s. f.) impulsa proyectos orientados a la gestión eficiente de los recursos urbanos, la movilidad sostenible y la participación ciudadana mediante el uso de *big data*, sensores IoT y plataformas de gestión en tiempo real.

En este contexto, los *embeddings* propuestos no pretenden sustituir la información geográfica existente en estas plataformas, sino complementarla como una nueva capa analítica. Cada municipio (o unidad territorial) podría incorporar un vector latente asociado que resumiera sus características socioeconómicas, funcionales y espaciales. Dichos vectores permitirían realizar consultas por similitud territorial, detectar municipios anómalos, identificar áreas funcionalmente equivalentes aunque no sean geográficamente contiguas, así como alimentar modelos predictivos dentro de los propios gemelos digitales. De este modo, el embedding actuaría como una representación sintética del comportamiento territorial que enriquecería los modelos espaciales ya existentes.

2. Metodología

2.1. Antecedentes

El análisis espacial urbano ha evolucionado desde los Sistemas de Información Geográfica (SIG) clásicos (Gutiérrez Puebla & Rubio Barroso 1997), que se centraban en la cartografía y el análisis descriptivo, hasta modelos predictivos y representaciones funcionales de las ciudades. En los últimos años, esta evolución ha incorporado enfoques basados en *embeddings*, aprovechando avances en redes neuronales y grafos para modelar relaciones espaciales más profundas.

Estos enfoques, conocidos como *graph embeddings*, poseen el objetivo principal de capturar patrones tanto locales como globales de la estructura de un grafo. Dependiendo de la técnica, los *embeddings* pueden aprenderse de forma transductiva, donde los vectores sólo son válidos para los nodos presentes en el entrenamiento, o inductiva, donde el modelo generaliza a nodos no vistos previamente. Los métodos de generación de *embeddings* pueden clasificarse en tres grandes categorías: basados en recorridos aleatorios, en factorización matricial y en aprendizaje profundo.

En primer lugar, dentro de los métodos **basados en recorridos aleatorios** aparecen técnicas como DeepWalk y Node2Vec (Qiu *et al.* 2017), que realizan exploraciones locales del grafo y aplican modelos de lenguaje para aprender vectores que codifican la coocurrencia de nodos en trayectorias. DeepWalk utiliza recorridos aleatorios simples y uniformes para generar secuencias de nodos, que luego son tratadas como “frases” sobre las que aplica Word2Vec (Goldberg & Levy 2014) para aprender representaciones vectoriales que capturan la proximidad en el grafo. Node2Vec, en cambio, introduce recorridos aleatorios sesgados mediante parámetros que controlan la probabilidad de explorar vecinos cercanos o alejarse en el grafo. Esta flexibilidad le permite capturar tanto la homofilia (nodos similares conectados) como roles estructurales (nodos con funciones análogas en distintas partes del grafo).

Por otro lado, en el ámbito de la **factorización matricial**, métodos como *Large-scale Information Network Embedding* (LINE) (Qiu *et al.* 2017) o *High-Order Proximity preserved Embedding* (HOPE) (Ou *et al.* 2016) buscan preservar medidas de proximidad globales (por ejemplo, la similitud de primer o segundo orden) mediante la descomposición de matrices derivadas de la adyacencia. La similitud de primer orden se refiere a la relación directa entre dos nodos conectados por una arista, de modo que nodos con enlaces fuertes o frecuentes tienden a representarse próximos en el espacio latente. Por otro lado, la similitud de segundo orden captura la semejanza estructural en los patrones de conexión de los nodos, es decir, dos nodos pueden considerarse similares aunque no estén directamente conectados si comparten vecinos o distribuciones de enlaces análogas. De este modo, los *embeddings* no solo reflejan conexiones inmediatas, sino también la homogeneidad de su contexto.

Por último, dentro de los métodos **basados en aprendizaje profundo**, algoritmos como *Graph Convolutional Network* (GCN) (Wu *et al.* 2021), *Graph Attention Network* (GAT) (Veličković *et al.* 2017) o *Graph AutoEncoder* (GAE) (Kipf & Welling 2016) aprenden representaciones latentes mediante el paso de mensajes entre nodos, en el que cada nodo actualiza su *embedding* agregando la información de sus vecinos y propagándola a través de múltiples capas. Este mecanismo permite capturar tanto las características propias de cada nodo como el contexto local definido por su vecindad. En particular, los GCN emplean operadores convolucionales sobre el grafo basados en el espectro del Laplaciano, los GAT introducen mecanismos de atención que ponderan de forma diferenciada la importancia de cada vecino en la agregación, y los GAE utilizan arquitecturas *encoder-decoder* para reconstruir la estructura del grafo a partir de representaciones comprimidas.

Adicionalmente, como incorporación más reciente al estado del arte, destacan los *Graph Transformers* (Shehzad *et al.* 2024), que trasladan el paradigma de los *Transformers* al dominio de los grafos. En lugar de procesar secuencias lineales, aplican mecanismos de *self-attention* directamente entre los nodos, lo que permite modelar dependencias globales más allá de la vecindad inmediata. Para ello, se complementan con codificaciones posicionales o estructurales (basadas, por ejemplo, en el espectro del Laplaciano, en distancias de camino más corto o en distribuciones derivadas de recorridos aleatorios), que incorporan explícitamente información sobre la topología del grafo y facilitan que el modelo distinga la posición relativa de los nodos dentro de la red.

En lo que respecta a la sinergia entre la representación urbana y las técnicas mencionadas, Barthélemy (2011) sentó las bases conceptuales al mostrar que una ciudad puede ser representada como una red compleja donde barrios son nodos y conexiones como infraestructuras o flujos son aristas. Introdujo métricas fundamentales de centralidad, conectividad y patrones estructurales que sirven como referencia para construir grafos urbanos formales.

Años más tarde (Yao *et al.* 2018) se introdujo Region2Vec, un modelo que aprende *embeddings* de regiones combinando grafos de movilidad en taxi con redes neuronales profundas, permitiendo capturar patrones espaciotemporales en la demanda de transporte. Esto es especialmente relevante, ya que demuestra que las representaciones vectoriales de áreas urbanas pueden originarse de variables reales. Un enfoque similar pero con una dimensión multimodal más pronunciada fue el expuesto en Urban2Vec (Wang *et al.* 2020), que incorporaba datos visuales obtenidos de imágenes de *street view* junto con atributos funcionales de los barrios, como la presencia de comercios, áreas verdes o instalaciones educativas.

De forma complementaria, Zhang *et al.* (2017) propusieron un modelo que integra características sociodemográficas, movilidad y servicios, usando redes residuales para aprender representaciones latentes que predicen flujos de personas entre barrios. Si bien es un enfoque eficaz en contextos de flujo, este no ha sido aplicado con un objetivo general de representación de unidades urbanas.

Además, se realizan estudios exhaustivos de *Graph Neural Networks* (GNN) (Wu *et al.* 2021) y se destacan aplicaciones relevantes en ciudades inteligentes, aunque señalando limitaciones como falta de interpretabilidad y sensibilidad a la calidad del grafo inicial. Estas limitaciones fortalecen la necesidad de implementar en nuestro caso mecanismos de validación cuantitativa.

Adicionalmente, añadiendo un factor espacio-temporal, surgieron las *Spatio-Temporal Graph Convolutional Networks* (ST-GCN) (Al Sahili & Awad 2023) y las *Graph Convolutional Network for Location-to-Vector* (GCN-L2V) (Tian *et al.* 2021), las cuales permiten modelar entornos dinámicos en los que, además de la estructura espacial del grafo, es necesario capturar la evolución temporal y los patrones de movilidad. Mientras que las ST-GCN representan los datos como secuencias de grafos o grafos con atributos temporales, combinando convoluciones espaciales y temporales para aprender cómo cambian las relaciones entre nodos a lo largo del tiempo, GCN-L2V integra información geográfica y funcional mediante un grafo dual construido a partir de proximidades y flujos de transporte, generando *embeddings* que reflejan simultáneamente la localización y el rol estructural de cada nodo en la red urbana.

Todas estas técnicas comparten la idea central de que las relaciones espaciales, funcionales o temporales entre municipios pueden ser representadas de forma explícita mediante grafos, y que los *embeddings* generados a partir de éstas son una herramienta poderosa para analizar fenómenos urbanos complejos. Sin embargo, la mayoría de los trabajos previos se han focalizado en dominios muy específicos como la movilidad o los servicios urbanos. Nuestro enfoque introduce una innovación al aplicar estos modelos a un nivel municipal completo, tomando como base la Comunidad de Madrid y considerando un conjunto heterogéneo de variables: desde características sociodemográficas y económicas hasta indicadores de servicios, movilidad y conectividad territorial.

De este modo, la aportación no reside únicamente en trasladar metodologías consolidadas al contexto regional, sino también en proponer una representación vectorial de municipios que integra relaciones de distinta índole (estáticas, temporales, matriciales direccionadas y no direccionadas). Además, se incorpora un proceso de validación cuantitativa de los *embeddings* para contrastar si las representaciones latentes capturan adecuadamente la estructura socioeconómica y espacial del territorio. En conjunto, este planteamiento ofrece una perspectiva novedosa para el análisis urbano-regional y abre la puerta a aplicaciones en gobernanza territorial y planificación inmobiliaria que van más allá de los enfoques ya expuestos en la literatura.

2.2. Datos, preprocesamiento y modelado del grafo

El éxito del enfoque que proponemos dependerá en gran medida de la cantidad y diversidad de los datos disponibles, así como de la calidad de los mecanismos de fusión empleados, ya que a mayor riqueza y heterogeneidad en las fuentes, más robusta y expresiva será la representación obtenida.

