



Universidad de Valladolid

FACULTAD DE CIENCIAS

TRABAJO FIN DE MÁSTER EN MATEMÁTICAS

***ANÁLISIS CLUSTER EN
PREFERENCIAS APROBATORIAS***

Autor: Víctor R. Díez Recio
Tutor: José Luis García Lapresta

16 de septiembre de 2025

Resumen

Las preferencias aprobatorias combinan la clasificación de un conjunto de alternativas en dos categorías: aceptables e inaceptables, con las preferencias sobre las alternativas de cada una de las categorías mediante órdenes débiles. En este trabajo se analizará la dispersión, la diversidad y la polarización de los agentes que muestran sus opiniones sobre un conjunto de alternativas a través de preferencias aprobatorias. Para ello se considerarán distancias entre dicho tipo de preferencias y los análisis cluster correspondientes.

Palabras clave: preferencias aprobatorias, medidas de distancia, análisis cluster, polarización, consenso y diversidad.

Abstract

Approval preferences combine the classification of a set of alternatives into two categories: acceptable and unacceptable, with weak orders within each category. This thesis analyses the dispersion, diversity and polarization of agents who express their opinions on a set of alternatives through approval preferences. To this end, several distance measures between such preferences and the corresponding clustering analyses are considered.

Key words: approval preferences, distance measures, clustering analysis, polarization, consensus and diversity.

Índice general

1	Introducción	7
2	Preferencias aprobatorias	8
2.1	Motivación y contexto	8
2.2	Definición formal	12
2.3	Propiedades lógicas de las preferencias aprobatorias	15
2.4	Codificación de las preferencias aprobatorias	18
2.5	Tamaño del espacio de preferencias aprobatorias	23
3	Distancias entre preferencias aprobatorias	26
3.1	Motivación y contexto	26
3.2	Distancias sobre la parte de aprobación	28
3.3	Distancias sobre la parte ordinal	30
3.4	Combinación convexa de distancias	34
3.5	Distancia D_r lambda entre preferencias aprobatorias	35
4	Métodos de análisis clúster	40
4.1	Tipos de métodos de agrupamiento	40
4.1.1	Ranked K-Medoids (RKM)	46
4.2	Elección de la distancia y del método según el objetivo	52
5	Análisis empírico	55
5.1	Elección del método	57
5.2	Clustering de votantes	57
5.2.1	Diferencias internas entre los dos clústeres	59
5.2.2	Relación con variables sustantivas	61
6	Conclusiones y otras líneas de investigación	63

Lista de gráficos

3.1	Evolución de D_{λ}^1 en función de λ	37
3.2	Evolución de $D_{0,5}^r$ en función de r	38
3.3	Mapa de calor de D_{λ}^r en función de λ y r , con isocuantas.	39
4.1	Dendrograma obtenido mediante clustering jerárquico aglomerativo.	45
5.1	Silhouette media en función del número de clústeres k	58
5.2	Representación en dos dimensiones (MDS) de los clústeres de votantes obtenidos con $k = 2$ y distancia D_{λ}^r	59
5.3	Diez alternativas con mayor diferencia en la proporción de aprobación (<i>Cluster 1</i> – <i>Cluster 2</i>). Las barras positivas indican mayor aprobación en el clúster 1 y las negativas, en el clúster 2.	60
5.4	Diez alternativas con mayor diferencia en la posición media (<i>Cluster 1</i> – <i>Cluster 2</i>). Los valores negativos indican que la alternativa se sitúa en posiciones más altas en el clúster 2 y los positivos en el clúster 1.	61
5.5	Distribución de la ideología política (escala de cinco puntos) en cada clúster.	62

Lista de tablas

2.1	Número de estructuras posibles según el modelo de representación de preferencias, para distintos valores de n , [1].	24
5.1	Variables de termómetros de favorabilidad seleccionadas.	55
5.2	Variables sociodemográficas y político-electorales utilizadas como comparativas en el análisis.	56

Capítulo 1

Introducción

A la hora de tomar una decisión colectiva, especialmente en contextos sociales, políticos o estratégicos, es muy frecuente que los individuos no solo manifiesten un orden de preferencia sobre las alternativas disponibles, sino que también indiquen explícitamente cuáles de ellas son aceptables o inaceptables. Este tipo de juicios mixtos, que combinan información de orden con un umbral de aceptabilidad, se llaman *preferencias aprobatorias*, [7, 8, 12].

En este Trabajo de Fin de Máster se estudiará sobre este tipo de preferencias, con el objetivo de analizar, desde una perspectiva matemática, cómo se puede medir la similitud o disimilitud entre diferentes agentes que expresan opiniones a través de preferencias aprobatorias, y cómo dichas medidas permiten llevar a cabo técnicas de agrupamiento (clustering), facilitando así el análisis de la diversidad, polarización o consenso en los datos.

El trabajo recorre desde los fundamentos teóricos hasta la aplicación empírica del modelo de preferencias aprobatorias en un conjunto de datos real. Se estructura de la siguiente forma. En el capítulo 2 se presentan estas estructuras, explicando definiciones y resultados necesarios para el análisis. También se analizan sus propiedades, su representación computacional y su relación con otros modelos de preferencias. El capítulo 3 introduce las distintas medidas de distancia que se usarán para comparar las preferencias aprobatorias. Se consideran distancias para la parte ordinal y para la parte de aprobación, así como combinaciones de ambas. Estas medidas permiten medir el desacuerdo entre agentes y servirán como base para los métodos de agrupamiento. En el capítulo 4 se describen diferentes técnicas de análisis cluster aplicables a este contexto. Se presentan el clustering jerárquico y el algoritmo Ranked K-Medoids, que permiten agrupar individuos o alternativas en función de la similitud en su estructura. El capítulo 5 está dedicado al estudio empírico, donde aplicaremos estos métodos sobre un conjunto de datos —. Se codificarán las preferencias, se calcularán distancias, y se realizará el agrupamiento, interpretando los resultados obtenidos, para estudiar si existe diversidad, polarización o consenso en los datos. Por último, en el capítulo 6 se comentarán las conclusiones principales del trabajo y se plantearán posibles extensiones y líneas futuras de investigación.

Esta estructura permite pasar de la teoría a la práctica, aplicando el modelo para identificar y estudiar patrones de consenso, diversidad y polarización en los datos.

Capítulo 2

Preferencias aprobatorias

En este capítulo se presenta el marco teórico de las *preferencias aprobatorias*, una estructura que combina simultáneamente la información ordinal (relaciones de orden entre alternativas) con la información evaluativa (umbrales de aceptabilidad). Este modelo surge como una extensión natural de los sistemas clásicos de votación y constituye un punto intermedio entre los órdenes lineales y los sistemas de aprobación pura. Esta dualidad lo convierte en una herramienta flexible para describir cómo los individuos expresan tanto sus jerarquías de preferencia como sus límites de aceptación.

El capítulo se organiza de manera progresiva. Comienza con una motivación general y una comparación con otros modelos de representación de preferencias, para situar el modelo en el contexto de la literatura. A continuación se presenta la definición formal de las preferencias aprobatorias, destacando sus propiedades lógicas y su clausura estructural. Posteriormente se introduce la codificación computacional mediante vectores ordinales y de aprobación, lo que permite cuantificar el grado de similitud entre preferencias. Finalmente, se revisan ejemplos y aplicaciones que ilustran la utilidad práctica de este marco. Esta base teórica servirá en los capítulos siguientes para desarrollar métricas de distancia, estudiar medidas de consenso y aplicar métodos de análisis clúster.

2.1. Motivación y contexto

Cuando abordamos un problema de decisión colectiva, los individuos no sólo pueden manifestar gustos diferentes, sino también distintos niveles de tolerancia hacia las opciones disponibles. En algunos contextos basta con saber el orden en que cada persona clasificaría las alternativas, desde la más preferida hasta la que menos. Sin embargo, en muchas situaciones sociales, políticas o estratégicas, ese orden no es suficiente. Además de saber qué opción está en primer lugar, interesa conocer si esa alternativa es al menos aceptable para el individuo o si por el contrario es completamente inaceptable.

De este modo, surge la necesidad de contar con modelos con dos tipos de información: por un lado, la información *ordinal*, que recoge cómo se ordenan las alternativas; y por otro, la información *evaluativa*, que contiene un umbral entre lo que se considera aceptable y lo que no lo es. La combinación de estas dos características nos proporciona una representación más detallada de las preferencias individuales y, por tanto, permite un análisis más ajustado a la realidad. Esta situación es particularmente relevante en votaciones públicas, procesos de deliberación en comités o contextos estratégicos, donde las preferencias no son solo qué alternativa está más arriba o más abajo en un ranking, sino también a qué opciones tienen el menor consenso o generan una inaceptabilidad total. De este modo, se puede definir el modelo de *preferencias*

aprobatorias como una fusión entre el ordenamiento clásico de las alternativas y la posibilidad de declarar explícitamente qué opciones son aceptables.

Una vez que ya hemos tratado la necesidad de añadir esta información a las votaciones, veamos que ventajas y limitaciones tienen algunos modelos clásicos de votación frente a las preferencias aprobatorias. Cada sistema puede favorecer a unos tipos de candidatos o alternativas, y se pueden producir resultados muy distintos según la forma en que los individuos expresen sus preferencias, como son los casos descritos a continuación:

- En primer lugar, los modelos *ordinales*, como los métodos de Condorcet o la regla de Borda, [5], se basan solo en el orden en que cada persona clasifica las opciones, sin tener en cuenta cuánto le gusta cada una. Esto permite captar ciertos matices. Por ejemplo, un candidato puede ser ligeramente preferido sobre otro, o aparecer mucho entre los primeros puestos, lo que le da una ventaja frente a otros que tienen opiniones más divididas. Sin embargo, este tipo de enfoque no distingue entre opciones que son apenas aceptables y otras que, aunque estén en segundo o tercer lugar, en realidad resultan inaceptables para una gran parte de los votantes. Por eso, los modelos ordinales suelen favorecer a quienes tienen un buen promedio en los rankings, aunque eso no signifique que sean bien recibidos por la mayoría.
- Por otro lado, los modelos *binarios de aprobación*, como el *Voto Aprobatorio* propuesto por Brams y Fishburn (1978), [6], permiten a los votantes indicar simplemente qué alternativas son aceptables y cuáles no. La ventaja de este sistema es su simplicidad, pues se ve el grado de aceptación mínima de cada opción y evita resultados en los que gane un candidato inaceptable por gran parte de los votantes. Sin embargo, este sistema pierde información, pues no se diferencia entre la opción favorita absoluta y otra que, aunque sea aceptable, se considere de menor valor. Por esta razón, este sistema de aprobación tiende a favorecer a candidatos de *consenso moderado*, que acumulan muchas aceptaciones, frente a candidatos más intensamente preferidos por un grupo pero considerados inaceptables por otros.

Ejemplo 2.1. *Veamos con un ejemplo las diferencias entre ambos enfoques. Consideremos tres alternativas $X = \{A, B, C\}$ y tres votantes. Las preferencias de los votantes son las siguientes:*

- *Votante 1:* $A \succ B \succ C$.
- *Votante 2:* $A \succ B \succ C$.
- *Votante 3:* $B \succ C \succ A$.
- *Votante 4:* $C \succ B \succ A$.
- *Votante 5:* $C \succ A \succ B$.

Escenario 1: método de Borda (ordinal). *Cada votante asigna 2 puntos a su primera opción, 1 punto a la segunda y 0 a la última. El resultado agregado es:*

- $A: 2 + 2 + 0 + 0 + 1 = 5$ puntos.
- $B: 1 + 1 + 2 + 1 + 0 = 5$ puntos.
- $C: 0 + 0 + 1 + 2 + 2 = 5$ puntos.

Por lo tanto, por este método, todas las opciones empatan.

Escenario 2: votación por aprobación. Supongamos ahora que cada votante acepta únicamente sus dos primeras opciones y no acepta la última. El recuento es:

- Votante 1: acepta A y B.
- Votante 2: acepta A y B.
- Votante 3: acepta B y C.
- Votante 4: acepta C y B.
- Votante 5: acepta C y A.

El total de veces que ha sido aceptable cada opción es:

- A: $1 + 1 + 0 + 0 + 1 = 3$ veces.
- B: $1 + 1 + 1 + 1 + 0 = 4$ veces.
- C: $0 + 0 + 1 + 1 + 1 = 3$ veces.

En este sistema el ganador es B, que ha sido aceptable por más votantes que el resto de alternativas, a pesar de no haber sido el elegido en el escenario de Borda.

Este ejemplo es sencillo, pero podemos ver que incluso en un caso simplificado se evidencia cómo el sistema ordinal de Borda puede mostrar un empate, mientras que la votación por aprobación premia a una opción de consenso, aceptable para una mayoría de votantes aunque no sea su favorita absoluta.

En resumen, los modelos *ordinales* capturan solamente la posición relativa de las alternativas en los rankings, mientras que los modelos de *aprobación* resaltan la aceptabilidad mínima, es decir, qué alternativas superan un umbral común. Ambos enfoques son útiles, pero también parciales: pueden conducir a resultados opuestos dependiendo del contexto. Precisamente de la interacción entre ambos surge la motivación del modelo de *preferencias aprobatorias*, que integra a la vez el orden y la aceptación para conseguir un análisis más equilibrado de la decisión colectiva.

Un segundo ejemplo sobre cómo un tipo de votación puede favorecer a un candidato pese a no ser el más querido por la mayoría, se encuentra en las elecciones presidenciales de Estados Unidos en el año 2000. En dicha votación, los principales candidatos fueron George W. Bush (Partido Republicano), Al Gore (Partido Demócrata) y Ralph Nader (Partido Verde). La polémica surgió por lo ocurrido en el estado de Florida, cuyo resultado terminó siendo decisivo para el ganador.

En Florida la diferencia final entre Bush y Gore fue de apenas unos cientos de votos, en un censo de millones. Sin embargo, el sistema electoral era de mayoría simple: cada persona marcaba un único candidato y el que tuviera más votos en el Estado se llevaba todos los electores. En este contexto, los votos obtenidos por Ralph Nader, aunque mucho menores en número, tuvieron mucho más peso. Muchos analistas sostienen que una parte considerable de quienes votaron por Nader habrían preferido a Gore antes que a Bush en una comparación directa, [10]. Sin embargo, como el sistema no permitía reflejar esa preferencia secundaria, esos votos quedaron contabilizados exclusivamente para el candidato verde.

El resultado fue que Bush ganó en Florida por un margen mínimo, lo que le dio la victoria final en el colegio electoral, pese a que Gore obtuvo aproximadamente medio millón más de votos a nivel nacional. La controversia se hizo más grande porque hubo disputas sobre el recuento, papeletas mal diseñadas y procesos judiciales que llegaron hasta la Corte Suprema. Al final, el sistema mayoritario favoreció a Bush en un Estado clave. Podemos ver, por tanto, cómo una regla de votación concreta puede favorecer a un candidato de manera decisiva.

Aplicaciones de las preferencias aprobatorias

El modelo de preferencias aprobatorias ha demostrado ser útil en diferentes contextos de toma de decisiones colectivas y análisis de datos. Al combinar simultáneamente información ordinal (jerarquía de alternativas) y evaluativa (umbrales de aceptabilidad), estas estructuras permiten capturar matices que se pierden en modelos más simples. A continuación veremos algunas aplicaciones.

Un primer contexto natural es el de las votaciones en comités, consejos o asambleas. En estos entornos, los participantes no solo deben establecer un orden de preferencia entre las alternativas (por ejemplo, distintos proyectos, candidatos o políticas), sino también declarar cuáles consideran suficientemente aceptables como para apoyarlas.

Un ejemplo destacado en este ámbito es el *Eurobarómetro*, donde se estudian regularmente las actitudes políticas y sociales de los ciudadanos europeos, [2]. En tales encuestas, la información de qué propuestas o instituciones son aceptables resulta esencial para distinguir entre candidatos que generan un amplio consenso mínimo (aceptabilidad) y aquellos que, aunque ocupen buenas posiciones en ciertos rankings, son inaceptables por grandes segmentos de la población. Esta capacidad de identificar consenso o polarización es particularmente relevante en contextos de gobernanza multinacional, donde la estabilidad de las decisiones depende de alcanzar un mínimo denominador común.

En el ámbito de los sistemas de recomendación (como los utilizados por plataformas de música, cine o comercio electrónico), no basta con ofrecer al usuario el ítem que maximiza su ranking de preferencias. Es crucial asegurar que las recomendaciones sean, al menos, *aceptables*.

El modelo de preferencias aprobatorias ofrece una solución natural. Esta consiste en dos vectores. Un vector ordinal que captura las jerarquías de gusto, y otro vector de aprobación que asegura que los ítems sugeridos superen un umbral de aceptabilidad. Esto es especialmente importante para evitar recomendar alternativas que, aunque podrían situarse en un rango intermedio del ranking del usuario, resulten en realidad inaceptables (por ejemplo, un género musical o literario que para el usuario sea explícitamente inaceptable). Al incorporar ambas dimensiones, los algoritmos pueden equilibrar diversidad y satisfacción mínima en las recomendaciones.

En procesos de negociación entre partes (por ejemplo, acuerdos laborales, tratados internacionales o mediaciones en conflictos), cada votante dispone de un conjunto de propuestas que valora de manera jerárquica, pero no todas las alternativas son igualmente viables. Lo esencial es identificar qué propuestas son aceptables, aunque no necesariamente las más deseadas.

Existen además múltiples ámbitos donde las preferencias aprobatorias pueden aportar valor añadido:

- **Asignación de recursos y proyectos:** en entornos como la investigación científica o la distribución de presupuestos, es habitual que los evaluadores expresen tanto un ranking de prioridades como un umbral de proyectos aceptables para financiación.

- **Selección en procesos educativos o laborales:** en admisiones universitarias o contrataciones, los evaluadores ordenan a los candidatos, pero también fijan un nivel mínimo de aceptabilidad. Este modelo ayuda a integrar ambas perspectivas.
- **Planificación colectiva y políticas públicas:** en la elaboración de planes de acción, los ciudadanos o los expertos pueden declarar qué medidas consideran aceptables, además de priorizarlas, lo que aporta información clave para detectar consensos sociales.

En definitiva, las preferencias aprobatorias ofrecen un marco flexible y versátil para representar juicios colectivos, ya que combinan jerarquías ordinales con umbrales de aceptabilidad. Esta doble dimensión enriquece el análisis de la diversidad, el consenso y la polarización. En los capítulos siguientes se profundizará en cómo estas estructuras permiten definir medidas de distancia, formular métricas de consenso y aplicar técnicas de agrupamiento tanto sobre individuos como sobre alternativas, con el fin de interpretar patrones colectivos en datos reales.

2.2. Definición formal

Antes de continuar con el análisis de similitudes, distancias y agrupamientos entre votantes, es necesario definir primero conceptos importantes en este contexto sobre preferencias aprobatorias.

Definición 2.2.1. Sea $X = \{x_1, x_2, \dots, x_n\}$ un conjunto finito de alternativas, con $n \geq 2$. Una relación de orden débil R sobre X es una relación binaria que cumple dos propiedades fundamentales:

- **Completitud:** para todo par $x_i, x_j \in X$, se cumple que $x_i R x_j$ o $x_j R x_i$ (o ambos).
- **Transitividad:** para cuales quiera $x_i, x_j, x_k \in X$, si $x_i R x_j$ y $x_j R x_k$, entonces $x_i R x_k$.

Definición 2.2.2. Denotemos por $W(X)$ al conjunto de todos los órdenes débiles sobre X . A partir de $R \in W(X)$ se definen dos relaciones asociadas:

- La **preferencia estricta:**

$$x_i \succ_R x_j \iff x_i R x_j \text{ y no } x_j R x_i.$$

- La **indiferencia:**

$$x_i \sim_R x_j \iff x_i R x_j \text{ y } x_j R x_i.$$

Además del orden relativo, nos interesa saber cuales son las alternativas que llamamos *aceptables* para un votante, es decir, las que considera una opción viable. Para ello, se introduce un subconjunto $A \subseteq X$, llamado *conjunto de alternativas aceptables*. La idea es que el votante establece un umbral de aceptabilidad y declara qué opciones superan dicho umbral. La combinación de estas dos dimensiones (orden débil y conjunto de alternativas aceptables) da lugar a la definición 2.2.3.

Definición 2.2.3. Denotamos por $\mathcal{P}(X)$ el conjunto de partes de X . Una preferencia aprobatoria sobre un conjunto X es un par $(R, A) \in W(X) \times \mathcal{P}(X)$ tal que se cumple la condición de consistencia:

$$(x_i R x_j \wedge x_j \in A) \implies x_i \in A, \quad \forall x_i, x_j \in X.$$

En particular, esta condición implica que:

- Si x_j es aceptada y $x_i \succ_R x_j$, entonces x_i también debe ser aceptada.
- Si x_j es aceptada y $x_i \sim_R x_j$, entonces x_i también debe ser aceptada.

Por tanto, la condición de consistencia asegura que, si una alternativa es aceptada, entonces todas las que sean estrictamente mejores o indiferentes con ella también lo son.

Definición 2.2.4. Denotamos por $R(X)$ al conjunto de todas las preferencias aprobatorias sobre X , es decir,

$$R(X) = \{(R, A) \in W(X) \times \mathcal{P}(X) : (R, A) \text{ cumple la condición de consistencia}\}.$$

Notación 2.1. Representaremos un elemento

$$(R, A) = (x_1 \succ_R \cdots \succ_R x_n \succ_R x_{n+1} \succ_R \cdots \succ_R x_m, \{x_1, \dots, x_n\}) \in R(X)$$

situando las alternativas aceptables en la parte superior, y las inaceptables en la parte inferior, separadas por una línea horizontal (umbral). A esta notación la llamaremos representación en bloques:

$$\begin{array}{c} x_1 \\ x_2 \\ \vdots \\ x_n \\ \hline x_{n+1} \\ \vdots \\ x_m \end{array}$$

donde $\{x_1, \dots, x_n\}$ son las alternativas aceptables (A) y $\{x_{n+1}, \dots, x_m\}$ las inaceptables ($X \setminus A$).

En cada bloque (arriba o abajo del umbral) las alternativas pueden aparecer en distintos niveles de preferencia de acuerdo con la relación R :

- si dos alternativas son indiferentes, se representan en la misma fila.
- si una es estrictamente preferida a otra, se colocan en filas sucesivas de arriba hacia abajo.

Así, la representación muestra simultáneamente el orden débil R y el subconjunto aprobado A , cumpliendo la condición de consistencia.

Veamos a continuación un ejemplo.

Ejemplo 2.2. Consideremos $X = \{x_1, x_2, x_3, x_4\}$ y la relación de orden débil

$$R = x_1 \succ_R x_2 \sim_R x_3 \succ_R x_4.$$

Si el conjunto aprobado es $A = \{x_1, x_2, x_3\}$, entonces (R, A) constituye una preferencia aprobatoria. La interpretación es la siguiente: el agente acepta x_1 , x_2 y x_3 , mientras que x_4 es inaceptable. Obsérvese que se cumple la condición de consistencia:

- como x_2 es aprobado, también lo es x_1 , que se encuentra por encima en el orden.
- como x_2 es aprobado, también lo es x_3 , que es indiferente en el orden.

En este caso, la representación en bloques sería la siguiente:

$$\begin{array}{c} x_1 \\ x_2 \quad x_3 \\ \hline x_4 \end{array}$$

Para entender mejor su alcance, resulta útil compararlas con otros modelos habituales de representación de preferencias en teoría de la elección social.

Órdenes lineales

Un *orden lineal* sobre X es un orden débil $R \in W(X)$ que es además antisimétrico, es decir,

$$x_i \sim_R x_j \implies x_i = x_j.$$

En este modelo, cada agente proporciona una jerarquía completa de todas las alternativas. La ventaja es que se obtiene una comparación exhaustiva, pero se pierde información sobre qué alternativas son aceptables y cuáles inaceptables en un sentido absoluto. Así, un candidato situado en última posición podría ser completamente inaceptable o simplemente considerado una opción débil, pero el orden lineal no distingue entre ambas situaciones.

Aprobación pura

En el modelo de *aprobación pura*, cada votante especifica únicamente un subconjunto $A \subseteq X$ de alternativas aceptadas, sin proporcionar información sobre el orden relativo entre ellas. Este modelo elimina la dimensión ordinal y conserva únicamente la información binaria sobre aceptabilidad, [6]. Su simplicidad lo hace atractivo en la práctica, pero presenta limitaciones claras, como que no permite discernir entre una alternativa favorita y otra que es apenas tolerada, siempre que ambas sean aceptables.

Juicio mayoritario y puntuaciones

En el *juicio mayoritario*, cada agente asigna a cada alternativa una valoración en una escala común de categorías lingüísticas (por ejemplo, excelente, bueno, aceptable, insuficiente), [4, 3]. El resultado se determina a partir de las distribuciones de estas valoraciones.

En los *modelos de puntuación*, cada agente asigna a cada alternativa un número dentro de un rango prefijado (por ejemplo, de 0 a 10). El resultado suele agregarse mediante medias o sumas.

Estos modelos introducen la noción de intensidad de preferencia. Sin embargo, la aceptabilidad de las alternativas no está estructurada de manera binaria ni sometida a una condición de consistencia, pues un candidato puede recibir puntuaciones dispersas sin que exista un umbral común de aprobación. Por tanto, son modelos más expresivos en cuanto a grados de preferencia, pero menos normativos en cuanto a la distinción entre lo que es aceptable y lo que no.

En conclusión, las preferencias aprobatorias generalizan y conectan dos modelos ampliamente estudiados en teoría de la elección social:

- **Órdenes clásicos:** si $A = X$ o $A = \emptyset$, entonces (R, A) equivale simplemente a un orden débil tradicional, ya que todas las alternativas son consideradas aceptables o inaceptables.

- **Voto Aprobatorio:** si se ignora el orden R y se conserva únicamente el conjunto aprobado $A \subseteq X$, cada votante se limita a indicar qué alternativas aceptables y cuáles inaceptables.

Esto las convierte en una herramienta flexible para estudiar cómo se forman consensos, dónde aparecen divisiones y hasta qué punto existe diversidad en las decisiones colectivas.