Los datos demográficos, tales como la distribución de edad, nivel educativo o densidad de población, se emplean habitualmente como atributos nodales en los grafos, ofreciendo información sobre la estructura social de cada municipio. Los datos de movilidad, obtenidos a partir de trayectorias GPS o flujos de tráfico, son fundamentales para definir aristas ponderadas que capturen la intensidad de las interacciones entre municipios. Por otro lado, los indicadores económicos, como renta per cápita, tasas de empleo o tipología de uso del suelo, aportan contexto sobre las dinámicas productivas locales. La disponibilidad de información sobre servicios públicos y puntos de interés, como número de hospitales, centros educativos o estaciones de transporte, enriquece aún más la representación de las regiones, permitiendo modelar la accesibilidad y la funcionalidad de los territorios. En algunos casos, se incorporan imágenes satelitales o de *street view*, que buscan capturar también el aspecto físico del espacio urbano.

Con el propósito de conformar una representación heterogénea de cada municipio, la información utilizada se deberá encontrar agregada a nivel municipal y anual. Para este fin se han empleado exclusivamente fuentes de datos de carácter público:

- El **Instituto de Estadística de la Comunidad de Madrid** (IECM), a través de sus bancos de datos BANVI y ALMUDENA, constituye una referencia esencial para el estudio del territorio madrileño. El Banco de Vivienda (BANVI) forma parte del Banco de Datos Estructurales y recoge información detallada sobre el parque de viviendas de la región, mientras que el Banco de Datos Municipal y Zonal (ALMUDENA) ofrece una base multitemática a nivel municipal que permite acceder a información sociodemográfica, económica y territorial.
- El **Instituto Nacional de Estadística** (INE) es el organismo encargado de la elaboración y difusión de estadísticas oficiales de carácter demográfico, económico y social en todo el territorio español, aportando un marco de referencia para el análisis a diferentes escalas.
- El **Centro Nacional de Información Geográfica** (CNIG) gestiona y difunde la información geoespacial oficial en España, proporcionando cartografía, datos topográficos y recursos geográficos esenciales para los análisis territoriales y espaciales.
- El **Consorcio Regional de Transportes de Madrid** (CRTM) coordina y gestiona el sistema de transporte público en la Comunidad de Madrid, ofreciendo información sobre la red de transportes y la movilidad regional, lo que permite valorar la accesibilidad y conectividad de los municipios.

La selección de variables no se realizó de forma arbitraria, sino siguiendo tres criterios metodológicos. En primer lugar, se priorizaron variables ampliamente utilizadas en la literatura de analítica urbana y representación territorial, por haber demostrado capacidad para caracterizar la estructura funcional de las ciudades. En segundo lugar, únicamente se incorporaron variables disponibles de forma homogénea para todos los municipios de la Comunidad de Madrid y con suficiente continuidad temporal, garantizando así la comparabilidad espacial y temporal. Finalmente, se buscó construir una representación multimodal del territorio integrando información de naturaleza diversa (sociodemográfica, económica, movilidad, conectividad y características físicas), con el objetivo de maximizar la capacidad descriptiva del grafo sin introducir redundancias excesivas.

Sin embargo, estos enfoques plantean retos significativos, como la necesidad de desarrollar mecanismos de fusión que integren datos heterogéneos, la dificultad de escalar los modelos a grandes áreas urbanas o cuestiones como la interpretabilidad de los *embeddings*. Por este motivo, resulta necesario establecer un conjunto de metodologías de fusión, escalado y validación que permitan homogeneizar la información en función de los requisitos del modelo a entrenar.

Por último, se debe definir la estructura del grafo urbano, en la que los municipios serán los nodos a representar vectorialmente y las variables nodales e internodales (aristas), las características que lo van a definir. En la Figura 2 se observa una posible representación de la relación entre un par de nodos en un grafo multimodal.

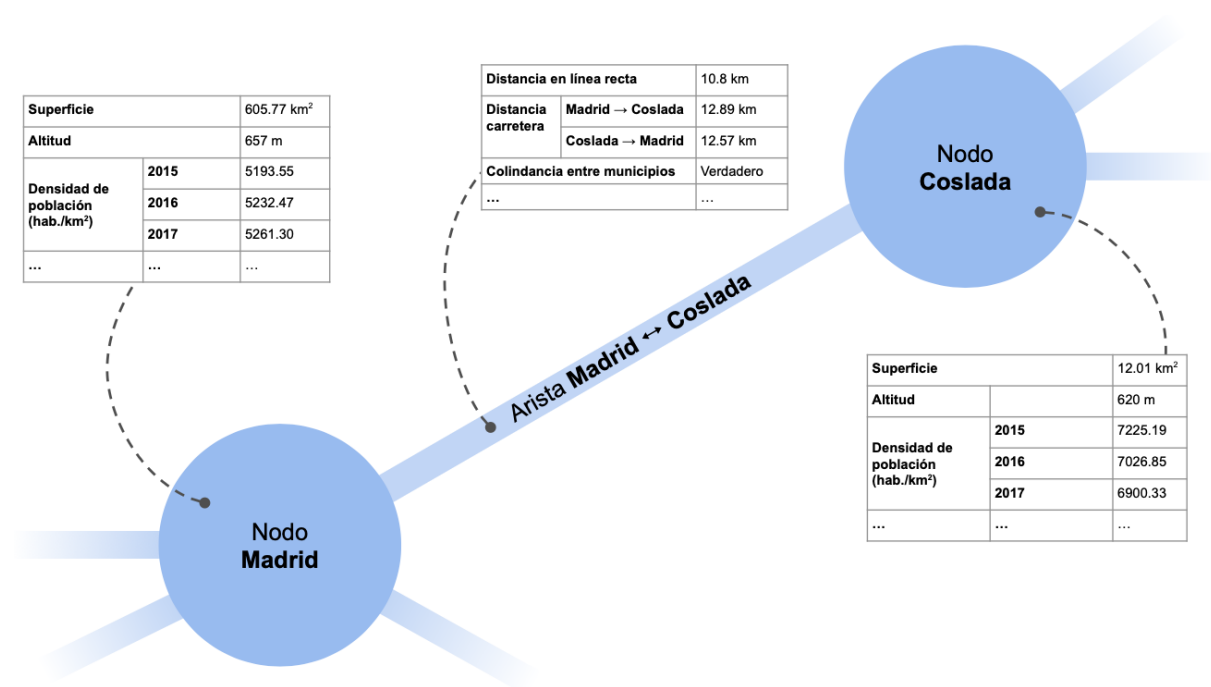


Figura 2. Representación de dos nodos (Madrid y Coslada) en un grafo urbano

2.3. Espacio latente

El espacio latente resultante de un *graph embedding* es una proyección en la que cada municipio se representa como un vector. La cercanía entre dos nodos en este espacio implica una similitud estructural o semántica según los criterios aprendidos por el modelo. Por ejemplo, dos municipios cercanos pueden compartir características sociodemográficas o estar fuertemente conectados en términos de movilidad. La noción de cercanía entre dos representaciones vectoriales puede evaluarse mediante distintas medidas de distancia, entre las cuales este trabajo emplea las siguientes:

- **Distancia Euclídea:** la cual mide la longitud del segmento de línea recta que une los dos puntos en un espacio d-dimensional. Es la métrica más utilizada en espacios vectoriales por su simplicidad e interpretación geométrica directa. Esta métrica es simétrica, satisface la desigualdad triangular y depende de la escala de las variables, por lo que puede ser necesario normalizar previamente los datos.
- **Distancia Coseno:** la cual no mide magnitudes absolutas, sino el ángulo entre dos vectores, capturando su direccionalidad relativa. Esta métrica es especialmente útil cuando interesa comparar la orientación de los vectores, ignorando su módulo.

2.4. Exploración de modelos

Una vez abordadas las tareas de recolección, preprocesado y fusión de datos, se procede a la exploración de posibles modelos de generación de *embeddings* de municipios. Para todos los casos, se realizará el análisis incluyendo y excluyendo la ciudad de Madrid dentro del entrenamiento, ya que ésta se presenta como posible valor atípico (también conocido como *outlier*) global dentro del conjunto de municipios, al ser la capital y concentrar la mayoría de población, servicios y actividad económica de la comunidad.

A su vez, se hablará de “año objetivo” como el año sobre el que se representen los *embeddings* de cada uno de los municipios (esto solo aplica a variables con carácter temporal). Una vez se realicen dos entrenamientos (con y sin la ciudad de Madrid) para cada tipología de modelo, se compararán, validarán y se seleccionará cuál de ellos es el que genera una mejor representación de los municipios.

2.4.1. GTMAE

Dentro de los modelos de aprendizaje profundo, se van a estudiar varias arquitecturas, la mayoría de ellas basadas en *autoencoders*. El primero de ellos es el **Graph Transformer-based Multitask AutoEncoder** (GTMAE) (Singh 2025). Este se basa en un *autoencoder* con un *encoder* de tipo *Transformer* sobre grafos (*edge-aware*) y con dos *decoders*: uno de aristas (regresión de atributos de arista) y otro de nodos (reconstrucción de atributos nodales).

En lo que a preprocesado de datos respecta, el punto de partida es definir el universo de nodos como la unión de todos los municipios que aparecen en las cuatro fuentes: variables nodales estáticas y temporales y variables matriciales dirigidas y no dirigidas. Una vez fijado ese conjunto, se ordena y se establece una codificación nodo-índice para poder trabajar con representaciones tensoriales.

La matriz de atributos de cada nodo se compone combinando, por columnas, un bloque estático y un bloque temporal. En el bloque estático se alinean los atributos tabulares estandarizando mediante *z-score*. En el bloque temporal se toma la foto del año objetivo y se aplica la misma estandarización. La concatenación de ambos bloques garantiza que cada nodo tenga una fila completa.