2.3. Propiedades lógicas de las preferencias aprobatorias

A partir de la definición 2.2.3, se obtienen de manera natural algunas propiedades sencillas, pero muy útiles, que nos ayudan a entender mejor cómo funcionan estas estructuras. Dichas propiedades garantizan que las preferencias sean coherentes y que no aparezcan contradicciones, pues lo que se acepta o no se transmite de manera lógica a lo largo del orden de las alternativas. Además, estas propiedades permiten establecer relaciones de simetría que más adelante utilizaremos para comparar distintos perfiles de votantes. De forma general, podemos destacar tres principios básicos:

- **Condición de consistencia:** ya explicada en la definición 2.2.3. Si una alternativa es aceptable, entonces cualquier otra que sea estrictamente mejor o al menos indiferente con respecto a ella también debe ser aceptable.
- **Herencia hacia abajo:** de manera dual, si una alternativa es inaceptable, todas las que se sitúan en posiciones inferiores o indiferentes en el orden también deben ser inaceptables.
- **Inversión:** dada una preferencia aprobatoria (R, A) , es posible definir su opuesto natural $(R^{-1}, X \setminus A)$, lo que establece una simetría que más adelante permitirá formalizar medidas de desacuerdo y de polarización.

Estas tres propiedades constituyen lo que denominamos *clausura estructural* de las preferencias aprobatorias, y serán fundamentales para los desarrollos posteriores en torno a distancias, consenso y agrupamiento.

Proposición 2.3.1 (Herencia hacia abajo). *Sea $(R, A) \in R(X)$ una preferencia aprobatoria. Si $x_j \notin A$ y $x_i \prec_R x_j$ o $x_i \sim_R x_j$, entonces $x_i \notin A$.*

Demostración. Si $x_j \notin A$ y supusiéramos que $x_i \in A$, entonces por la condición de consistencia tendríamos $x_j \in A$, lo que contradice la hipótesis. Por tanto, necesariamente $x_i \notin A$. El caso $x_i \sim_R x_j$ se deduce de forma análoga. $\square$

Proposición 2.3.2. *Sea $(R, A) \in R(X)$ una preferencia aprobatoria. Son equivalentes:*

- (a) **Condición de consistencia:** para todo $x_i, x_j \in X$,

$$(x_j \in A \wedge x_i R x_j) \implies x_i \in A.$$

Es decir, si x_j es aceptada, entonces cualquier x_i que sea mejor o indiferente que x_j (esto es, $x_i \succ_R x_j$ o $x_i \sim_R x_j$) también es aceptada.

- (b) **Herencia hacia abajo:** para todo $x_i, x_j \in X$,

$$(x_j \notin A \wedge x_j R x_i) \implies x_i \notin A,$$

donde $x_i \preceq_R x_j$ significa $x_j R x_i$ (equivalentemente, $x_i \prec_R x_j$ o $x_i \sim_R x_j$). Es decir, si x_j es inaceptable, entonces cualquier x_i peor o indiferente que x_j también es inaceptable.

Demostración. $(a \Rightarrow b)$. Supongamos (a). Sean $x_i, x_j \in X$ con $x_j \notin A$ y $x_i \preceq_R x_j$, esto es, $x_j R x_i$. Si por absurdo $x_i \in A$, entonces, aplicando (a) con x_j en el papel de “alternativa mejor o indiferente” respecto de x_i (pues $x_j R x_i$), obtendríamos $x_j \in A$, contradicción. Luego $x_i \notin A$.

$(b \Rightarrow a)$. Supongamos (b). Sean $x_i, x_j \in X$ con $x_j \in A$ y $x_i R x_j$, esto es, $x_j \preceq_R x_i$. Si por absurdo $x_i \notin A$, entonces, aplicando (b) con x_i como alternativa inaceptable y $x_j \preceq_R x_i$, concluiríamos $x_j \notin A$, contradicción. Por tanto, $x_i \in A$. $\square$

Por lo tanto, por el Teorema 2.3.2, se tiene que la condición de consistencia es equivalente a la herencia hacia abajo.

Proposición 2.3.3. *Sea $(R, A) \in R(X)$ una preferencia aprobatoria. Entonces A coincide exactamente con un conjunto superior en el orden R , es decir:*

$$A = \{x \in X \mid \exists y \in A \text{ con } x \succ_R y \text{ o } x \sim_R y\}.$$

Demostración. La condición de consistencia implica que toda alternativa mejor o indiferente que una aceptada también debe estar en A , por lo que A es un conjunto superior en R . Recíprocamente, todo elemento de A pertenece por definición al lado derecho de la igualdad. De este modo, se obtiene la caracterización deseada. $\square$

Demostración. Denotemos

$$\text{Up}_R(A) := \{x \in X \mid \exists y \in A \text{ con } x \succ_R y \text{ o } x \sim_R y\}.$$

Probaremos las dos inclusiones $\text{Up}_R(A) \subseteq A$ y $A \subseteq \text{Up}_R(A)$.

(i) $\text{Up}_R(A) \subseteq A$. Sea $x \in \text{Up}_R(A)$. Entonces existe $y \in A$ tal que $x \succ_R y$ o $x \sim_R y$, lo que equivale a $x R y$. Por la condición de consistencia de (R, A) , de $y \in A$ y $x R y$ se deduce $x \in A$. Luego $\text{Up}_R(A) \subseteq A$.

(ii) $A \subseteq \text{Up}_R(A)$. Sea $x \in A$. Como R es un orden débil (completo y transitivo), es en particular reflexivo, luego para todo $x \in X$ se tiene $x R x$, y por tanto $x \sim_R x$. Tomando $y = x \in A$, obtenemos que x verifica la condición que define $\text{Up}_R(A)$; por tanto $x \in \text{Up}_R(A)$. Así, $A \subseteq \text{Up}_R(A)$.

De (i) y (ii) concluimos $A = \text{Up}_R(A)$, como se quería. $\square$

Corolario 2.3.4. *Sea $(R, A) \in R(X)$. Entonces A es unión de clases de indiferencia completas y, más aún, existe una clase umbral C^* de $X/\sim_R$ tal que*

$$A = \bigcup_{C \in X/\sim_R, C \succeq_{\tilde{R}} C^*} C,$$

donde el orden en el cociente $X/\sim_R$ viene dado por

$$C \succeq_{\tilde{R}} D \iff x R y \text{ para cualesquiera } x \in C, y \in D.$$

En particular, A no corta ninguna clase de indiferencia: o bien contiene una clase entera, o bien la excluye por completo.

Demostración. (1) Orden inducido en el cociente).

La relación de indiferencia $\sim_R$ es una relación de equivalencia, de modo que el cociente $X/\sim_R$ está bien definido, siendo sus elementos son las clases de indiferencia $[x]_{\sim_R} = \{y \in X : y \sim_R x\}$.

Sobre este cociente definimos el orden inducido $\succeq_{\tilde{R}}$ por

$$[x]_{\sim_R} \succeq_{\tilde{R}} [y]_{\sim_R} \iff xRy.$$

Esta definición no depende de los representantes, pues si $x \sim_R x'$ y $y \sim_R y'$, entonces xRy implica $x'Ry'$ (por transitividad de R y propiedades de $\sim_R$). Además, como R es completo y transitivo, la relación $\succeq_{\tilde{R}}$ es un orden lineal sobre las clases de indiferencia.

(2) *A es unión de clases*).

Si $x \in A$ y $y \sim_R x$, entonces por la condición de consistencia $y \in A$. Así, A es unión de clases de $\sim_R$. Denotemos por

$$\mathcal{A}^\sim := \{ [x]_{\sim_R} \in X/\sim_R : x \in A \}$$

el conjunto de clases contenidas en A .

(3) *$\mathcal{A}^\sim$ es un conjunto superior*).

Sea $C \in \mathcal{A}^\sim$ y $D \in X/\sim_R$ con $D \succeq_{\tilde{R}} C$. Tome $x \in C \cap A$ y $y \in D$. De $D \succeq_{\tilde{R}} C$ se sigue yRx ; como $x \in A$, por consistencia $y \in A$, luego $D \in \mathcal{A}^\sim$. Por tanto, $\mathcal{A}^\sim$ es superior en el orden lineal $(X/\sim_R, \succeq_{\tilde{R}})$.

(4) *Estructura de conjuntos superiores en órdenes lineales finitos*).

Como X es finito, también lo es el cociente $X/\sim_R$. En un orden lineal finito, todo conjunto superior no vacío es la “cola” a partir de su mínima clase. Así, si $A = \emptyset$ entonces $\mathcal{A}^\sim = \emptyset$; si $A = X$ entonces $\mathcal{A}^\sim = X/\sim_R$; y en caso contrario existe una *clase umbral* C^* mínima en $\mathcal{A}^\sim$ tal que

$$\mathcal{A}^\sim = \{ C \in X/\sim_R : C \succeq_{\tilde{R}} C^* \}.$$

Puesto que A es la unión de las clases que contiene, obtenemos

$$A = \bigcup_{C \in \mathcal{A}^\sim} C,$$

lo que, en particular, muestra que A no corta ninguna clase de indiferencia.

Esto prueba la afirmación y, en particular, que A es unión de clases completas. $\square$

Una consecuencia de la definición de preferencias aprobatorias es que el conjunto $R(X)$ posee una serie de propiedades de clausura que lo hacen estable frente a operaciones naturales como la inversión del orden o la restricción a subconjuntos de alternativas. Estas propiedades, que denominamos *clausura estructural*, resultan fundamentales para el estudio de distancias y medidas de consenso en capítulos posteriores.

Proposición 2.3.5 (Clausura por restricción). *Sea $(R, A) \in R(X)$ una preferencia aprobatoria y sea $Y \subseteq X$ un subconjunto no vacío de alternativas. Definimos la restricción de (R, A) a Y como*

$$(R|_Y, A \cap Y),$$

donde $R|_Y$ es la relación inducida por R sobre Y . Entonces se cumple que $(R|_Y, A \cap Y) \in R(Y)$.

Demostración. Como R es un orden débil sobre X , su restricción $R|_Y$ es un orden débil sobre Y . Además, si $y_j \in A \cap Y$ y $y_i R|_Y y_j$, entonces $y_i R y_j$ en X . Por consistencia de (R, A) , se deduce que $y_i \in A$, y por tanto $y_i \in A \cap Y$. De ello se sigue que $(R|_Y, A \cap Y)$ cumple la condición de consistencia y pertenece a $R(Y)$. $\square$

Proposición 2.3.6 (Clausura por inversión). *Sea $(R, A) \in R(X)$ una preferencia aprobatoria. Definimos la preferencia inversa como*

$$(R, A)^{-1} = (R^{-1}, X \setminus A),$$

donde R^{-1} es el orden débil inverso de R . Entonces se cumple que $(R, A)^{-1} \in R(X)$.

Demostración. El inverso R^{-1} de un orden débil es también un orden débil. Sea $x_j \in X \setminus A$ y $x_i R^{-1} x_j$, lo cual equivale a $x_j R x_i$. Dado que $x_j \notin A$, la condición de consistencia en (R, A) implica que $x_i \notin A$. Por tanto, $x_i \in X \setminus A$. Así, $(R^{-1}, X \setminus A)$ cumple la condición de consistencia y pertenece a $R(X)$. $\square$

Observación 2.3.7. *La clausura por inversión es especialmente relevante para la definición de medidas de disensión: permite comparar perfiles de votantes no solo en términos de similitud, sino también en términos de oposición o simetría. Asimismo, la clausura por restricción es clave para el análisis de subconjuntos de alternativas, como ocurre al reducir la dimensión del problema en técnicas de clustering.*

2.4. Codificación de las preferencias aprobatorias

Desde el punto de vista computacional, es muy útil disponer de una representación numérica de las preferencias aprobatorias que permita aplicar técnicas de análisis cuantitativo, tales como el cálculo de distancias o la aplicación de algoritmos de agrupamiento. Para ello, cada preferencia aprobatoria $(R, A) \in R(X)$ puede codificarse mediante dos vectores, uno asociado al orden débil R y otro al conjunto de alternativas aceptables A .

Sea $R \in W(X)$ un orden débil sobre $X = \{x_1, \dots, x_n\}$. Definimos la función de posición de R como

$$PR : X \longrightarrow [1, n],$$

donde $PR(x_j)$ representa la *posición media* de la alternativa x_j dentro del orden R . Esta posición se calcula teniendo en cuenta dos factores:

- cuántas alternativas están estrictamente por debajo de x_j .
- cuántas son indiferentes a ella.

Definición 2.4.1. *Dado $R \in W(X)$ y $x_j \in X$, definimos*

$$PR(x_j) = n - \#\{x_k \in X : x_j \succ_R x_k\} - \frac{1}{2} \#\{x_k \in X \setminus \{x_j\} : x_j \sim_R x_k\}.$$

La interpretación es la siguiente:

- El primer término n es el número total de alternativas.
- El segundo término cuenta cuántas alternativas son estrictamente peores que x_j , es decir, cuántas se encuentran por debajo en el orden.
- El tercer término ajusta la posición en caso de que x_j sea indiferente con otras alternativas. En lugar de colocarlas en posiciones sucesivas, se asigna a todas la posición media del bloque de indiferencia.