El conjunto de aristas se construye a partir de las relaciones dirigidas y no dirigidas. Para ello, todas las aristas se representan como dirigidas, pero se incorpora un atributo adicional que permite distinguir entre aquellas que provienen de variables intrínsecamente dirigidas y las que provienen de relaciones no dirigidas. En el caso de las primeras, cada par ordenado (u, v) y (v, u) se mantiene como arista distinta, respetando la posible asimetría de sus atributos. En el caso de las segundas, cada par de nodos se duplica explícitamente en ambas direcciones (u, v) y (v, u) , de modo que el grafo resultante es completamente dirigido, pero conserva la información de que se trata de una relación simétrica.

Este diseño permite posteriormente realizar un particionado de aristas específico para la tarea de predicción de enlaces. En particular, para las aristas dirigidas, cada dirección se considera de manera independiente y puede caer en diferentes particiones de entrenamiento, validación o prueba. Sin embargo, para las aristas no dirigidas se realiza un particionado agrupado: el par (u, v) y su inverso (v, u) siempre se asignan al mismo conjunto, evitando fugas de información entre splits. De esta forma se garantiza que el modelo no pueda memorizar la existencia de una arista en una dirección y beneficiarse de ella en la validación o el test en la dirección contraria, lo que reforzaría un sesgo no deseado. Con este esquema, se obtiene un grafo capaz de respetar la naturaleza diferenciada de las relaciones que lo componen y de ofrecer un entrenamiento y evaluación más coherentes. Con estos elementos se obtiene una representación homogénea del grafo formada por:

- $X \in \mathbb{R}^{N \times F}$ es la matriz de atributos nodales.
- $E_{idx} \in \mathbb{N}^{2 \times |E|}$ es el listado de aristas como pares de índices.
- $E_{attr} \in \mathbb{R}^{|E| \times K}$ es la matriz de atributos de arista, alineada con el orden de las aristas.

En lo que respecta al modelo, se divide en tres piezas, un *encoder* y dos *decoders*:

- El *encoder* aplica dos capas *Transformer* que realizan propagación de mensajes con atención e incorporan explícitamente los atributos de arista. En la primera capa, cada nodo agrega la información de sus vecinos ponderada por coeficientes de atención que dependen tanto de las representaciones de los nodos como de los vectores de atributos que los conectan; el resultado pasa por ReLU y *dropout*. La segunda capa vuelve a propagar mensajes (también *edge-aware*) y proyecta a la dimensión latente deseada. El resultado de la segunda capa es el *embedding* de los nodos.
- El *decoder* de aristas toma un conjunto de pares de nodos y calcula una puntuación para cada uno. Para ello, no solo utiliza los *embeddings* de los nodos, sino que también añade información explícita sobre cómo se relacionan entre sí esos nodos. Esta información adicional se llama atributos de par densos y consiste en medidas continuas que reflejan el grado de similitud o proximidad entre ambos nodos. En este caso se utiliza la similitud coseno, que indica lo parecidos que son sus vectores de representación. De este modo, el modelo combina la

representación de cada nodo con una medida cuantitativa de su relación, generando un vector final que intenta reconstruir los atributos de la arista.

- El *decoder* de nodos únicamente toma los *embeddings* de los nodos y trata de reconstruir un subconjunto de atributos nodales seleccionados.

La función de pérdida (en adelante, LOSS) es multitarea y combina una pérdida Huber de regresión de aristas, otra pérdida Huber de reconstrucción nodal sobre las columnas objetivo seleccionadas, y opcionalmente, un término de ranking sobre aristas para mejorar el orden relativo de puntuaciones. Además, se puede aplicar un enmascarado parcial de las columnas nodales objetivo en la entrada para evitar fugas triviales (el modelo no puede “copiar” directamente esas columnas). La pérdida total se obtiene como suma ponderada de estos términos.

2.4.2. GTMVAE

En el **Graph Transformer-based Multitask Variational AutoEncoder** (GTMVAE) se replica el mismo proceso que en el GTMAE sustituyendo el *autoencoder* simple por un *variational autoencoder*. En este caso, el *encoder* estima por nodo los parámetros de una distribución latente (media y varianza) e introduce el muestreo mediante el truco de reparametrización, el cual permite “aprender” aunque haya azar: en vez de muestrear directamente un valor aleatorio, el modelo genera un valor base (media) y cuánto debe variar (desviación), y luego le añade un ruido fijo y conocido mediante el producto de Hadamard (producto elemento a elemento). Los *decoders* permanecen análogos a los del GTMAE. El objetivo de entrenamiento es una LOSS multitarea con regularización variacional:

$$LOSS_{GTMVAE} = LOSS_{GTMAE} + \beta_{KL} KL(q_{\phi}(z|x)||p(z)),$$

donde:

- $LOSS_{GTMAE}$ equivale a la función de pérdida del modelo GTMAE.
- $KL(q||p)$ es la divergencia de Kullback–Leibler (KL) (Shlens 2014), la cual mide cuánta información extra, o cuán “sorprendente”, es la creencia del posterior aproximado del *encoder* q respecto al prior gaussiano p : si q y p son casi iguales, la KL es pequeña; si q se aleja (medias grandes o varianzas muy pequeñas), la KL crece y penaliza el entrenamiento.
- β pondera la divergencia KL y controla el compromiso de reconstrucción-regularización. Puede annealarse en el tiempo para introducir el valor de KL de forma progresiva (proceso también denominado *warm-up*).

2.4.3. E2A-SAGE-MAE

Como variación de los modelos anteriores, el **Edge-Attribute Aware GraphSAGE Multitask AutoEncoder** (E2A-SAGE-MAE) (Hamilton *et al.* 2017, Lo *et al.* 2022) mantiene la misma estructura general que el GTMAE, pero sustituye el *encoder* basado en *Transformer* por uno construido sobre un operador EdgeSAGE. Este operador no emplea mecanismos de atención multi-cabeza, sino que agrega información vecinal mediante una media ponderada lineal de dos componentes: una transformación de las características del nodo vecino y otra de los atributos de la arista.

A nivel de regularización y estabilización, el E2A-SAGE-MAE introduce elementos que no aparecen en el GTMAE, como una capa de *batch normalization* tras la primera convolución y la opción de normalizar en norma L2 las salidas de cada capa. Junto al uso de *dropout*, estas modificaciones buscan controlar la varianza interna de las activaciones y mejorar la generalización en escenarios con ruido o con relaciones heterogéneas. En resumen, el E2E-SAGE-MAE replica la misma lógica general del GTMAE en cuanto a *pipeline* y *decoder*, pero se distingue por un *encoder* más ligero, estable y explícitamente lineal en la incorporación de atributos de arista y representaciones vecinales.

2.4.4. TGT

Añadiendo la dimensión temporal, se desarrolla el **Temporal Graph Transformer** (TGT) (Yu *et al.* 2023), el cual trabaja con un grafo heterogéneo a partir de las cuatro fuentes de datos estudiadas:

atributos nodales estáticos, atributos nodales temporales, aristas direccionadas y aristas no direccionadas.

En primer lugar, las aristas se preprocesan del mismo modo que en apartados anteriores: las dirigidas se codifican como origen-destino más un tensor de los valores de unión, y las no dirigidas se duplican. El modelo procesará estos conjuntos con parámetros distintos, reflejando la naturaleza de cada relación.

Por otro lado, al aportarse una secuencia temporal de atributos por municipio y año con coberturas parciales (algunas series empiezan o acaban en años distintos), se reorganiza la información en un tensor denso de forma $[T, N, D]$, donde T es el número de años distintos, N el número de municipios y D el número de variables temporales. Los huecos no observados se tratan de dos maneras complementarias:

1. Se aplica una imputación ligera por municipio para estabilizar la escala y evitar propagar nulos en el cómputo.
2. Se genera un canal de máscara de presencia del mismo tamaño que marca 1 cuando el dato exista y 0 cuando se imputa. Esta máscara se concatena a las features temporales para que el modelo “vea” explícitamente la incertidumbre.

En cada año t se construye la entrada nodal como la concatenación de tres bloques: atributos estáticos, atributos temporales del año t y máscara de presencia del año t . Esta concatenación se proyecta a una dimensión intermedia con una capa lineal para homogeneizar escalas. A continuación, se aplica un codificador espacial heterogéneo compuesto por dos *Transformer* convolucionales (uno por relación: dirigida y no dirigida), que realizan *self-attention* sobre el grafo usando tanto las variables nodales proyectadas como los atributos de arista. Las dos salidas se fusionan por concatenación y una proyección lineal, produciendo un *embedding* por año:

$$h \in \mathbb{R}^{N \times H}$$

Para introducir la información temporal absoluta, cada año se “etiqueta” con un *temporal encoding sinusoidal* centrado en el año objetivo. Este vector de codificación temporal se proyecta a la misma dimensión H y se suma a h , de modo que el modelo pueda distinguir explícitamente el año que está procesando y a qué distancia queda del año objetivo.

Después, todos los *embeddings* anuales se apilan y se pasan a un *Transformer* temporal que modela dependencias de largo alcance entre años para cada municipio. La salida del *Transformer* temporal es una secuencia refinada por municipio. Acto seguido, se realiza una agregación ponderada hacia el año objetivo: se calculan pesos que priorizan los años más cercanos al año objetivo, se normalizan para que sumen 1 y se hace un promedio ponderado. Este vector agregado se proyecta finalmente con una capa lineal de lectura para producir el *embedding* final por municipio, que se interpreta como “la foto del año objetivo informada por la historia”.

2.4.5. HGT-TE

Por último, buscando un enfoque más general y heterogéneo que el del TGT, se desarrolla un *Heterogeneous Graph Transformer with Temporal Encoding* (HGT-TE) (Hu *et al.* 2020, Wang 2025, Li 2025). A nivel conceptual, el TGT trata la heterogeneidad “a mano”: separa explícitamente las dos relaciones, aplica por cada una un *Transformer edge-aware*, fusiona sus salidas y luego añade *temporal encoding* y un peso exponencial hacia el año objetivo; es decir, el énfasis está en explotar los atributos asociados a cada arista y en priorizar el año objetivo con una agregación diseñada. En cambio, el HGT-TE desplaza el foco a la heterogeneidad de tipos (*type-aware*), permitiendo que distintos tipos de nodos y relaciones sean tratados de manera diferenciada mediante parámetros y atenciones específicas para cada tipo.

En términos de codificación temporal, ambos modelos incorporan representaciones del paso del tiempo en los *embeddings* para capturar la evolución dinámica de las relaciones, pero en el caso del

HGT-TE esta información se integra conjuntamente con la naturaleza heterogénea del grafo, ajustando las atenciones según el tipo de arista y el instante temporal en que ocurre la interacción.

2.5. Validación

La validación de *graph embeddings* presenta un desafío particular cuando no existe una tarea supervisada explícita. Entre las estrategias más comunes se encuentra la visualización mediante técnicas de reducción de dimensionalidad como *t-distributed Stochastic Neighbor Embedding* (t-SNE) (Cai & Ma 2021) o *Uniform Manifold Approximation and Projection* (UMAP) (Ghojogh *et al.* 2021), que permiten proyectar las representaciones latentes a menos dimensiones preservando las relaciones locales o globales. Matemáticamente, estas técnicas buscan minimizar una divergencia entre distribuciones de proximidad en el espacio original y el espacio proyectado.

El t-SNE trata de proyectar los puntos del espacio original preservando la probabilidad de vecindad local (distorsionando por otro lado de manera negativa las distancias globales), mientras que UMAP asume que los datos viven sobre una variedad y trata de preservar la estructura topológica local y global. A través de estas técnicas se pueden analizar visualmente grupos de municipios que compartan propiedades similares o contrastar las proyecciones con datos externos para evaluar la coherencia de las agrupaciones; sin embargo, resultan limitadas para determinar si los *embeddings* representan correctamente a los municipios, dado que en la reducción de dimensionalidad se pierde una parte significativa de la información.

Adicionalmente, una estrategia complementaria consiste en trasladar al ámbito de los grafos la lógica de evaluaciones estandarizadas como el *Massive Text Embedding Benchmark* (MTEB) (Muennighoff *et al.* 2022), cuyo principio fundamental es validar la calidad de un *embedding* a través de un conjunto amplio de tareas externas. Inspirándose en este enfoque, la validación se plantea como la predicción de variables incluidas y no incluidas en la construcción del grafo, de manera que la capacidad de los *embeddings* para generalizar y capturar información latente pueda medirse indirectamente a través de su rendimiento en dichas tareas. La validación se articula mediante la predicción de cuatro grupos de variables:

- **Seis variables estáticas nodales conocidas por el modelo:** renta media bruta, número de afiliados a la seguridad social por habitante, precio medio de vivienda, tamaño medio de vivienda, población e indicador dicotómico de colindancia con Guadalajara.
- **Dos variables estáticas nodales desconocidas para el modelo:** número de documentos notariales de créditos, préstamos y garantías hipotecarias y número de líneas telefónicas.
- **Dos variables estáticas relacionales conocidas por el modelo:** distancia en línea recta y grado de unión a través de autobús entre municipios.
- **Nueve variables temporales conocidas por los modelos de carácter temporal:** diferencia media en los ejercicios registrados, diferencia en el último ejercicio e indicador dicotómico de si la última diferencia es positiva o negativa para las variables de edad media, número de unidades urbanas y densidad de población.

Para todas ellas se realiza una predicción utilizando un modelo *eXtreme Gradient Boosting* (XGB), evaluado con R^2 o F1-score según la tipología del dato. Este diseño permite contrastar la robustez de las representaciones latentes bajo distintos paradigmas de validación y, en línea con la filosofía de *benchmarks* como MTEB, ofrece una medida indirecta pero objetiva de la calidad de los *embeddings* generados.

Para comparar el rendimiento de las diferentes predicciones se emplearán diferentes tests estadísticos. En primer lugar, se aplicará el **test de Friedman** (Friedman 1937), el cual contrasta la hipótesis nula de que todos los modelos de *embeddings* presentan un rendimiento equivalente a lo largo de las distintas validaciones. Este test, basado en rangos, resulta adecuado cuando se comparan múltiples algoritmos sobre los mismos conjuntos de validación y no requiere asumir normalidad en los datos.

Una vez detectadas diferencias globales significativas, se recurrirá al **test de Nemenyi** (Nemenyi 1963), el cual permite identificar qué pares de modelos presentan diferencias estadísticamente relevantes en sus rangos medios, y se complementará con el **test de Wilcoxon pareado con corrección de Holm**

(Wilcoxon 1945), que ofrece un contraste más detallado de cada modelo frente al candidato ganador, controlando el error de comparaciones múltiples.

Finalmente, para integrar la incertidumbre asociada a los intervalos de confianza de cada validación, se implementará un **bootstrap paramétrico** (Efron & Tibshirani 1993). A partir de los límites de los intervalos de confianza se estimará la variabilidad de cada métrica y se simularán múltiples replicaciones, lo que permitirá calcular tanto distribuciones de rangos medios como probabilidades de que cada modelo sea el mejor.

Adicional a lo comentado, se realiza una **evaluación intrínseca de los embeddings** a partir de una **métrica de coherencia de vecinos**. Ésta se basa en identificar, para cada nodo, sus K vecinos más próximos según la distancia coseno y comparar su similitud respecto a variables externas de referencia. Si la variable considerada es categórica, se toma directamente su valor; si es numérica, se discretiza previamente en intervalos definidos por cuantiles, transformándola en categorías comparables. Posteriormente, se calcula la proporción de coincidencia entre el nodo de referencia y sus vecinos, dividiendo el número de vecinos que comparten la misma categoría o cuantil entre el total de vecinos que disponen de valor para dicha variable.

3. Análisis de resultados

En primer lugar, a partir de la métrica de coherencia de vecinos, se examina cómo se comportan los *embeddings* en relación con las variables de referencia. El análisis busca determinar si los nodos con características similares tienden a situarse próximos en el espacio latente, lo que evidenciaría que el modelo ha capturado correctamente las relaciones estructurales subyacentes. Los resultados en la Tabla 1 muestran que, a priori, los tres modelos con mayor coherencia son GTMAE, GTMVAE y E2A-SAGE-MAE, con un acierto mayor al 76.5 %.

Tabla 1. Media de proporciones de vecinos coherentes para cada modelo de embeddings

<i>Embedding</i>	Proporción	IC 95 %
GTMAE	78.44 %	[77.98 %, 78.90 %]
GTMVAE	78.02 %	[77.57 %, 78.47 %]
E2A-SAGE-MAE	76.59 %	[76.13 %, 77.05 %]
TGT	66.71 %	[66.20 %, 67.22 %]
HGT-TE	62.92 %	[62.43 %, 63.41 %]

Acto seguido, se calculan los test estadísticos de comparación de métricas de predicción (mostradas en el Apéndice 1) para los cinco modelos candidatos.

Tabla 2. Resultados del test de Friedman

<i>Embedding</i>	Rango Medio
E2A-SAGE-MAE	1.842105
GTMAE	2.657895
GTMVAE	2.763158
HGT-TE	3.815789
TGT	3.921053

El test no paramétrico de Friedman (Tabla 2) revela diferencias estadísticamente significativas entre los cinco modelos de generación de *embeddings* (Estatístico = 46.02, p-valor = 2.44×10^{-9} , N° Bloques = 38). Esto indica que el rendimiento medio de los modelos no es equivalente en las distintas tareas y configuraciones evaluadas. El ranking medio muestra una clara jerarquía de desempeño: **E2A-SAGE-MAE** obtuvo el mejor rango promedio (1.84), seguido por GTMAE (2.66) y GTMVAE (2.76), mientras que los modelos temporales HGT-TE (3.82) y TGT (3.92) presentaron los peores resultados globales.

Estos resultados confirman de forma estadísticamente robusta que los modelos estáticos, especialmente E2A-SAGE-MAE, superan de manera consistente a las variantes temporales en el conjunto de predicciones analizadas.

Tabla 3. Resultados del test de Nemenyi

<i>Embedding A</i>	<i>Embedding B</i>	Rango A	Rango B	$ \Delta\text{Rango} $	Significativo
E2A-SAGE-MAE	GTMAE	1.842105	2.657895	0.815789	FALSE
E2A-SAGE-MAE	GTMVAE	1.842105	2.763158	0.921053	FALSE
E2A-SAGE-MAE	HGT-TE	1.842105	3.815789	1.973684	TRUE
E2A-SAGE-MAE	TGT	1.842105	3.921053	2.078947	TRUE
GTMAE	GTMVAE	2.657895	2.763158	0.105263	FALSE
GTMAE	HGT-TE	2.657895	3.815789	1.157895	FALSE
GTMAE	TGT	2.657895	3.921053	1.263158	TRUE
GTMVAE	HGT-TE	2.763158	3.815789	1.052632	FALSE
GTMVAE	TGT	2.763158	3.921053	1.157895	FALSE
HGT-TE	TGT	3.815789	3.921053	0.105263	FALSE

De manera complementaria, el contraste de Nemenyi mostrado en la Tabla 3 permite identificar en qué pares de modelos se concentran las diferencias observadas en el Friedman. Los resultados muestran que **E2A-SAGE-MAE presenta diferencias significativas respecto a HGT-TE y TGT**, con una distancia de rangos superior al umbral crítico ($|\Delta\text{Rango}| = 1.97$ y 2.08 frente a una diferencia crítica de 1.20). En cambio, las comparaciones entre E2A-SAGE-MAE y los otros modelos estáticos (GTMAE y GTMVAE) no resultan significativas, lo que indica un rendimiento similar entre ellos. Asimismo, la diferencia entre GTMAE y TGT también alcanza significación, reforzando la inferioridad de los modelos temporales. En conjunto, el test de Nemenyi confirma que las diferencias detectadas por el Friedman se deben principalmente a la peor actuación de las arquitecturas temporales frente a las variantes estáticas, con E2A-SAGE-MAE destacando como el modelo con mejor comportamiento global.