Así, $PR(x_j) = 1$ significa que x_j ocupa la primera posición del orden (es la mejor alternativa), mientras que $PR(x_j) = n$ indica que es la peor. Valores intermedios reflejan la posición relativa teniendo en cuenta posibles empates.

Definición 2.4.2. *El vector ordinal asociado a un orden $R \in W(X)$ se define como*

$$p_R = (PR(x_1), PR(x_2), \dots, PR(x_n)) \in [1, n]^n.$$

Por otra parte, la información de aceptabilidad contenida en el conjunto $A \subseteq X$ puede representarse de forma binaria.

Definición 2.4.3. *Dado $A \subseteq X$, definimos la función*

$$IA : X \longrightarrow \{0, 1\}, \quad IA(x_j) = \begin{cases} 1, & \text{si } x_j \in A, \\ 0, & \text{si } x_j \notin A. \end{cases}$$

El vector de aprobación asociado a $A \subseteq X$ se define como

$$i_A = (IA(x_1), IA(x_2), \dots, IA(x_n)) \in \{0, 1\}^n.$$

Con ambas representaciones, cada preferencia aprobatoria $(R, A) \in R(X)$ puede codificarse como el par

$$(p_R, i_A).$$

Proposición 2.4.1. *La codificación (p_R, i_A) preserva toda la información de una preferencia aprobatoria (R, A) . Es decir, a partir de (p_R, i_A) puede reconstruirse de manera unívoca tanto el orden débil R como el conjunto aprobado A .*

Demostración. Sea $\Phi : R(X) \rightarrow [1, n]^n \times \{0, 1\}^n$ la aplicación dada por $\Phi(R, A) = (p_R, i_A)$. Basta con probar que Φ es inyectiva. Sea $(R_1, A_1), (R_2, A_2) \in R(X)$ tales que

$$(p_{R_1}, i_{A_1}) = (p_{R_2}, i_{A_2}).$$

Primero, como $i_{A_1} = i_{A_2}$, se tiene que para cada $x_j \in X$:

$$i_{A_1}(x_j) = 1 \iff x_j \in A_1, \quad i_{A_2}(x_j) = 1 \iff x_j \in A_2.$$

De aquí se deduce $A_1 = A_2$.

Segundo, como $p_{R_1} = p_{R_2}$, se cumple que para cada $x_i, x_j \in X$:

$$PR_1(x_i) < PR_1(x_j) \iff PR_2(x_i) < PR_2(x_j),$$

y

$$PR_1(x_i) = PR_1(x_j) \iff PR_2(x_i) = PR_2(x_j).$$

Por la definición de $PR(\cdot)$, la primera equivalencia implica que

$$x_i \succ_{R_1} x_j \iff x_i \succ_{R_2} x_j,$$

y la segunda que

$$x_i \sim_{R_1} x_j \iff x_i \sim_{R_2} x_j.$$

Por tanto, $R_1 = R_2$.

Concluimos, por tanto, que $(R_1, A_1) = (R_2, A_2)$. Esto demuestra la inyectividad de la codificación. $\square$

Ejemplo 2.3. Sea $X = \{x_1, x_2, x_3, x_4, x_5, x_6\}$. Un votante declara la siguiente preferencia aprobatoria:

$$x_1 \succ_R x_2 \sim_R x_3 \succ_R x_4 \sim_R x_5 \succ_R x_6, \quad A = \{x_1, x_2, x_3\}.$$

La representación en bloques es

$$\begin{array}{cc} & x_1 \\ x_2 & x_3 \\ \hline x_4 & x_5 \\ & x_6 \end{array},$$

donde las alternativas situadas por encima de la línea son las aceptables y las situadas por debajo son las inaceptables. Obsérvese que x_2 y x_3 comparten un mismo nivel de indiferencia, al igual que x_4 y x_5 .

Condición de consistencia. Se cumple que si una alternativa es aceptada, entonces todas las que son mejores o indiferentes también lo son. En efecto, como $x_2 \in A$ y $x_1 \succ_R x_2$, tenemos $x_1 \in A$; además, como $x_2 \in A$ y $x_2 \sim_R x_3$, se cumple $x_3 \in A$. Por tanto, (R, A) es una preferencia aprobatoria válida.

Cálculo de la codificación. Con $n = 6$, se obtienen las siguientes posiciones medias de las alternativas:

$$PR(x_1) = 1, \quad PR(x_2) = PR(x_3) = 2,5, \quad PR(x_4) = PR(x_5) = 4,5, \quad PR(x_6) = 6.$$

Por tanto,

$$p_R = (1, 2,5, 2,5, 4,5, 4,5, 6).$$

Dado que $A = \{x_1, x_2, x_3\}$, el vector indicador es:

$$i_A = (1, 1, 1, 0, 0, 0).$$

En conclusión, la preferencia aprobatoria queda codificada por el par:

$$(R, A) \longleftrightarrow (p_R, i_A) = ((1, 2,5, 2,5, 4,5, 4,5, 6), (1, 1, 1, 0, 0, 0)).$$

Observación 2.4.2. Esta codificación es fundamental porque permite trabajar con estructuras aprobatorias en un espacio vectorial, lo que posibilita definir distancias (Capítulo 3) y aplicar técnicas de análisis de datos como el clustering (Capítulo 4). En particular, el vector p_R facilita el uso de distancias ordinales (como la distancia de Kemeny), mientras que el vector i_A es la base para distancias binarias (como la distancia de Hamming).

Proposición 2.4.3. Sea $\pi : X \rightarrow X$ una permutación. Definimos la acción de π sobre una preferencia aprobatoria (R, A) por

$$\pi(R) := \{(\pi(x), \pi(y)) : (x, y) \in R\}, \quad \pi(A) := \{\pi(x) : x \in A\}.$$

Entonces, si $\Phi(R, A) = (p_R, i_A)$ denota la codificación, se cumple

$$\Phi(\pi(R), \pi(A)) = (p_{\pi(R)}, i_{\pi(A)}), \quad \text{con } p_{\pi(R)}(\pi(x)) = p_R(x), \quad i_{\pi(A)}(\pi(x)) = i_A(x) \quad \forall x \in X.$$

En particular, la codificación se transforma por permutación coordinada de componentes y, por tanto, es neutral respecto al renombrado de alternativas.

Demostración. Por definición, $x \pi(R) y \iff \pi^{-1}(x) R \pi^{-1}(y)$. Para PR calculado con R y PR^π con $\pi(R)$, se tiene:

$$PR^\pi(\pi(x)) = n - \#\{z : \pi(x) \succ_{\pi(R)} z\} - \frac{1}{2} \#\{z \neq \pi(x) : \pi(x) \sim_{\pi(R)} z\}.$$

Usando la definición de $\pi(R)$, los conjuntos contados son imágenes por π de los análogos con x bajo R , por lo que sus cardinales coinciden con los correspondientes para x en R . De aquí $PR^\pi(\pi(x)) = PR(x)$, es decir, $p_{\pi(R)}(\pi(x)) = p_R(x)$ para todo x . Del mismo modo, $i_{\pi(A)}(\pi(x)) = \mathbf{1}_{\pi(A)}(\pi(x)) = \mathbf{1}_A(x) = i_A(x)$. $\square$

Proposición 2.4.4. *Para todo $R \in W(X)$ y $x \in X$ se cumple $PR(x) \in [1, n]$. Además,*

$$\sum_{x \in X} PR(x) = \frac{n(n+1)}{2} \quad \text{y, por tanto,} \quad \frac{1}{n} \sum_{x \in X} PR(x) = \frac{n+1}{2}.$$

Demostración. Sea la partición de X en clases de indiferencia $E_1 \succeq_{\tilde{R}} \cdots \succeq_{\tilde{R}} E_k$ (de mejor a peor), con $s_r := |E_r|$ y $\sum_{r=1}^k s_r = n$. Para $x \in E_r$ se tiene

$$PR(x) = \sum_{t < r} s_t + \frac{s_r + 1}{2}.$$

Entonces,

- para $r = 1$, $PR(x) = (s_1 + 1)/2 \geq 1$.
- para $r = k$, $PR(x) = \sum_{t < k} s_t + (s_k + 1)/2 = n - \frac{s_k - 1}{2} \leq n$.

Por tanto se deduce que $1 \leq PR(x) \leq n$. Sumando sobre $x \in X$ y reagrupando por clases,

$$\sum_{x \in X} PR(x) = \sum_{r=1}^k \sum_{x \in E_r} \left(\sum_{t < r} s_t + \frac{s_r + 1}{2} \right) = \sum_{r=1}^k \left(s_r \sum_{t < r} s_t \right) + \sum_{r=1}^k s_r \frac{s_r + 1}{2}.$$

Para el primer sumando, se obtiene

$$\sum_{r=1}^k s_r \sum_{t < r} s_t = \sum_{1 \leq t < r \leq k} s_t s_r,$$

ya que al expandir se recorren todas las parejas (t, r) con $t < r$. Por simetría,

$$\sum_{t \neq r} s_t s_r = 2 \sum_{1 \leq t < r \leq k} s_t s_r,$$

y como

$$\left(\sum_{r=1}^k s_r \right)^2 = \sum_{r=1}^k s_r^2 + \sum_{t \neq r} s_t s_r,$$

se deduce

$$\sum_{r=1}^k s_r \sum_{t < r} s_t = \frac{1}{2} \left(n^2 - \sum_{r=1}^k s_r^2 \right).$$

Para el segundo sumando,

$$\sum_{r=1}^k s_r \frac{s_r + 1}{2} = \frac{1}{2} \sum_r s_r^2 + \frac{1}{2} \sum_r s_r = \frac{1}{2} \sum_r s_r^2 + \frac{n}{2}.$$

Sumando ambas expresiones:

$$\sum_{x \in X} PR(x) = \frac{1}{2} (n^2 - \sum_r s_r^2) + \frac{1}{2} \sum_r s_r^2 + \frac{n}{2} = \frac{n^2 + n}{2} = \frac{n(n+1)}{2}.$$

□

Proposición 2.4.5. Sea $(R, A) \in W(X) \times \mathcal{P}(X)$ y sea (p_R, i_A) su codificación, con A no vacío. Entonces las siguientes afirmaciones son equivalentes:

1. $(R, A) \in R(X)$.
2. A es unión de las primeras clases de indiferencia de R . Es decir, existe $r_0 \in \{1, \dots, k\}$ tal que $A = \bigcup_{r=1}^{r_0} E_r$.
3. Existe $\tau \in [1, n]$ tal que

$$A = \{x \in X : PR(x) \leq \tau\} \quad \text{equivalente a} \quad IA(x) = \mathbf{1}_{\{PR(x) \leq \tau\}} \quad \forall x \in X.$$

Demostración. (1) $\Rightarrow$ (2): Sea $E_1 \succeq_{\tilde{R}} \dots \succeq_{\tilde{R}} E_k$ la partición de X en clases de indiferencia de R . La consistencia implica que si $x \in A$ y $y \sim_R x$, entonces $y \in A$. Luego, cada clase E_r está completamente contenida en A o completamente fuera. Además, si $E_r \subseteq A$ y $s < r$, entonces todo $y \in E_s$ satisface $y \succ_R x$ para cualquier $x \in E_r$, por lo que $y \in A$. Por tanto, existe r_0 tal que $A = \bigcup_{r=1}^{r_0} E_r$.

(2) $\Rightarrow$ (3): Si $A = \bigcup_{r=1}^{r_0} E_r$, tomemos

$$\tau := \sum_{t < r_0} s_t + \frac{s_{r_0} + 1}{2}$$

(para $r_0 = 0$ basta tomar, por ejemplo, $\tau < \frac{1}{2}$). Para $x \in E_r$, $PR(x) = \sum_{t < r} s_t + \frac{s_r + 1}{2}$. Entonces $PR(x) \leq \tau$ si y solo si $r \leq r_0$, que solo ocurre si y solo si $x \in A$. Por lo tanto $A = \{x : PR(x) \leq \tau\}$.

(3) $\Rightarrow$ (1): Si $A = \{x : PR(x) \leq \tau\}$, sea $x_j \in A$ y $x_i R x_j$. Entonces $PR(x_i) \leq PR(x_j) \leq \tau$, de donde $x_i \in A$. Por tanto, (R, A) cumple la condición de consistencia. □

Observación 2.4.6. En la proposición 2.4.5 hemos supuesto que $A \neq \emptyset$, de modo que $r_0 \geq 1$. Sin embargo, si se permite $A = \emptyset$, el resultado sigue siendo válido. Para ello basta considerar el caso $r_0 = 0$, donde por convenio

$$\bigcup_{r=1}^0 E_r = \emptyset.$$

Esto corresponde a un votante que no acepta ninguna alternativa. De manera análoga, el caso extremo $r_0 = k$ da lugar a $A = X$, en el que todas las alternativas resultan aceptables. Por tanto, la caracterización es general y no requiere la hipótesis $A \neq \emptyset$.

La caracterización anterior muestra que, en la codificación, la consistencia equivale a que i_A sea el indicador de un “corte por umbral” sobre p_R . En términos prácticos, esta forma estructurada de A resultará útil más adelante al definir medidas de consenso y al diseñar procedimientos de preprocesado (por ejemplo, filtrado por umbrales) en el análisis empírico.

Observación 2.4.7. *La aplicación $\Phi : R(X) \rightarrow [1, n]^n \times \{0, 1\}^n$, $\Phi(R, A) = (p_R, i_A)$, es inyectiva, pero su imagen es un subconjunto propio, pues no todo par (p, i) es realizable. En efecto, p debe provenir de un orden débil (constancia por clases e incremento estricto entre clases) y i debe ser el indicador de un corte por umbral sobre p , como en la proposición 2.4.5.*

2.5. Tamaño del espacio de preferencias aprobatorias

El número de preferencias aprobatorias posibles sobre un conjunto de alternativas crece de manera mucho más rápida que en los modelos tradicionales de representación de preferencias. Mientras que los órdenes lineales y los órdenes débiles producen espacios ya de por sí grandes, el hecho de añadir la dimensión evaluativa de aceptabilidad multiplica exponencialmente la complejidad.

- El número de *órdenes lineales* (o permutaciones completas) sobre n alternativas es exactamente $n!$, [5].
- El número de *órdenes débiles* (o preórdenes completos) sobre n alternativas corresponde a los números de Fubini (o números de Bell ordenados), conocidos en combinatoria clásica.
- En cambio, el número de *preferencias aprobatorias* sobre n alternativas, denotado $|R(X)|$, incluye además todas las posibles elecciones de subconjuntos aprobados $A \subseteq X$, siempre sujetos a la condición de consistencia. Estudios recientes (véase [9, 1]) han mostrado que este número crece de forma combinatorial mucho más intensa que los órdenes débiles, debido a la multiplicación entre las posibles particiones de X en bloques indiferentes y la elección de un umbral de aceptabilidad compatible.

En otras palabras, las preferencias aprobatorias contienen simultáneamente la riqueza estructural de los órdenes débiles y la flexibilidad de la aprobación binaria, lo que genera un espacio de representaciones exponencialmente mayor.

Para ilustrar el crecimiento, consideremos los primeros valores de n . En la tabla 2.1 se comparan el número de órdenes lineales, órdenes débiles y preferencias aprobatorias para $n = 2, \dots, 10$, [1].

n	Aprobaciones 2^n	Órdenes lineales $n!$	Órdenes débiles	Preferencias aprobatorias $ R(X) $
2	4	2	3	8
3	8	6	13	44
4	16	24	75	308
5	32	120	541	2 612
6	64	720	4 683	25 988
7	128	5 040	47 293	296 564
8	256	40 320	545 835	3 816 548
9	512	362 880	7 087 261	54 667 412
10	1 024	3 628 800	102 247 563	862 440 068

Tabla 2.1: Número de estructuras posibles según el modelo de representación de preferencias, para distintos valores de n , [1].

El número de órdenes débiles y de preferencias aprobatorias es conocido en forma general y puede obtenerse para cualquier valor de n , [11].

Se puede observar que mientras los órdenes lineales crecen factorialmente y los órdenes débiles lo hacen de forma aún más acelerada, el tamaño del espacio de preferencias aprobatorias aumenta todavía más rápidamente. Por ejemplo, con tan solo $n = 5$ alternativas, el espacio de preferencias aprobatorias ya supera las dos mil posibilidades distintas, frente a las 541 de los órdenes débiles y las 120 de los órdenes lineales.

El rápido crecimiento del espacio de preferencias aprobatorias tiene dos consecuencias principales:

1. **Complejidad computacional:** los algoritmos que operan sobre estos espacios (por ejemplo, los que calculan distancias o ejecutan procedimientos de agregación) deben estar diseñados para manejar dominios muy grandes. El uso de representaciones vectoriales, como se ha explicado en la sección anterior, resulta indispensable para reducir la complejidad práctica de estos cálculos.
2. **Capacidad expresiva:** este crecimiento refleja la enorme variedad de juicios que los individuos pueden expresar en un modelo aprobatorio. Frente a la rigidez de un orden lineal o la simplicidad de una votación por aprobación, las preferencias aprobatorias permiten capturar una gama mucho más rica de actitudes, desde el consenso parcial hasta la inaceptabilidad colectiva, pasando por situaciones de fuerte polarización.

En conclusión, el espacio $R(X)$ tiene un crecimiento enorme desde el punto de vista combinatorio, incluso en problemas con pocas alternativas. Esto hace difícil tratar el problema y motiva el uso de herramientas matemáticas y computacionales que permitan *resumir*, *comparar* y *agrupar* preferencias. En particular:

- en el capítulo 3 se introducirán medidas de distancia que permiten cuantificar el desacuerdo entre dos preferencias aprobatorias, combinando información ordinal y de aprobación, siguiendo ideas como las de [9];
- en el capítulo 4 se estudiarán técnicas de agrupamiento (clustering) que, a partir de dichas distancias, posibilitan analizar perfiles numerosos sin necesidad de recorrer todo el espacio de preferencias [2].

De este modo, la complejidad combinatoria se podrá tratar mediante distancias para reducir la comparación a números manejables y algoritmos de clustering para detectar patrones de consenso, diversidad o polarización.

Capítulo 3

Distancias entre preferencias aprobatorias

En el capítulo 2 se introdujo el marco teórico de las *preferencias aprobatorias*, definiendo formalmente estas estructuras, sus propiedades lógicas y su representación computacional. Dicho marco nos permite describir cómo los individuos expresan simultáneamente un orden de preferencia y un umbral de aceptabilidad sobre un conjunto de alternativas.

El siguiente paso natural es preguntarse cómo *comparar* dos preferencias aprobatorias. En contextos de decisión colectiva no basta con conocer las opiniones individuales de forma aislada, sino que es esencial medir el grado de *similitud o disimilitud* entre distintos agentes. Estas comparaciones son la base para estudiar fenómenos como el consenso, la disensión, la diversidad o la polarización en grupos sociales, políticos o estratégicos. Para ello introduciremos en este capítulo distintas nociones de *distancia*, entendidas como funciones que miden la diferencia entre dos estructuras aprobatorias. Dichas distancias se apoyan en la codificación presentada en el capítulo 2, distinguiendo entre la componente ordinal y la componente binaria de aceptación.

El capítulo se organiza de la siguiente forma. En la sección 3.1 se expone la motivación general y las propiedades deseables de una distancia en este contexto. En las secciones 3.2 y 3.3 se presentan, respectivamente, distancias que actúan únicamente sobre la parte de aprobación y sobre la parte ordinal. Posteriormente, en la sección 3.4 se describe una primera propuesta de distancia combinada mediante una combinación convexa de ambas componentes. Finalmente, en la sección 3.5 se introduce la familia de distancias D_λ^r , que permite integrar de manera flexible la información ordinal y evaluativa.

Estas herramientas constituirán la base matemática para el capítulo 4, dedicado al análisis *clúster*, donde dichas distancias se emplearán para agrupar individuos o alternativas en función de su grado de similitud.

3.1. Motivación y contexto

En el capítulo anterior se introdujo el concepto de preferencia aprobatoria como un par $(R, A) \in W(X) \times P(X)$, donde R representa un orden débil sobre el conjunto finito de alternativas X y $A \subseteq X$ indica el subconjunto de alternativas consideradas aceptables. La combinación de estas dos dimensiones (información ordinal y evaluativa) nos permite describir de forma más realista las opiniones de los individuos en un proceso de decisión colectiva.

Una vez fijada la estructura, surge la cuestión de cómo *comparar* dos preferencias aprobatorias $(R_1, A_1), (R_2, A_2) \in R(X)$. Comparar dos preferencias no es solo un ejercicio teórico,

ya que nos podemos encontrar en tal escenario en contextos sociales, políticos o estratégicos. Por esta razón es fundamental que los responsables de la toma de decisiones puedan evaluar hasta qué punto los individuos coinciden, discrepan o incluso se polarizan en sus votos. Para ello es imprescindible contar con una medida cuantitativa de la *similitud* o *disimilitud* entre preferencias, que permita:

- **Medir consenso social:** si la distancia entre todas las preferencias de un grupo es pequeña, podemos interpretar que existe un núcleo común de aceptación que facilita la toma de decisiones colectivas.
- **Analizar diversidad y disensión:** las distancias grandes entre preferencias reflejan posturas alejadas o incluso opuestas, lo que permite detectar grupos con visiones contrarias dentro de un mismo colectivo.
- **Aplicar técnicas de agrupamiento:** tal como se estudiará en el capítulo 4, los algoritmos de análisis *clúster* necesitan haber definido una distancia para agrupar individuos con opiniones similares o, alternatively, para identificar bloques de alternativas que reciben patrones de apoyo comparables.

Definición 3.1.1. Una distancia sobre el espacio de preferencias aprobatorias $R(X)$ es una función

$$d : R(X) \times R(X) \longrightarrow [0, \infty)$$

que mide el grado de disimilitud entre dos preferencias $(R_1, A_1), (R_2, A_2) \in R(X)$.

En la práctica, no siempre se exige que d cumpla todas las propiedades de una distancia. Recordemos que una función d es una *distancia* si satisface, para todo $(R_1, A_1), (R_2, A_2), (R_3, A_3) \in R(X)$:

1. $d((R_1, A_1), (R_2, A_2)) \geq 0$ (no negatividad).
2. $d((R_1, A_1), (R_2, A_2)) = 0 \iff (R_1, A_1) = (R_2, A_2)$.
3. $d((R_1, A_1), (R_2, A_2)) = d((R_2, A_2), (R_1, A_1))$ (simetría).
4. $d((R_1, A_1), (R_3, A_3)) \leq d((R_1, A_1), (R_2, A_2)) + d((R_2, A_2), (R_3, A_3))$ (desigualdad triangular).

En algunos casos basta con relajar la condición (2), obteniendo así una *pseudodistancia*. Este tipo de flexibilización es útil en contextos donde dos preferencias distintas pueden ser indistinguibles desde el punto de vista de la distancia escogida.

El diseño de distancias adecuadas para estructuras mixtas como las preferencias aprobatorias no es trivial, ya que deben equilibrar dos componentes heterogéneas: la parte ordinal y la parte de aprobación. En la literatura reciente se han desarrollado diversas aproximaciones para combinar información de este tipo (véase, por ejemplo, [9]). En esencia, la construcción de distancias en este marco tiene dos objetivos complementarios:

- **Fidelidad representativa.** La distancia debe reflejar con precisión las discrepancias tanto en el orden de las alternativas como en el conjunto aprobado.
- **Operatividad computacional.** El cálculo de distancias debe ser factible en escenarios con un gran número de alternativas o votantes, de modo que pueda integrarse en algoritmos de clustering o en medidas de consenso.

Una vez justificada la necesidad de disponer de medidas de distancia en el contexto de las preferencias aprobatorias, el siguiente paso consiste en clasificar dichas medidas según la componente sobre la que actúan. En primer lugar, en la Sección 3.2 consideraremos distancias que se definen únicamente sobre la parte binaria de aprobación. A continuación, en la Sección 3.3 analizaremos distancias que trabajan exclusivamente sobre la parte ordinal. Posteriormente, en la Sección 3.4, veremos cómo combinar ambas perspectivas en una única medida mediante una combinación convexa. Finalmente, en la Sección 3.5, presentaremos la familia de distancias D_λ^r , que constituye un enfoque unificado y flexible para capturar simultáneamente discrepancias ordinales y aprobatorias.

3.2. Distancias sobre la parte de aprobación

En esta sección nos centraremos en la componente binaria de una preferencia aprobatoria, esto es, en el conjunto $A \subseteq X$ de alternativas aceptadas. Como se definió en el capítulo 2, esta parte puede representarse mediante la función característica

$$i_A : X \longrightarrow \{0, 1\}, \quad i_A(x) = \begin{cases} 1 & \text{si } x \in A, \\ 0 & \text{si } x \notin A. \end{cases}$$

De este modo, cada preferencia aprobatoria (R, A) puede codificarse, en lo relativo a la aprobación, como un vector binario de dimensión $n = |X|$.

Comparar dos preferencias aprobatorias desde esta perspectiva equivale, por tanto, a comparar dos vectores binarios. Existen diferentes medidas de distancia en este contexto, siendo las más habituales la distancia de Hamming y la distancia de Jaccard.

Definición 3.2.1. Sea $u = (u_1, \dots, u_n)$ y $v = (v_1, \dots, v_n)$ dos vectores en $\{0, 1\}^n$. La distancia de Hamming entre u y v se define como

$$d_H(u, v) = \sum_{k=1}^n |u_k - v_k|.$$

Es decir, $d_H(u, v)$ cuenta el número de posiciones en las que u y v difieren.

En nuestro contexto, cada aprobación $A \subseteq X$ asociada a una preferencia aprobatoria $(R, A) \in R(X)$ puede representarse mediante el vector binario $i_A = (i_A(x_1), \dots, i_A(x_n)) \in \{0, 1\}^n$, donde

$$i_A(x) = \begin{cases} 1 & \text{si } x \in A, \\ 0 & \text{si } x \notin A. \end{cases}$$

De este modo, para dos preferencias aprobatorias (R_1, A_1) y (R_2, A_2) definimos

$$d_H((R_1, A_1), (R_2, A_2)) = d_H(i_{A_1}, i_{A_2}) = \sum_{x \in X} |i_{A_1}(x) - i_{A_2}(x)|.$$

La distancia de Hamming entre aprobaciones mide, por tanto, el número de alternativas para las que los individuos discrepan en su aceptabilidad.