Tabla 4. Resultados del test de Wilcoxon-Holm

<i>Embedding A</i>	<i>Embedding B</i>	p-valor	p-valor Ajustado	Significativo
E2A-SAGE-MAE	TGT	5.54E-03	5.54E-02	TRUE
E2A-SAGE-MAE	HGT-TE	2.02E-01	1.82E+00	TRUE
E2A-SAGE-MAE	GTMVAE	5.40E+01	4.32E+02	TRUE
GTMAE	HGT-TE	1.83E+03	1.28E+04	TRUE
E2A-SAGE-MAE	GTMAE	1.30E+04	7.83E+04	FALSE
GTMVAE	HGT-TE	1.30E+04	6.52E+04	FALSE
GTMAE	TGT	2.43E+04	9.73E+04	FALSE
GTMVAE	TGT	7.92E+04	2.38E+05	FALSE
HGT-TE	TGT	1.53E+05	3.05E+05	FALSE
GTMAE	GTMVAE	5.47E+05	5.47E+05	FALSE

El contraste pareado de Wilcoxon con corrección de Holm (Tabla 4) ratifica los hallazgos anteriores, mostrando diferencias estadísticamente significativas en múltiples comparaciones bilaterales. En particular, **E2A-SAGE-MAE superó significativamente a TGT, HGT-TE y GTMVAE** ($p < 10^{-6}$ en todos los casos) y mostró una ventaja marginal frente a GTMAE ($p\text{-valor} = 0.013$, $p\text{-valor ajustado} = 0.078$). Asimismo, se observó una diferencia significativa entre GTMAE y HGT-TE ($p\text{-valor} = 0.0018$, $p\text{-valor ajustado} = 0.0128$), lo que refuerza la tendencia de menor rendimiento de los modelos temporales. En conjunto, los resultados del Wilcoxon-Holm consolidan la superioridad estadística del modelo E2A-SAGE-MAE, al tiempo que confirman que las variantes basadas en transformadores temporales (TGT y HGT-TE) se sitúan consistentemente por debajo en las métricas de predicción evaluadas.

Tabla 5. Resultados del bootstrap paramétrico

<i>Embedding A</i>	<i>Embedding B</i>	Diferencia Media	p-valor Bilateral	CI95%
E2A-SAGE-MAE	HGT-TE	24.314729	0.0137	[5.319, 43.372]
E2A-SAGE-MAE	GTMAE	8.585144	0.0311	[0.703, 16.519]
HGT-TE	TGT	-20.464974	0.0400	[-39.874, -1.011]
E2A-SAGE-MAE	GTMVAE	5.567229	0.0403	[0.264, 10.994]
E2A-SAGE-MAE	TGT	3.849755	0.0480	[0.034, 7.772]
GTMVAE	HGT-TE	18.747500	0.0615	[-1.023, 38.644]
GTMAE	HGT-TE	15.729585	0.1294	[-4.892, 36.581]
GTMAE	TGT	-4.735390	0.2984	[-13.424, 4.120]
GTMAE	GTMVAE	-3.017915	0.5400	[-12.453, 6.572]
GTMVAE	TGT	-1.717475	0.6110	[-8.262, 4.993]

Por último, el análisis mediante *bootstrap* paramétrico mostrado en la Tabla 5 permite estimar diferencias medias y sus intervalos de confianza. Los resultados confirman diferencias significativas entre E2A-SAGE-MAE y los modelos HGT-TE ($\Delta = +24.31$, p-valor = 0.0137, IC95% [5.3, 43.4]), GTMAE ($\Delta = +8.59$, p-valor = 0.0311), GTMVAE ($\Delta = +5.57$, p-valor = 0.0403) y TGT ($\Delta = +3.85$, p-valor = 0.048). Además, el modelo TGT muestra peor rendimiento que HGT-TE ($\Delta = -20.46$, p-valor = 0.04), corroborando la debilidad de las variantes temporales. Estos resultados, coherentes con los obtenidos en los contrastes no paramétricos, evidencian que **E2A-SAGE-MAE no solo obtiene mejores rangos, sino también diferencias medias positivas estadísticamente significativas** en magnitud de rendimiento frente al resto de modelos comparados.

Los resultados obtenidos con E2A-SAGE-MAE demuestran que la combinación de información de nodos y aristas junto con una arquitectura multitarea ofrece una ventaja diferencial frente a modelos que integran información nodal e internodal pero carecen de un diseño multitarea, como GTMAE. Del mismo modo, supera a arquitecturas temporales que, si bien destacan en la literatura por su capacidad para capturar dinámicas longitudinales (Al Sahili & Awad 2023, Hu *et al.* 2020), en este caso, los resultados sugieren que la información histórica no introduce una ganancia sustancial en la capacidad representativa de los *embeddings*. Los modelos temporales tienden a comportarse como una versión suavizada de los modelos estáticos, lo que indica que la dinámica interanual es limitada y que la estructura espacial domina sobre la temporal en la configuración del territorio.

Por otro lado, en el Apéndice 1 se puede observar que, realizando una media de las métricas obtenidas para cada uno de los modelos, **los *embeddings* generados excluyendo el municipio de la ciudad de Madrid alcanzan valores cinco veces mejores que los obtenidos incluyendo dicho municipio**, lo cual valida la hipótesis planteada anteriormente de que Madrid actúa como un *outlier* dentro del grafo. Este efecto no solo se refleja en la media global, sino también en casos concretos: en E2A-SAGE-MAE, la métrica de predicción del tamaño medio de vivienda mejora de 0.732 a 0.761, y la población pasa de 0.697 a 0.890 al excluir Madrid. La mejora se manifiesta de forma transversal en otros modelos como GTMAE o HGT-TE, lo que indica que no se trata de un sesgo metodológico, sino de un rasgo estructural del grafo.

Para el caso específico de E2A-SAGE-MAE, se calcula la distancia de Cook en el espacio latente con el objetivo de evaluar la influencia del nodo de la ciudad de Madrid en la estructura de los *embeddings*. Esta medida, inspirada en la estadística clásica, permite estimar cuánto varía la configuración global cuando se elimina un municipio específico del conjunto de entrenamiento.

Los resultados globales muestran que **la eliminación de la ciudad de Madrid genera un efecto de gran magnitud**. El desplazamiento medio relativo alcanza un valor de 1.33 veces la escala típica del *embedding*, acompañado de un cambio medio en la similitud coseno del 6.3 %. Ambos indicadores se clasifican como altos, lo que evidencia que Madrid actúa como un nodo central que estructura de manera decisiva las posiciones y relaciones del resto de municipios en el espacio latente.

La interpretación de estos resultados sugiere que la ausencia de Madrid no solo reubica de forma significativa los vectores de representación, sino que también altera las redes de proximidad entre municipios. Este comportamiento es coherente con la posición de Madrid como polo demográfico, económico y de conectividad, cuya supresión obliga al modelo a redistribuir y reajustar la organización interna del *embedding*.

A partir de los *embeddings* generados por el E2A-SAGE-MAE (incluyendo la ciudad de Madrid para una representación global de la comunidad), y calculando la similitud coseno entre vectores, es posible identificar qué municipios presentan mayor cercanía en el espacio latente. El análisis de estas similitudes revela patrones territoriales muy coherentes con la realidad funcional y socioeconómica de la región.

El cinturón sur (Alcorcón, Leganés, Getafe, Móstoles y Fuenlabrada) aparece como un bloque muy denso, con valores de similitud superiores a 0.90 en muchos casos, lo que refleja su fuerte integración metropolitana y un perfil común de ciudades residenciales e industriales de gran tamaño. Con umbrales ligeramente más bajos se incorporan Coslada, Parla y Pinto, lo que refuerza la imagen de un corredor metropolitano muy cohesionado.

En el noroeste, los municipios de renta alta como Pozuelo, Majadahonda, Las Rozas y Boadilla forman un clúster igualmente compacto, con parejas muy próximas entre sí y conexiones evidentes con Villaviciosa de Odón y Villanueva de la Cañada. Este bloque representa la diferenciación clara de las áreas residenciales de alto estatus frente al resto del territorio.

El eje del Henares también queda bien representado. Alcalá de Henares y Torrejón de Ardoz conforman un dúo estrechamente vinculado, mientras que Coslada, San Fernando, Arganda y Rivas-Vaciamadrid se alinean con valores altos de similitud, reflejando tanto su proximidad geográfica como su especialización productiva. Llama la atención que Coslada muestre conexiones fuertes con municipios del sur (como Alcorcón), lo que sugiere que el modelo capta no solo cercanía espacial, sino también similitudes funcionales derivadas de su carácter logístico e industrial.

En el sur profundo y vegas del Tajo destacan Aranjuez, Ciempozuelos, Valdemoro, Pinto y San Martín de la Vega, todos ellos estrechamente vinculados. Este conjunto refleja un patrón mixto de función residencial, logística y agrícola, articulado en torno al eje de la A-4. De forma paralela, en la Sierra se aprecian pares y tríadas de municipios con alta similitud, como Cercedilla-Los Molinos, Guadalix-Pedrezuela o El Escorial-Galapagar, además de conjuntos más amplios en torno a Lozoya y Rascafría. Ello evidencia que el modelo distingue adecuadamente las dinámicas de segunda residencia y turismo propias de estas áreas.