Ejemplo 3.1. Consideremos el conjunto de alternativas $X = \{x_1, x_2, x_3, x_4\}$. Sean (R_1, A_1) y (R_2, A_2) dos preferencias aprobatorias con la siguiente representación en bloques:

$$(R_1, A_1) = \frac{\begin{matrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{matrix}}{\quad} \quad (R_2, A_2) = \frac{\begin{matrix} x_2 \\ x_1 \ x_3 \\ x_4 \end{matrix}}{\quad}$$

Por tanto,

$$A_1 = \{x_1, x_2\}, \quad A_2 = \{x_2, x_1, x_3\}.$$

Los vectores de aprobación correspondientes son

$$i_{A_1} = (1, 1, 0, 0), \quad i_{A_2} = (1, 1, 1, 0).$$

Si nos centramos únicamente en la parte binaria, la distancia de Hamming entre estas aprobaciones es

$$d_H((R_1, A_1), (R_2, A_2)) = |1 - 1| + |1 - 1| + |0 - 1| + |0 - 0| = 1.$$

Esto significa que las dos preferencias difieren en una sola alternativa, que en nuestro caso es x_3 , ya que es aceptable en A_2 pero inaceptable en A_1 . En el resto de alternativas coinciden.

Definición 3.2.2 (Distancia de Jaccard general). Sean $A, B \subseteq X$ dos subconjuntos de un conjunto finito X . La distancia de Jaccard entre A y B se define como

$$d_J(A, B) = 1 - \frac{|A \cap B|}{|A \cup B|}.$$

Por convenio, $d_J(A, B) = 0$ si $A = B = \emptyset$.

Observación 3.2.1. La distancia de Jaccard toma valores en $[0, 1]$. Dos aprobaciones son idénticas si y sólo si $d_J = 0$, mientras que $d_J = 1$ indica que no hay ninguna alternativa aceptable en común.

La distancia de Jaccard proporciona una medida relativa de disimilitud, normalizando respecto al tamaño de la unión. En el contexto de las preferencias aprobatorias, si (R_1, A_1) y (R_2, A_2) son dos preferencias en $R(X)$, definimos

$$d_J((R_1, A_1), (R_2, A_2)) = d_J(A_1, A_2) = 1 - \frac{|A_1 \cap A_2|}{|A_1 \cup A_2|}.$$

Ejemplo 3.2. Retomemos las preferencias aprobatorias del ejemplo 3.1 sobre $X = \{x_1, x_2, x_3, x_4\}$, donde $A_1 = \{x_1, x_2\}$ e $A_2 = \{x_1, x_2, x_3\}$. Aplicando la definición, obtenemos

$$d_J((R_1, A_1), (R_2, A_2)) = 1 - \frac{|A_1 \cap A_2|}{|A_1 \cup A_2|} = 1 - \frac{|\{x_1, x_2\}|}{|\{x_1, x_2, x_3\}|} = 1 - \frac{2}{3} = \frac{1}{3}.$$

Observación 3.2.2. Ambas distancias reflejan la disimilitud entre los conjuntos, aunque en escalas diferentes, pues Hamming mide en unidades absolutas de discrepancia, mientras que Jaccard lo hace de forma relativa al tamaño de la unión.

Ambas distancias presentan propiedades útiles en el análisis de datos:

- La distancia de Hamming es adecuada cuando interesa contabilizar el número exacto de discrepancias en la aceptación de alternativas.

- La distancia de Jaccard resulta más apropiada cuando los tamaños de los conjuntos aprobados son muy diferentes, ya que normaliza por la unión y proporciona una medida relativa de similitud.

En conclusión, estas medidas sobre la parte binaria permiten estudiar la cercanía entre preferencias aprobatorias solo fijándonos en el umbral de aceptabilidad, sin considerar la información ordinal. Como veremos en las secciones siguientes, este enfoque resulta limitado cuando las preferencias difieren en el orden relativo de las alternativas pero no en el conjunto aprobado, lo que hace necesario introducir distancias sobre la parte ordinal (Sección 3.3) y posteriormente la combinación de ambas perspectivas (Sección 3.4).

3.3. Distancias sobre la parte ordinal

Pasemos ahora a definir distancias sobre la otra componente: la parte ordinal. El problema se reduce a comparar dos órdenes (que pueden incluir empates). Para este fin existen diversas distancias en la literatura sobre teoría de rankings, entre las que destacan la distancia de Kendall-Tau modificada y la distancia de Kemeny.

Distancia de Kendall–Tau

Al analizar únicamente la parte ordinal de las preferencias aprobatorias, una herramienta clásica en la comparación de órdenes es el coeficiente de *Kendall–Tau*, introducido originalmente en estadística para medir el grado de correlación entre dos rankings. Este coeficiente se ha utilizado de forma extensa en análisis de rankings y teoría de decisión, pues proporciona una medida intuitiva de la concordancia relativa entre dos órdenes.

Dado un conjunto X de n alternativas y dos órdenes lineales (sin empates) R_1 y R_2 , consideramos todos los pares no ordenados de alternativas $\{x, y\} \subseteq X$. Decimos que un par es

- *concordante* si R_1 y R_2 coinciden en el orden relativo de x e y ,
- *discordante* si R_1 y R_2 invierten dicho orden.

Denotando por C el número de pares concordantes y por D el número de discordantes, con $C + D = \binom{n}{2}$, el coeficiente de Kendall–Tau se define como

$$\tau(R_1, R_2) = \frac{C - D}{\binom{n}{2}}.$$

Por construcción, τ toma valores en el intervalo $[-1, 1]$, ya que vale 1 cuando los órdenes coinciden plenamente, -1 cuando son opuestos, y 0 cuando no existe correlación lineal entre ambos. De esta expresión se obtiene fácilmente que

$$\tau = 1 - \frac{2D}{\binom{n}{2}},$$

lo cual muestra la relación entre τ y el número de pares discordantes D . Este último es, en realidad, una medida de *distancia*, ya que cuenta cuántas inversiones son necesarias para transformar un orden en el otro.

Definición 3.3.1. *Dados dos órdenes lineales R_1 y R_2 , definimos*

$$d_{KT}(R_1, R_2) := D,$$

el número de pares discordantes. Su versión normalizada en $[0, 1]$ es

$$\tilde{d}_{KT}(R_1, R_2) := \frac{D}{\binom{n}{2}} = \frac{1 - \tau(R_1, R_2)}{2}.$$

En el contexto de las preferencias aprobatorias, los órdenes pueden presentar empates debido a los órdenes débiles. Para tratar esta situación, se utiliza la generalización conocida como *Kendall–Tau-b*, que ajusta el denominador para tener en cuenta los pares empatados en uno u otro orden. Denotando por T_1 el número de pares empatados en R_1 pero no en R_2 , y T_2 de forma análoga, se define

$$\tau_b(R_1, R_2) = \frac{C - D}{\sqrt{(C + D + T_1)(C + D + T_2)}}.$$

De nuevo, $\tau_b \in [-1, 1]$ y coincide con τ cuando no hay empates. A partir de este coeficiente se construye una distancia ordinal normalizada que llamaremos *distancia de Kendall–Tau modificada*, que es la más conveniente de aplicar en el contexto de las preferencias aprobatorias.

Definición 3.3.2. *La distancia de Kendall–Tau-modificada entre dos órdenes R_1, R_2 se define como*

$$d_{KT}(R_1, R_2) := \frac{1 - \tau_b(R_1, R_2)}{2}.$$

Ejemplo 3.3. *Retomemos el ejemplo 3.1. En la parte ordinal, los órdenes inducidos son*

$$R_1 : x_1 \succ x_2 \succ x_3 \succ x_4, \quad R_2 : x_2 \succ (x_1 \sim x_3) \succ x_4.$$

Existen $n = \binom{4}{2} = 6$ pares de alternativas a comparar. Clasificándolos:

- (x_1, x_2) : R_1 ordena $x_1 \succ x_2$, mientras que R_2 ordena $x_2 \succ x_1 \Rightarrow$ **discordante**.
- (x_1, x_3) : R_1 ordena $x_1 \succ x_3$, mientras que R_2 los sitúa como $x_1 \sim x_3 \Rightarrow$ **cuenta como empate frente a orden**.
- (x_1, x_4) : ambos órdenes coinciden en $x_1 \succ x_4 \Rightarrow$ **concordante**.
- (x_2, x_3) : ambos coinciden en $x_2 \succ x_3 \Rightarrow$ **concordante**.
- (x_2, x_4) : ambos coinciden en $x_2 \succ x_4 \Rightarrow$ **concordante**.
- (x_3, x_4) : ambos coinciden en $x_3 \succ x_4 \Rightarrow$ **concordante**.

En total, tenemos $C = 4$ pares concordantes, $D = 1$ discordante y $T_2 = 1$ par con empate en R_2 pero no en R_1 . No hay ningún empate en R_1 , por lo que $T_1 = 0$. Aplicando la definición 3.3.2, tenemos que

$$\tau_b(R_1, R_2) = \frac{C - D}{\sqrt{(C + D + T_1)(C + D + T_2)}} = \frac{4 - 1}{\sqrt{(5 + 0 + 0)(5 + 1 + 0)}} = \frac{3}{\sqrt{30}} \approx 0,547.$$

Finalmente, la distancia normalizada es

$$d_{KT}(R_1, R_2) = \frac{1 - \tau_b(R_1, R_2)}{2} = \frac{1 - 0,547}{2} \approx 0,226.$$

Este valor indica que los dos órdenes son relativamente próximos (pues d_{KT} es bajo), aunque no idénticos, ya que existe una inversión clara en el par (x_1, x_2) y una discrepancia adicional por el empate en (x_1, x_3) .

Distancia de Kemeny

Otra medida clásica de para medir la diferencia en la parte ordinal es la *distancia de Kemeny*. En órdenes lineales (sin empates) se interpreta como el número de pares de alternativas cuyo orden relativo difiere en los dos rankings. Para órdenes *débiles* (con empates), una formulación habitual consiste en contar, para cada par, si se mantiene el orden, se invierte o hay o no un empate, [9].

Notación 3.1. Cuando escribimos sumatorios del tipo

$$\sum_{1 \leq i < j \leq n} f(i, j),$$

entendemos que i y j recorren el conjunto $\{1, \dots, n\}$ y que cada par no ordenado $\{i, j\}$ con $i \neq j$ se cuenta una única vez, imponiendo la condición $i < j$. De este modo, el número total de términos es exactamente $\binom{n}{2}$, correspondiente al número de pares distintos de elementos de $\{1, \dots, n\}$.

Definición 3.3.3. Sea $X = \{x_1, \dots, x_n\}$, con $n \geq 2$, un conjunto de alternativas finito y para cada orden débil $R \in W(X)$, sea $PR : X \rightarrow [1, n]$ la función de posición media vista en la definición 2.4.1. Denotamos por $\text{sgn} : \mathbb{R} \rightarrow \{-1, 0, 1\}$ la función signo,

$$\text{sgn}(a) = \begin{cases} 1, & a > 0, \\ 0, & a = 0, \\ -1, & a < 0. \end{cases}$$

Para $R_1, R_2 \in W(X)$ definimos la distancia de Kemeny como

$$\tilde{d}_K(R_1, R_2) := \sum_{1 \leq i < j \leq n} \left| \text{sgn}(PR_1(x_i) - PR_1(x_j)) - \text{sgn}(PR_2(x_i) - PR_2(x_j)) \right|.$$

De manera más visual, cada par no ordenado $\{x_i, x_j\}$ aporta:

juicio en R_1	juicio en R_2	contribución
igual	igual	0
estricto	empate	1
empate	estricto	1
estricto ($x_i \succ x_j$)	estricto ($x_j \succ x_i$)	2

En otras palabras, la distancia de Kemeny mide en qué medida las comparaciones par a par entre alternativas coinciden, se contradicen o se ven alteradas entre los dos órdenes.

Proposición 3.3.1. Para todo $R_1, R_2 \in W(X)$,

$$0 \leq d_K(R_1, R_2) \leq 2 \binom{n}{2} = n(n-1).$$

La cota superior se alcanza, por ejemplo, cuando R_1 y R_2 son órdenes lineales inversos (para cada par hay inversión estricta y cada par contribuye 2).

Demostración. Sea $X = \{x_1, \dots, x_n\}$ y sea, para $k = 1, 2$, $PR_k : X \rightarrow [1, n]$ la función de posiciones medias asociada al orden débil R_k . Para cada par no ordenado $\{x_i, x_j\}$ con $i < j$, definimos

$$a_{ij} := \text{sgn}(PR_1(x_i) - PR_1(x_j)), \quad b_{ij} := \text{sgn}(PR_2(x_i) - PR_2(x_j)),$$

de modo que $a_{ij}, b_{ij} \in \{-1, 0, 1\}$. Por definición,

$$\tilde{d}_K(R_1, R_2) = \sum_{1 \leq i < j \leq n} |a_{ij} - b_{ij}|.$$

(i) *Cota inferior.* Para todo $i < j$, $|a_{ij} - b_{ij}| \geq 0$, luego $d_K(R_1, R_2) \geq 0$ por ser suma de términos no negativos. Además, si $R_1 = R_2$ entonces $a_{ij} = b_{ij}$ para todo $i < j$ y, por tanto, $d_K(R_1, R_2) = 0$. En consecuencia, 0 es la cota inferior y es alcanzable.