Otros casos igualmente consistentes son los de Alcobendas y San Sebastián de los Reyes, que aparecen como “gemelos metropolitanos” con una similitud superior al 0.80, o el de Tres Cantos y Colmenar Viejo, que se asocian por su perfil residencial y tecnológico en el eje de la M-607. Por su parte, Madrid capital presenta similitudes moderadas con los grandes municipios del sur (0.40–0.48), lo que es lógico: mantiene la conexión funcional con su área metropolitana más inmediata, pero se diferencia claramente por su escala y centralidad.

El análisis de similitudes se acompaña de una reducción de la dimensionalidad (UMAP) de los *embeddings* plasmada en 2D (Figura 3), la cual confirma estas lecturas. El cinturón sur se representa como un bloque muy compacto en el diagrama, mientras que el noroeste de renta alta aparece igualmente agrupado. El eje del Henares se alinea en la parte oriental, con pares cercanos como Alcalá y Torrejón. La Sierra y el noreste rural se dispersan en nubes coherentes de municipios de menor tamaño, y Madrid aparece próxima al sur pero con una cierta separación, lo que ilustra su singularidad como centro de gravedad.

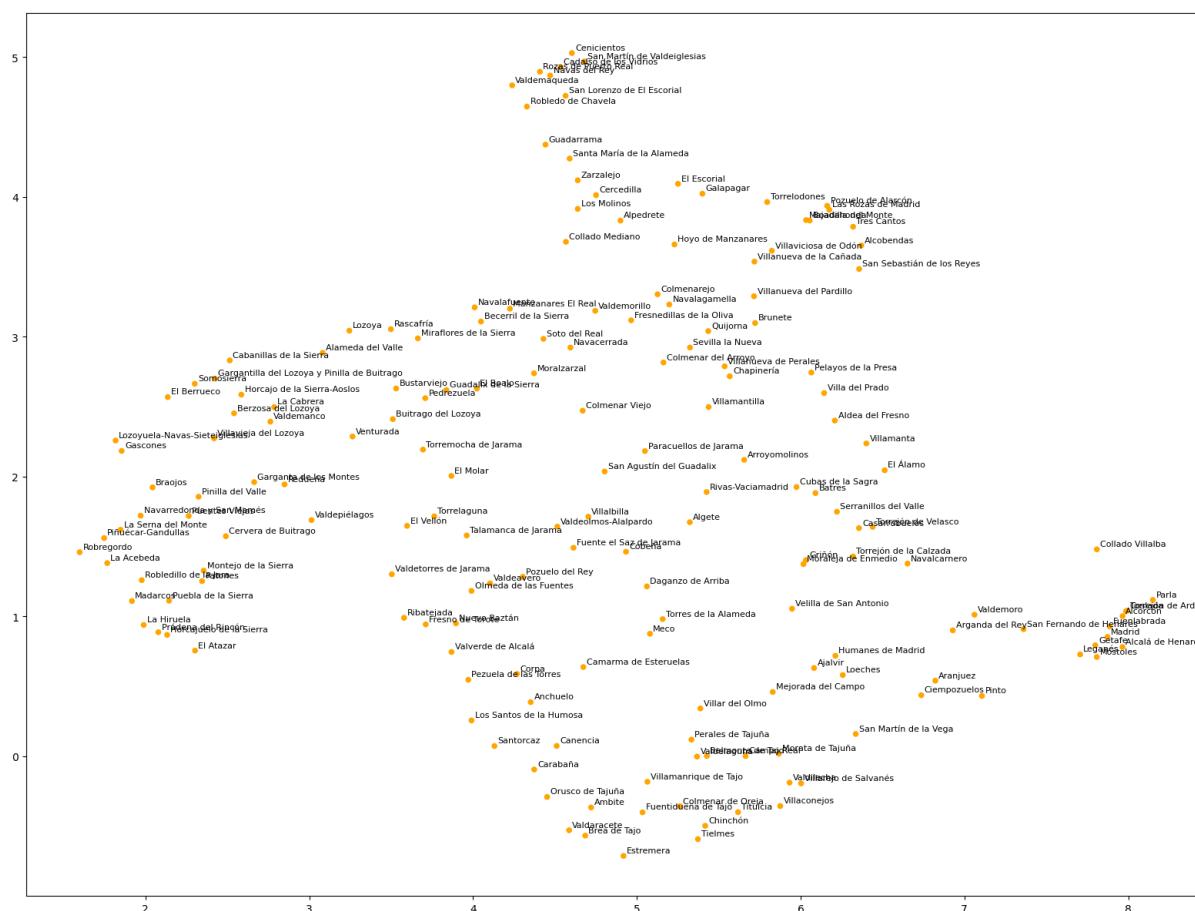


Figura 3. Representación de los *embeddings* de E2A-SAGE-MAE reducidos a 2D.

En conjunto, las representaciones vectoriales reproducen con notable fidelidad la geografía funcional de la Comunidad de Madrid. Captan no solo proximidad espacial, sino también factores socioeconómicos, productivos y de movilidad. Incluso algunos resultados llamativos, como la cercanía entre Coslada y municipios del sur o la posición de Collado Villalba respecto a grandes núcleos, son interpretables y plausibles, lo que refuerza la validez del modelo.

4. Conclusión

En conclusión, los resultados de la validación cruzada permiten afirmar que es posible representar vectorialmente unidades geográficas a partir de técnicas de *graph embeddings* y que el modelo E2A-SAGE-MAE, de entre los que se han probado, es el que mejor logra representar la estructura y dinámica de los municipios de la Comunidad de Madrid. Su capacidad para integrar información tanto de nodos como de aristas, junto con el enfoque multitarea, proporciona vectores más estables, evitando el sobreajuste a una sola variable y mejorando de manera consistente frente al resto de modelos candidatos.

Asimismo, el análisis evidencia que la exclusión de la ciudad de Madrid mejora de forma significativa la capacidad predictiva de los *embeddings*, lo que refuerza la idea de su papel como *outlier* en el conjunto de datos si el objetivo de la representación es analizar los municipios de manera particular y no requiriendo una foto completa de la comunidad. Además, aunque modelos temporales como el TGT o el HGT-TE no superan en media al E2A-SAGE-MAE, aportan un valor añadido en contextos donde la evolución temporal es clave, pudiendo ser muy útiles en la predicción a futuro de variables con carácter estacional.

Dentro de las futuras líneas de trabajo se plantea, en primer lugar, la involucración directa en casos de uso reales, siguiendo la línea de investigaciones previas (Kefato *et al.* 2021, Murray *et al.* 2023, Huang *et al.* 2024), como podría ser la predicción de índices económicos a partir de representaciones

vectoriales de los municipios, con el fin de evaluar la efectividad práctica de este enfoque. Por otro lado, se plantea validar si los *embeddings* resultantes de aumentar el rango geográfico de los datos agregando a niveles superiores mantienen un comportamiento coherente (Rodríguez-González *et al.* 2026). Por ejemplo, al pasar de municipios a provincias. La idea sería evaluar si esta agregación puede aportar más información que entrenar directamente un *embedding* de provincias a partir de las variables agregadas.

Adicionalmente, se propone la predicción de la evolución de los municipios a lo largo del tiempo, para lo cual sería necesario generar *embeddings* correspondientes a diferentes cortes temporales (varios años) y analizar el desplazamiento de cada municipio en el espacio latente. Este planteamiento permitiría anticipar qué municipios tenderán a acercarse o alejarse en términos de similitud socioeconómica. Desde el punto de vista metodológico, esta dinámica temporal podría modelarse mediante *Recurrent Neural Networks* (RNN), como LSTM o GRU, o bien con *Transformers* temporales, aprovechando su capacidad para capturar dependencias de largo plazo en secuencias (Pandhre *et al.* 2018, Dall’Amico, 2024).

Referencias

- Aalbers, M. B. (2016). *The financialization of housing: a political economy approach*. London, Routledge.
- Al Sahili, Z. A. & Awad, M. (2023). “Spatio-temporal graph neural networks: a survey”, *arXiv preprint arXiv:2301.10569*. <https://doi.org/10.48550/arXiv.2301.10569>
- Amsterdam Smart City (s. f.). *Amsterdam Smart City*. Amsterdam, Amsterdam Smart City. [accessed 11-09-2025]. Available at <https://amsterdamsmartcity.com>
- Barthélemy, M. (2011). “Spatial networks”, *Physics Reports*, 499(1-3), 1–101.
- Cai, T. T. & Ma, R. (2021). “Theoretical foundations of t-SNE for visualizing high-dimensional clustered data”. <https://doi.org/10.48550/arXiv.2105.07536>
- Castro, C. M. (1862). *Plan general de ensanche de Madrid*. Madrid, Ayuntamiento de Madrid.
- Ciudad de Madrid (s. f.). *Ciudad de Madrid: City Profile Smart Cities Report 2024*. London, SmartCitiesWorld. [accessed 11-09-2025]. Available at <https://www.smartcitiesworld.net/city-profile/smart-cities-reports/smartcitiesworld-city-profile-2024--city-of-madrid-10751>
- City of Helsinki (s. f.). *Helsinki 3D: 3D City Models of Helsinki*. Helsinki, City of Helsinki. [accessed 11-09-2025]. Available at <https://www.hel.fi/en/decision-making/information-on-helsinki/maps-and-geospatial-data/helsinki-3d>
- Dall’Amico, L., Barrat, A., Cattuto, C. (2024). *An embedding-based distance for temporal graphs*. *Nature Communications*, 15, Article 9954. <https://doi.org/10.1038/s41467-024-54280-4>
- Efron, B. & Tibshirani, R. J. (1993). *An introduction to the bootstrap*. New York, Chapman & Hall/CRC.
- Fernández, R. & Aalbers, M. B. (2016). “Financialization and housing: between globalization and varieties of capitalism”, *Competition & Change*, 20(2), 71–88.
- Friedman, M. (1937). “The use of ranks to avoid the assumption of normality implicit in the analysis of variance”, *Journal of the American Statistical Association*, 32(200), 675–701. <https://doi.org/10.1080/01621459.1937.10503522>
- Ghojogh, B., Ghodsi, A., Karray, F., Crowley, M. (2021). “Uniform manifold approximation and projection (UMAP) and its variants: tutorial and survey”. <https://doi.org/10.48550/arXiv.2109.02508>
- Goldberg, Y. & Levy, O. (2014). “word2vec explained: deriving Mikolov et al.’s negative-sampling word-embedding method”. <https://doi.org/10.48550/arXiv.1402.3722>
- Gutiérrez Puebla, J. & Rubio Barroso, A. (1997). “Los sistemas de información geográficos: origen y perspectivas”, *Revista General de Información y Documentación*, 7(1), 93–106.