(ii) *Cota superior.* Dado que $a_{ij}, b_{ij} \in \{-1, 0, 1\}$, los posibles valores de $|a_{ij} - b_{ij}|$ son $\{0, 1, 2\}$, y en particular $|a_{ij} - b_{ij}| \leq 2$ para todo $i < j$. Sumando sobre los $\binom{n}{2}$ pares distintos (cada uno considerado una única vez mediante la condición $i < j$), obtenemos

$$d_K(R_1, R_2) = \sum_{1 \leq i < j \leq n} |a_{ij} - b_{ij}| \leq \sum_{1 \leq i < j \leq n} 2 = 2 \binom{n}{2} = n(n-1).$$

□

Definición 3.3.4. Para $R_1, R_2 \in W(X)$ definimos la distancia de Kemeny normalizada como

$$d_K(R_1, R_2) := \frac{\tilde{d}_K(R_1, R_2)}{n(n-1)} \in [0, 1].$$

En este trabajo utilizaremos la versión *normalizada* de la distancia de Kemeny, puesto que facilita la comparación entre órdenes de distinto tamaño y su combinación con otras distancias también acotadas en $[0, 1]$, como la distancia de Hamming.

Ejemplo 3.4. Retomemos el ejemplo 3.1. La parte ordinal es

$$R_1 : x_1 \succ x_2 \succ x_3 \succ x_4, \quad R_2 : x_2 \succ x_1 \sim x_3 \succ x_4.$$

Las posiciones medias son

$$p_{R_1} = (1, 2, 3, 4), \quad p_{R_2} = (2, 5, 1, 2, 5, 4).$$

Para los $\binom{4}{2} = 6$ pares $\{x_i, x_j\}$ con $i < j$:

$\{x_i, x_j\}$	$\text{sgn}(PR_1(x_i) - PR_1(x_j))$	$\text{sgn}(PR_2(x_i) - PR_2(x_j))$	aportación
(x_1, x_2)	-1	+1	2
(x_1, x_3)	-1	0	1
(x_1, x_4)	-1	-1	0
(x_2, x_3)	-1	-1	0
(x_2, x_4)	-1	-1	0
(x_3, x_4)	-1	-1	0

Sumando, obtenemos

$$\tilde{d}_K(R_1, R_2) = 2 + 1 = 3.$$

Por tanto,

$$d_K(R_1, R_2) = \frac{3}{n(n-1)} = \frac{3}{12} = \frac{1}{4}.$$

3.4. Combinación convexa de distancias

En las secciones anteriores hemos definido distancias que actúan únicamente sobre una de las dos componentes de una preferencia aprobatoria $(R, A) \in R(X)$: las distancias entre los vectores ordinales p_R y las distancias entre los vectores de aprobación i_A . Aunque son útiles en ciertos contextos, estas distancias son parciales, ya que ignoran parte esencial de la información.

El objetivo es integrar simultáneamente las dimensiones ordinal y aprobatoria. Una forma sencilla y ampliamente utilizada en el análisis de datos consiste en realizar una combinación convexa de dos distancias previamente normalizadas.

Definición 3.4.1. Sean d_R una distancia definida sobre la parte ordinal y d_A una distancia definida sobre la parte binaria de aprobación. La distancia combinada convexa se define como

$$d_k((R_1, A_1), (R_2, A_2)) = k \cdot d_R(R_1, R_2) + (1 - k) \cdot d_A(A_1, A_2),$$

donde $k \in (0, 1)$ es un parámetro de ponderación.

Por construcción, d_k es una combinación convexa de d_R y d_A . Para que la definición sea coherente es necesario que ambas distancias estén previamente *normalizadas* al intervalo $[0, 1]$, de modo que d_k también tome valores en dicho rango.

Proposición 3.4.1. Si d_R y d_A son distancias normalizadas en $[0, 1]$, entonces d_k es también una distancia para todo $k \in (0, 1)$.

Demostración. La no negatividad y la simetría se heredan inmediatamente de d_R y d_A . Si $d_k((R_1, A_1), (R_2, A_2)) = 0$, entonces necesariamente $d_R(R_1, R_2) = 0$ y $d_A(A_1, A_2) = 0$, lo cual implica $R_1 = R_2$ y $A_1 = A_2$. Por tanto $(R_1, A_1) = (R_2, A_2)$, lo que prueba la propiedad de separación. Finalmente, la desigualdad triangular se cumple porque

$$d_k(P_1, P_3) \leq k(d_R(R_1, R_2) + d_R(R_2, R_3)) + (1 - k)(d_A(A_1, A_2) + d_A(A_2, A_3)),$$

y aplicando la desigualdad triangular de d_R y d_A se obtiene el resultado. $\square$

El parámetro k permite ajustar la importancia relativa de cada dimensión:

- Si $k \approx 1$, la distancia d_k se aproxima a d_R , se premian las discrepancias en el orden de las alternativas por encima de la información de aprobación.
- Si $k \approx 0$, la distancia d_k se aproxima a d_A , dando más peso a la coincidencia o diferencia en los conjuntos aprobados.
- Para valores intermedios, d_k equilibra ambas componentes.

Esta flexibilidad convierte a d_k en una herramienta versátil para distintos escenarios. En entornos donde el orden detallado de las alternativas es relevante (por ejemplo, en clasificaciones académicas o competiciones deportivas), conviene elegir k alto. En cambio, en contextos donde la aceptabilidad mínima es prioritaria (como en procesos políticos o negociaciones), puede ser preferible un k bajo.

Observación 3.4.2. La distancia d_k cumple propiedades deseables en el ámbito de la elección social, como neutralidad (no depende del renombrado de alternativas) y reciprocidad (simetría entre los agentes comparados), [9].

En resumen, la distancia combinada convexa d_k constituye un primer modelo unificado para capturar la similitud entre preferencias aprobatorias, integrando en una sola expresión la dimensión ordinal y la evaluativa. En la siguiente sección se introducirá una familia más sofisticada de distancias, D_λ^r , que generaliza esta idea mediante medias de potencia ponderadas, proporcionando un control más fino sobre la manera en que se agregan las discrepancias.

3.5. Distancia D_λ^r entre preferencias aprobatorias

La distancia combinada convexa d_k definida en la sección 3.4 es una primera forma de integrar la componente ordinal y la aprobatoria en una única expresión. Sin embargo, tiene una estructura lineal, la cual puede resultar demasiado rígida en ciertos contextos, al asignar un peso fijo k a cada componente. En este apartado presentamos una familia más general de distancias, que se basa en la noción de *discordancia entre pares de alternativas* y permite una agregación más flexible de ambas dimensiones, [1].

Sea $X = \{x_1, \dots, x_n\}$ un conjunto de alternativas finito y $(R_1, A_1), (R_2, A_2) \in R(X)$ dos preferencias aprobatorias. Para cada par de alternativas (x_i, x_j) con $i < j$, se consideran dos tipos de discordancia: posicional y de aprobación.

Definición 3.5.1. La discordancia posicional entre x_i y x_j relativa a dos órdenes R_1 y R_2 se define como

$$p_{ij} = \frac{|PR_1(x_i) - PR_1(x_j)| - |PR_2(x_i) - PR_2(x_j)|}{n - 1},$$

donde $PR_k : X \rightarrow [1, n]$ es la función de la definición 2.4.1.

En términos intuitivos, p_{ij} mide cuánto difiere la percepción relativa entre x_i y x_j en los dos órdenes considerados. Valores cercanos a 0 indican que los individuos colocan a x_i y x_j a distancias similares, mientras que valores altos reflejan un desacuerdo mayor respecto a su posición relativa.

Definición 3.5.2. La discordancia de aprobación entre x_i y x_j relativa a dos conjuntos de aprobaciones A_1, A_2 se define como

$$a_{ij} = |(IA_1(x_i) - IA_1(x_j)) - (IA_2(x_i) - IA_2(x_j))|,$$

donde IA es la función de la definición 2.4.3.

El valor a_{ij} toma valores en $\{0, 1\}$ y mide si existe discrepancia en el estatus de aceptabilidad relativa de x_i y x_j en ambas preferencias. Si ambos individuos coinciden en aceptar (o no aceptar) a x_i frente a x_j , entonces $a_{ij} = 0$; en caso contrario, $a_{ij} = 1$.

Una vez obtenidas las discordancias posicional p_{ij} y de aprobación a_{ij} , se combinan en un único valor a través de una media de potencia ponderada. Dado $r \geq 1$ y un parámetro de ponderación $\lambda \in (0, 1)$, se define

$$h(x, y) = (\lambda x^r + (1 - \lambda)y^r)^{1/r}, \quad x, y \in [0, 1]. \quad (3.5.1)$$

La función h es continua, simétrica y creciente en ambas variables. El parámetro λ determina la importancia relativa de las discordancias posicionales frente a las de aprobación, mientras que r controla la curvatura de la media. Para $r = 1$ se obtiene la media aritmética ponderada, mientras que para $r \rightarrow \infty$ se tiende al máximo ponderado de x y y .

Definición 3.5.3. La distancia D_λ^r entre dos preferencias aprobatorias $(R_1, A_1), (R_2, A_2) \in R(X)$ se define como

$$D_\lambda^r((R_1, A_1), (R_2, A_2)) = \frac{2}{n(n-1)} \sum_{1 \leq i < j \leq n} h(p_{ij}, a_{ij}).$$

De este modo, D_λ^r se obtiene como la media de las discordancias entre todos los pares de alternativas, combinando en cada caso la información posicional y la aprobatoria mediante h . La normalización por $\frac{2}{n(n-1)}$ garantiza que $D_\lambda^r \in [0, 1]$.

Proposición 3.5.1. Para $r \geq 1$ y $\lambda \in (0, 1)$, la aplicación D_λ^r es una distancia sobre $R(X)$.

Demostración. La no negatividad y simetría provienen de la definición de h . La identidad del indiscernible se cumple porque $D_\lambda^r((R_1, A_1), (R_2, A_2)) = 0$ implica $p_{ij} = 0$ y $a_{ij} = 0$ para todo par, lo cual equivale a $R_1 = R_2$ y $A_1 = A_2$. Finalmente, la desigualdad triangular se deriva de que h es una distancia L^r ponderada y la suma sobre los pares preserva la propiedad. $\square$

Observación 3.5.2. Como podemos observar en (3.5.1), necesitamos dos parámetros para construir h :

- El parámetro λ determina la importancia relativa de la parte ordinal y la parte aprobatoria. Si $\lambda = 0,5$, ambas dimensiones se consideran igualmente relevantes. Valores cercanos a 1 enfatizan el orden posicional, mientras que valores cercanos a 0 enfatizan la aceptación binaria.
- El parámetro r controla la sensibilidad de la distancia ante discrepancias grandes. Para $r = 1$ se obtiene un promedio lineal de discordancias; para $r > 1$ las discrepancias más elevadas tienen mayor peso relativo. En particular, en el límite $r \rightarrow \infty$, $h(x, y) = \max\{x, y\}$.

Ejemplo 3.5. Sea $X = \{x_1, x_2, x_3, x_4, x_5\}$. Consideramos las siguientes dos preferencias aprobatorias:

$$(R_1, A_1) : \begin{array}{c} x_1 \\ x_2 \ x_3 \\ x_4 \\ x_5 \end{array} \quad (R_2, A_2) : \begin{array}{c} x_2 \\ x_1 \\ x_3 \ x_4 \\ x_5 \end{array}$$

Se obtienen los siguientes vectores ordinales y de aprobación:

$$p_{R1} = (1, 2, 5, 2, 5, 4, 5), \quad p_{R2} = (2, 1, 3, 5, 3, 5, 5),$$

$$i_{A1} = (1, 1, 1, 0, 0), \quad i_{A2} = (1, 1, 0, 0, 0),$$

Para cada par $i < j$, definimos la discordancia posicional y aprobatoria, vistas en las definiciones 3.5.1 y 3.5.2 respectivamente.

(i, j)	$ PR_1(x_i) - PR_1(x_j) $	$ PR_2(x_i) - PR_2(x_j) $	p_{ij}	a_{ij}
(1, 2)	1,5	1,0	0,125	0
(1, 3)	1,5	1,5	0,000	$\frac{1}{2}$
(1, 4)	3,0	1,5	0,375	0
(1, 5)	4,0	3,0	0,250	0
(2, 3)	0,0	2,5	0,625	$\frac{1}{2}$
(2, 4)	1,5	2,5	0,250	0
(2, 5)	2,5	4,0	0,375	0
(3, 4)	1,5	0,0	0,375	$\frac{1}{2}$
(3, 5)	2,5	1,5	0,250	$\frac{1}{2}$
(4, 5)	1,0	1,5	0,125	0

Dada la media de potencia ponderada (3.5.1), la distancia total es

$$D_{\lambda}^r((R_1, A_1), (R_2, A_2)) = \frac{2}{n(n-1)} \sum_{1 \leq i < j \leq n} h(p_{ij}, a_{ij}) = \frac{1}{10} \sum_{1 \leq i < j \leq n} h(p_{ij}, a_{ij}).$$

Veamos ejemplos con algunos parámetros.