- Hamilton, W. L., Ying, R., Leskovec, J. (2017). “Inductive representation learning on large graphs”. <https://doi.org/10.48550/arXiv.1706.02216>
- Hu, Z., Dong, Y., Wang, K., Sun, Y. (2020). “Heterogeneous graph transformer”. <https://doi.org/10.48550/arXiv.2003.01332>
- Huang, C., Yu, F., Wan, Z., Li, F., Ji, H., Li, Y. (2024). “Knowledge graph confidence-aware embedding for recommendation”, *Neural Networks*, 180, 106601. <https://doi.org/10.1016/j.neunet.2024.106601>
- Idealista. (2025a). *Evolución del precio medio de la vivienda en venta en la Comunidad de Madrid*. Madrid, Idealista. [accessed 11-09-2025]. Available at <https://www.idealista.com/sala-de-prensa/informes-precio-vivienda/venta/madrid-comunidad/madrid-provincia/madrid/>
- Idealista. (2025b). *Evolución del precio medio de la vivienda en alquiler en la Comunidad de Madrid*. Madrid, Idealista. [accessed 11-09-2025]. Available at <https://www.idealista.com/sala-de-prensa/informes-precio-vivienda/alquiler/madrid-comunidad/madrid-provincia/madrid/>
- Kefato, Z. T., Girdzijauskas, S., Sheikh, N., Montresor, A. (2021). “Dynamic embeddings for interaction prediction”, in *Proceedings of The Web Conference (WWW)*. Ljubljana, Slovenia, 1–10. <https://doi.org/10.1145/3442381.3450020>
- Kipf, T. N. & Welling, M. (2016). “Variational graph auto-encoders”, in *Bayesian Deep Learning Workshop (NIPS 2016)*. Barcelona, Spain. <https://doi.org/10.48550/arXiv.1611.07308>
- Lees, L., Slater, T., Wyly, E. (2013). *Gentrification*. London, Routledge.
- Li, S. (2025). “Graph transformer-based heterogeneous graph neural networks enhanced by multiple meta-path adjacency matrices decomposition”, *Neurocomputing*, 629, 129604. <https://doi.org/10.1016/j.neucom.2025.129604>
- Liang, Y., Zhu, J., Ye, W., Gao, S. (2022). “Region2Vec: community detection on spatial networks using graph embedding with node attributes and spatial interactions”, *arXiv preprint arXiv:2210.08041*. <https://doi.org/10.1145/3557915.3560974>
- Lo, W. W., Layegghy, S., Sarhan, M., Gallagher, M., Portmann, M. (2022). “E-GraphSAGE: a graph neural network based intrusion detection system for IoT”, in *IEEE/IFIP Network Operations and Management Symposium (NOMS)*. Budapest, Hungary. <https://doi.org/10.1109/NOMS54207.2022.9789878>
- Muennighoff, N., Tazi, N., Magne, L., Reimers, N. (2022). “MTEB: massive text embedding benchmark”, *arXiv preprint arXiv:2210.07316*. <https://doi.org/10.18653/v1/2023.eacl-main.148>
- Murray, D., Yoon, J., Kojaku, S., Costas, R., Jung, W., Milojević, S., Ahn, Y. (2023). “Unsupervised embedding of trajectories captures the latent structure of scientific migration”, *Proceedings of the National Academy of Sciences*, 120, e2305414120. <https://doi.org/10.1073/pnas.2305414120>
- Nemenyi, P. (1963). *Distribution-free multiple comparisons*. Princeton, Princeton University Press.
- OECD-OPSI (2024). *Virtual Twin Singapore*. Paris, OECD/OPSI. [accessed 11-09-2025]. Available at <https://oecd-opsi.org/innovations/virtual-twin-singapore/>
- Ou, M., Cui, P., Pei, J., Zhang, Z., Zhu, W. (2016). “Asymmetric transitivity preserving graph embedding”, in *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD)*. San Francisco, United States, pp. 1105–1114. <https://doi.org/10.1145/2939672.2939751>
- Qiu, J., Dong, Y., Ma, H., Li, J., Wang, K., Tang, J. (2017). “Network embedding as matrix factorization: unifying DeepWalk, LINE, PTE, and node2vec”, *arXiv preprint arXiv:1710.02971*. <https://doi.org/10.1145/3159652.315970>
- Pandhre, S., Mittal, H., Gupta, M., Balasubramanian, V. N. (2018). *STWalk: Learning trajectory representations in temporal graphs*. En *Proceedings of CoDS-COMAD 2018* (pp. 210-219). ACM. <https://doi.org/10.1145/3152494.3152512>

- Rodríguez-González, A., Sisamón Serrano, I., Esplugues Conca, I., Sánchez Santolaya, D., Medina Sánchez, J. (2026). Enforcing Vector Space Stability in Embeddings with Evolving Vocabularies: A Web-App Behavior Perspective. In: Pfahringer, B., et al. Machine Learning and Knowledge Discovery in Databases. Research Track and Applied Data Science Track. ECML PKDD 2025. Lecture Notes in Computer Science, vol 16020. Springer, Berlin, Heidelberg. https://doi.org/10.1007/978-3-662-72243-5_31
- Shehzad, A., Xia, F., Abid, S., Peng, C., Yu, S., Zhang, D., Verspoor, K. (2024). “Graph transformers: a survey”, *arXiv preprint arXiv:2407.09777*. <https://doi.org/10.48550/arXiv.2407.09777>
- Shlens, J. (2014). “Notes on Kullback–Leibler divergence and likelihood”, *arXiv preprint arXiv:1404.2000*. <https://arxiv.org/abs/1404.2000>
- Singh, M. T. (2025). “Ensemble-enhanced graph autoencoder with GAT and transformer-based encoders for robust fault diagnosis”. <https://doi.org/10.48550/arXiv.2504.09427>
- Tian, C., Zhang, Y., Weng, Z., Gu, X., Chan, W. K. V. (2021). “Learning large-scale location embedding from human mobility trajectories with graphs”, *arXiv preprint arXiv:2103.00483*. <https://arxiv.org/abs/2103.00483>
- Veličković, P., Cucurull, G., Casanova, A., Romero, A., Liò, P., Bengio, Y. (2017). “Graph attention networks”. <https://doi.org/10.48550/arXiv.1710.10903>
- Wachsmuth, D. & Weisler, A. (2018). “Airbnb and the rent gap: gentrification through the sharing economy”, *Environment and Planning A: Economy and Space*, 50(6), 1147–1170. <https://doi.org/10.1177/0308518X1877803>
- Wang, Y. (2025). “HTGformer: heterogeneous temporal graph transformer”, in *Proceedings of the 48th ACM SIGIR Conference on Research and Development in Information Retrieval*. Padua, Italy. <https://doi.org/10.1145/3726302.3730209>
- Wang, Z., Li, H., Rajagopal, R. (2020). “Urban2Vec: incorporating street view imagery and POIs for multi-modal urban neighborhood embedding”. <https://doi.org/10.1609/aaai.v34i01.5450>
- Wilcoxon, F. (1945). “Individual comparisons by ranking methods”, *Biometrics Bulletin*, 1(6), 80–83. <https://doi.org/10.2307/3001968>
- Wu, Z., Pan, S., Chen, F., Long, G., Zhang, C., Philip, S. Y. (2021). “A comprehensive survey on graph neural networks”, *IEEE Transactions on Neural Networks and Learning Systems*, 32(1), 4–24. <https://doi.org/10.1109/TNNLS.2020.2978386>
- Yao, H., Wu, F., Ke, J., Tang, X., Jia, Y., Lu, S., Gong, P., Ye, J. (2018). “Deep multi-view spatial-temporal network for taxi demand prediction”, in *Proceedings of the Thirty-Second AAAI Conference on Artificial Intelligence (AAAI-18)*. New Orleans, United States, pp. 2588–2595. <https://doi.org/10.1609/aaai.v32i1.11836>
- Yu, L., Sun, L., Du, B., Lv, W. (2023). “Towards better dynamic graph learning: new architecture and unified library”, in *Advances in Neural Information Processing Systems 36 (NeurIPS 2023)*.
- Zhang, J., Zheng, Y., Qi, D. (2017). “Deep spatio-temporal residual networks for citywide crowd flows prediction”, in *Proceedings of the Thirty-First AAAI Conference on Artificial Intelligence (AAAI-17)*. San Francisco, United States, pp. 1655–1661. <https://doi.org/10.1609/aaai.v31i1.10735>
- Zhao, X., Wang, M., Zhao, X., Li, J., Zhou, S., Yin, D., Li, Q., Tang, J., Guo, R. (2023). “Embedding in recommender systems: a survey”, *arXiv preprint arXiv:2310.18608*. <https://arxiv.org/abs/2310.18608>

APÉNDICE 1

En el presente apéndice se muestran los resultados medios de las predicciones realizadas sobre las siguientes variables de validación, incluyendo y excluyendo la ciudad de Madrid de la construcción de los *embeddings* de municipio:

- St.Kn.1: Renta media bruta
- St.Kn.2: Número de afiliados a la seguridad social
- St.Kn.3: Precio medio de vivienda
- St.Kn.4: Tamaño medio de vivienda
- St.Kn.5: Población
- St.Kn.6: Colinda con Guadalajara
- St.Unk.1: Número de documentos notariales
- St.Unk.2: Número de líneas telefónicas
- Rel.1: Distancia en línea recta entre municipios
- Rel.2: Grado de unión a través de autobús entre municipios
- Temp.1.1: Diferencia media en los últimos ejercicios (1998-2022) para la variable de edad media
- Temp.1.2: Diferencia en el último ejercicio (2021-2022) para la variable de edad media
- Temp.1.3: Indicador dicotómico de si la última diferencia (2021-2022) es positiva o negativa para la variable de edad media
- Temp.2.1: Diferencia media en los últimos ejercicios (1992-2022) para la variable de número de unidades urbanas
- Temp.2.2: Diferencia en el último ejercicio (2021-2022) para la variable de número de unidades urbanas
- Temp.2.3: Indicador dicotómico de si la última diferencia (2021-2022) es positiva o negativa para la variable de número de unidades urbanas
- Temp.3.1: Diferencia media en los últimos ejercicios (1985-2022) para la variable de densidad de población
- Temp.3.2: Diferencia en el último ejercicio (2021-2022) para la variable de densidad de población
- Temp.3.3: Indicador dicotómico de si la última diferencia (2021-2022) es positiva o negativa para la variable de densidad de población

<i>Resultados de las predicciones <u>incluyendo</u> la ciudad de Madrid en la generación de los embeddings.</i>					
	GTMMAE	GTMVAE	E2A-SAGE-MAE	TGT	HGT-TE
Rel.1	0.970200 [0.96915, 0.97125]	0.977708 [0.97501, 0.98040]	0.970686 [0.96781, 0.97356]	0.496111 [0.48179, 0.51042]	0.939608 [0.93594, 0.94327]
Rel.2	0.809465 [0.78669, 0.83224]	0.820081 [0.79894, 0.84122]	0.802160 [0.77831, 0.82601]	0.572034 [0.55376, 0.59031]	0.761973 [0.72884, 0.79510]
St.Kn.1	0.747952 [0.64374, 0.85216]	0.634399 [0.35447, 0.91432]	0.707231 [0.55872, 0.85574]	0.303558 [0.00000, 1.00000]	0.600926 [0.36637, 0.83548]
St.Kn.2	0.383729 [0.12743, 0.64002]	0.397248 [0.04614, 0.74836]	0.459792 [0.22956, 0.69002]	neg	0.393545 [0.01185, 0.77523]
St.Kn.3	0.568108 [0.22736, 0.90886]	0.552198 [0.15477, 0.94962]	0.802275 [0.73389, 0.87066]	0.970664 [0.95411, 0.98722]	0.676892 [0.50419, 0.84959]
St.Kn.4	0.645757 [0.50344, 0.78807]	0.709733 [0.59865, 0.82081]	0.731570 [0.57886, 0.88428]	0.639177 [0.44951, 0.82884]	0.669641 [0.53394, 0.80534]
St.Kn.5	neg	neg	0.696506 [0.26696, 1.00000]	neg	neg
St.Kn.6	0.768485 [0.64051, 0.89645]	0.757778 [0.57258, 0.94297]	0.770707 [0.50802, 1.00000]	0.266752 [0.04367, 0.48983]	0.717619 [0.59566, 0.83957]

<i>Resultados de las predicciones <u>incluyendo</u> la ciudad de Madrid en la generación de los embeddings.</i>					
	GTMAE	GTMVAE	E2A-SAGE-MAE	TGT	HGT-TE
St.Unk.1	neg	neg	0.570085 [0.19267, 0.94750]	neg	neg
St.Unk.2	neg	neg	0.619423 [0.19753, 1.00000]	neg	neg
Temp.1.1	0.425983 [0.17982, 0.67215]	0.301559 [0.00000, 5 0.8236]	0.361985 [0.06097, 0.66300]	neg	0.327627 [0.15689, 0.49836]
Temp.1.2	neg	neg	neg	neg	neg
Temp.1.3	0.935424 [0.88183, 0.98902]	0.921032 [0.88492, 0.95714]	0.923134 [0.86966, 0.97661]	0.915843 [0.86406, 0.96762]	0.912175 [0.86869, 0.95565]
Temp.2.1	neg	neg	0.594989 [0.25504, 0.93493]	neg	neg
Temp.2.2	neg	neg	neg	neg	neg
Temp.2.3	0.791793 [0.72870, 0.85489]	0.806358 [0.74487, 0.86784]	0.794566 [0.72715, 0.86198]	0.810766 [0.77051, 0.85102]	0.846590 [0.78755, 0.90563]
Temp.3.1	0.524596 [0.26414, 0.78504]	0.433883 [0.00000, 8 0.95120]	0.419017 [0.00000, 9 0.8922]	neg	0.486568 [0.25461, 0.71853]
Temp.3.2	neg	neg	neg	neg	neg
Temp.3.3	0.864065 [0.80388, 0.92425]	0.893370 [0.83900, 0.94774]	0.903111 [0.88234, 0.92388]	0.846029 [0.78558, 0.90647]	0.894365 [0.86300, 0.92573]

<i>Resultados de las predicciones <u>excluyendo</u> la ciudad de Madrid en la generación de los embeddings.</i>					
	GTMAE	GTMVAE	E2A-SAGE-MAE	TGT	HGT-TE
Rel.1	0.971934 [0.96992, 0.97395]	0.976607 [0.97517, 0.97805]	0.967407 [0.96514, 0.96967]	0.469750 [0.46169, 0.47782]	0.697480 [0.68110, 0.71386]
Rel.2	0.800450 [0.78953, 0.81137]	0.803031 [0.78592, 0.82014]	0.801007 [0.78703, 0.81499]	0.552381 [0.51249, 0.59227]	0.682508 [0.66353, 0.70148]
St.Kn.1	0.574236 [0.01233, 1.00000]	0.711390 [0.54874, 0.87404]	0.685886 [0.49491, 0.87686]	0.525793 [0.04526, 1.00000]	neg
St.Kn.2	0.399949 [0.23127, 0.56863]	0.549644 [0.27417, 0.82512]	0.592646 [0.30349, 0.88180]	neg	neg
St.Kn.3	0.502454 [0.00000, 1.00000]	0.669175 [0.49515, 0.84321]	0.625383 [0.34821, 0.90256]	0.959714 [0.88374, 1.00000]	neg
St.Kn.4	0.625866 [0.45047, 0.80127]	0.697360 [0.54637, 0.84835]	0.760839 [0.70412, 0.81756]	0.298358 [0.09557, 0.50114]	neg
St.Kn.5	0.893040 [0.79148, 0.99460]	0.787684 [0.56712, 1.00000]	0.890042 [0.82744, 0.95264]	0.103918 [0.00000, 0.58553]	neg
St.Kn.6	0.540952 [0.33088, 0.75102]	0.548254 [0.25489, 0.84162]	0.769841 [0.49675, 1.00000]	0.127778 [0.00000, 0.27732]	0.000000 [0.00000, 0.00000]
St.Unk.1	0.691018 [0.44581, 0.93623]	0.612646 [0.29101, 0.93428]	0.744377 [0.57111, 0.91765]	0.376211 [0.00405, 0.74838]	neg
St.Unk.2	0.925299 [0.88122, 0.96937]	0.877000 [0.77631, 0.97769]	0.883546 [0.81006, 0.95704]	0.335587 [0.07359, 0.59759]	neg
Temp.1.1	0.352695 [0.12979, 0.57560]	0.366215 [0.13251, 0.59992]	0.377301 [0.21128, 0.54332]	0.016011 [0.00000, 0.22535]	neg
Temp.1.2	neg	neg	0.108583 [0.00000, 0.3577]	neg	neg
Temp.1.3	0.934139 [0.89515, 0.97313]	0.918102 [0.86197, 0.97423]	0.933921 [0.90218, 0.96567]	0.885476 [0.82180, 0.94915]	0.881084 [0.82092, 0.94125]
Temp.2.1	0.801828 [0.63518, 0.96848]	0.676327 [0.25433, 1.00000]	0.862066 [0.73584, 0.98829]	neg	neg

Resultados de las predicciones excluyendo la ciudad de Madrid en la generación de los embeddings.

	GTMAE	GTMVAE	E2A-SAGE-MAE	TGT	HGT-TE
Temp.2.2	neg	neg	neg	neg	neg
Temp.2.3	0.817944 [0.74463, 0.89126]	0.809721 [0.75170, 0.86774]	0.792683 [0.72985, 0.85552]	0.778899 [0.74528, 0.81252]	0.816709 [0.80243, 0.83099]
Temp.3.1	0.559508 [0.26545, 0.85357]	0.604831 [0.44262, 0.76704]	0.644628 [0.38927, 0.89999]	neg	neg
Temp.3.2	neg	neg	neg	neg	neg
Temp.3.3	0.880822 [0.83759, 0.92406]	0.874511 [0.81834, 0.93069]	0.880403 [0.82960, 0.93121]	0.836432 [0.77566, 0.89720]	0.893822 [0.85731, 0.93033]

Los valores “neg” representan métricas negativas.