Caso 1: $r = 1$, $\lambda = \frac{1}{2}$. Aquí $h(x, y) = \frac{1}{2}(x + y)$. Sumando los h de la tabla:

$$\sum_{1 \leq i < j \leq n} h(p_{ij}, a_{ij}) = \frac{1}{2} \sum_{1 \leq i < j \leq n} (p_{ij} + a_{ij}) = 2,375.$$

Finalmente obtenemos

$$D_{0,5}^1((R_1, A_1), (R_2, A_2)) = 0,2375.$$

Caso 2: efecto de λ (con $r = 1$ fijo). El gráfico 3.1 muestra la evolución de la distancia D_{λ}^1 al variar $\lambda \in [0, 1]$. Se observa un crecimiento lineal como cabría esperar con r fijo. Cuando se otorga mayor peso a la parte ordinal, la distancia aumenta, pues las discrepancias posicionales en este ejemplo son más significativas que las de aprobación. Por ejemplo,

$$D_{0,2}^1 \approx 0,215, \quad D_{0,8}^1 \approx 0,260.$$

Esto confirma que las inversiones y cambios en la separación relativa de las alternativas tienen un impacto notable en el desacuerdo global cuando se prioriza la dimensión ordinal.

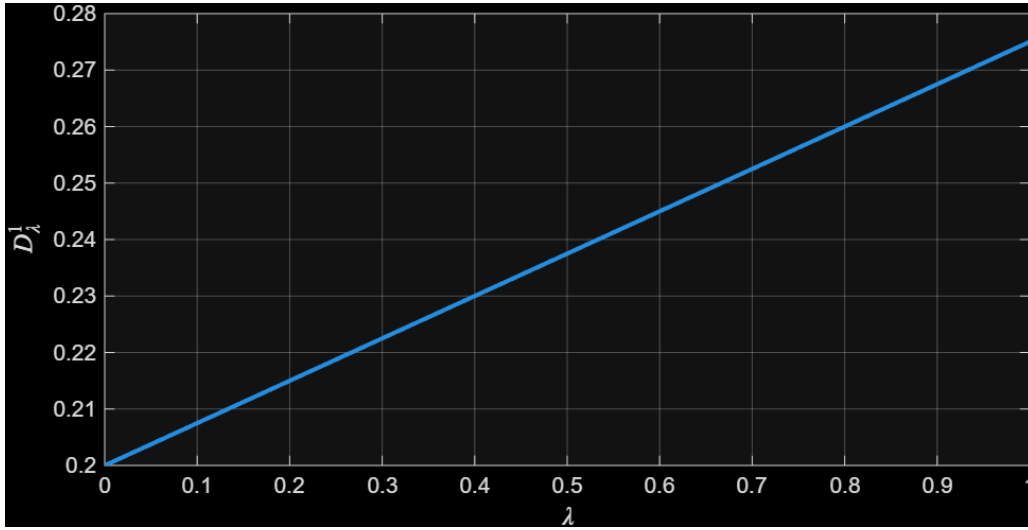
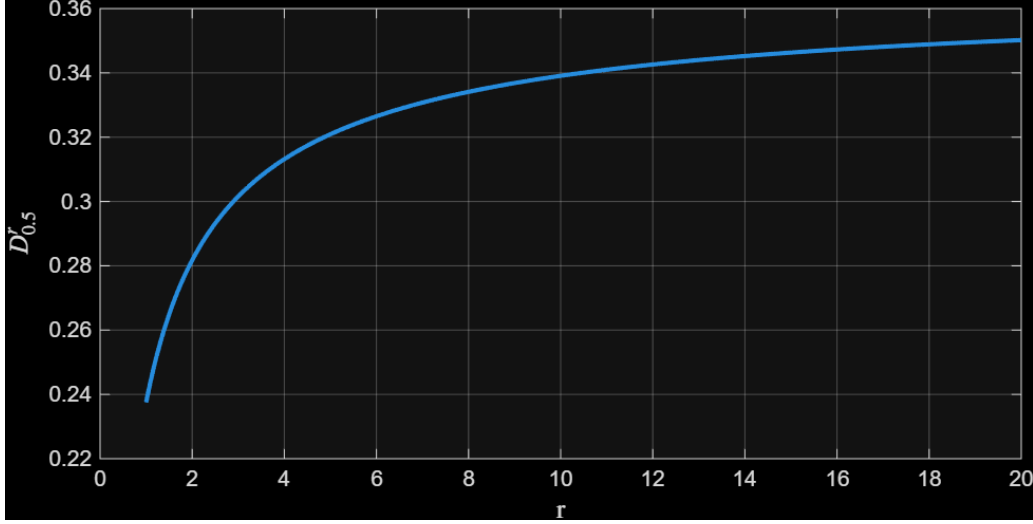


Gráfico 3.1: Evolución de D_{λ}^1 en función de λ .

Caso 3: efecto de r (curvatura, con $\lambda = 0,5$ fijo). El gráfico 3.2 muestra la distancia $D_{0,5}^r$ en función de r . Se aprecia un crecimiento monótono que tiende a estabilizarse. A medida que aumenta r , la media de potencia ponderada se aproxima al máximo entre las discordancias posicionales y de aprobación, penalizando de forma más severa los pares con mayor desacuerdo. Concretamente,

$$D_{0,5}^2 \approx 0,2817, \quad D_{0,5}^4 \approx 0,3132, \quad D_{0,5}^{10} \approx 0,3391.$$

De este modo, el parámetro r actúa como un regulador de la sensibilidad de la distancia, amplificando el peso de las discrepancias más grandes en la comparación global.

Gráfico 3.2: Evolución de $D_{0,5}^r$ en función de r .

Caso 4: visión conjunta ($D_{\lambda}^r(\lambda, r)$) mediante mapa de calor. El gráfico 3.3 muestra el valor de la distancia D_{λ}^r entre estas dos preferencias en función de los parámetros $\lambda \in [0, 1]$ (eje horizontal) y $r \in [1, 20]$ (eje vertical). El color representa el valor de la distancia y las curvas blancas son isocuantas, es decir, conjuntos de valores (λ, r) que producen la misma distancia.

- **Efecto de λ (peso relativo).** Para r fijo, D_{λ}^r crece linealmente con λ . Esto se debe a que en este ejemplo las discrepancias ordinales son más relevantes que las aprobatorias, ya que cuanto más peso se da al orden, mayor es la distancia. Si λ es cercano a 0 o a 1, una de las dos fuentes de discordancia se ignora, y la distancia es algo menor.
- **Efecto de r (curvatura).** Para λ fijo, D_{λ}^r aumenta con r , aunque con rendimientos decrecientes. Esto significa que al crecer r la media de potencia ponderada se acerca al máximo entre las discordancias posicionales y aprobatorias, penalizando con más fuerza los pares con mayor discrepancia, aunque el incremento adicional se suaviza.
- **Interacción (λ, r) .** Los valores más altos de la distancia aparecen cuando se da importancia al orden (λ intermedio o grande) y simultáneamente se elige un r alto, lo cual refleja que el desacuerdo entre estas dos preferencias es más notorio si se priorizan las discrepancias ordinales y, además, se penalizan con intensidad las más grandes.

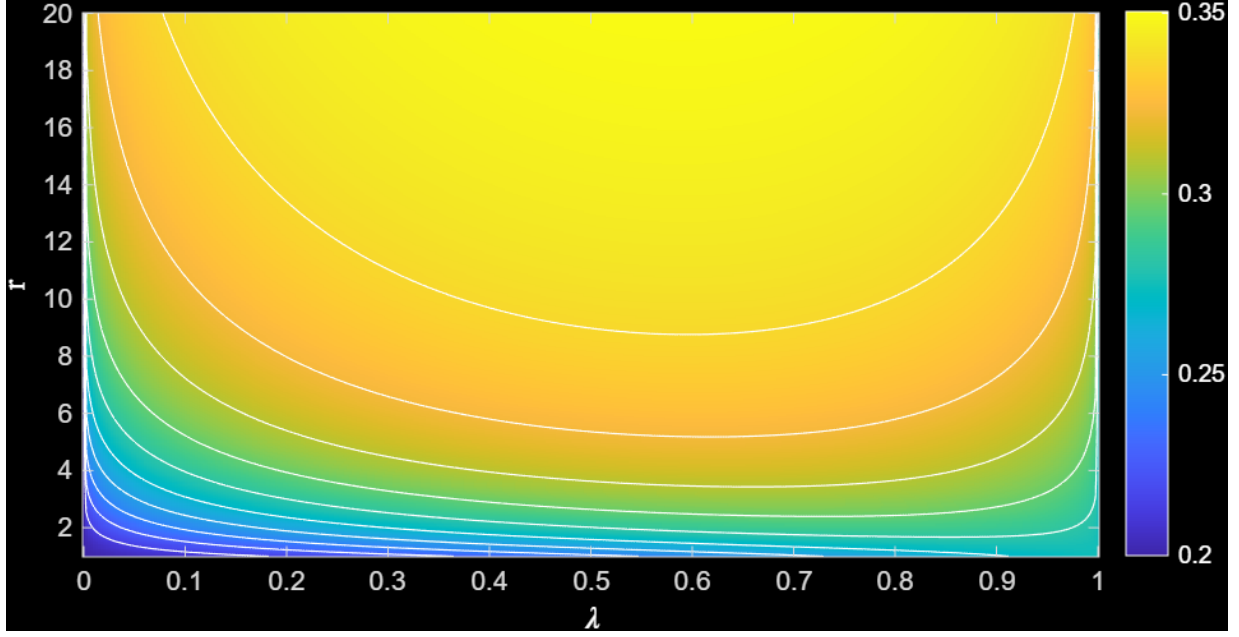


Gráfico 3.3: Mapa de calor de D_λ^r en función de λ y r , con isocuantas.

Esta dependencia de los parámetros es importante de cara al Capítulo 4, pues muestra que, según los valores de λ y r , los mismos datos pueden considerarse más próximos o más lejanos, lo que repercutirá en los resultados de los métodos de agrupamiento. La familia de distancias D_λ^r generaliza la métrica convexa d_k y proporciona un marco flexible para medir similitudes entre preferencias aprobatorias. En la práctica, el valor de λ se elige en función de la importancia relativa que se quiera asignar al orden o a la aprobación, mientras que r se ajusta según la sensibilidad deseada frente a discrepancias extremas.

En el análisis empírico de este trabajo utilizaremos principalmente las distancias D_λ^r con $r = 1$ y $\lambda = 0,5$, ya que ofrecen un compromiso equilibrado entre ambas dimensiones y son más fáciles de interpretar. Además, estas distancias constituyen la entrada fundamental para los algoritmos de análisis clúster, permitiendo agrupar individuos o alternativas según el grado de similitud en sus preferencias. De este modo, el marco teórico desarrollado aquí servirá para identificar patrones de consenso, diversidad y polarización en los datos.

Capítulo 4

Métodos de análisis clúster

En los capítulos anteriores se han introducido las *preferencias aprobatorias* y las distintas distancias que permiten medir la disimilitud entre los votantes o entre las alternativas. Este capítulo se centra en el uso de dichas distancias como base para aplicar *métodos de análisis clúster*, es decir, técnicas de aprendizaje no supervisado que buscan identificar estructuras en los datos a partir de relaciones de similitud.

El análisis cluster constituye una herramienta especialmente útil en el contexto de preferencias aprobatorias, ya que permite:

- explorar la *diversidad* de opiniones presentes en un colectivo.
- detectar posibles *bloques homogéneos* de votantes con preferencias similares.
- identificar fenómenos de *polarización*, cuando los individuos se agrupan en subconjuntos bien diferenciados.
- reconocer situaciones de *consenso parcial*, cuando surgen clústeres mayoritarios con grado bajo de disimilitud interna.

El capítulo se organiza de la siguiente manera. En la Sección 4.1 se presentan los principales métodos de agrupamiento aplicables en este contexto, que son técnicas particionales como *K-means* y sus variantes, algoritmos jerárquicos basados en dendrogramas, y propuestas específicas como el *Ranked K-Medoids (RKM)*, diseñadas para trabajar directamente con distancias definidas sobre preferencias aprobatorias. Además, se mencionan brevemente otras técnicas alternativas, como DBSCAN o el clustering espectral.

Posteriormente, en la sección 4.2 se discute la *elección de la distancia y del método* en función de los objetivos del análisis, resaltando cómo diferentes distancias capturan distintos aspectos de la estructura de las preferencias.

Más adelante, en el capítulo 5, aplicaremos de manera empírica los métodos de análisis cluster a un conjunto de datos real, con el fin de estudiar la dispersión, la diversidad y la polarización de las preferencias aprobatorias observadas.

4.1. Tipos de métodos de agrupamiento

Existen diversos métodos de análisis clúster, cada uno con fundamentos, supuestos y limitaciones distintas. Todos ellos tienen como objetivo dividir un conjunto de elementos en grupos internos lo más homogéneos posible, y lo más heterogéneos posible entre sí, a partir de alguna medida de similitud o distancia.

En el caso de las *preferencias aprobatorias*, estas técnicas se pueden aplicar gracias a las representaciones numéricas que hemos introducido en los capítulos anteriores. A partir de la codificación de cada votante (o de cada alternativa) y de las distancias definidas entre dichas representaciones, es posible construir una matriz de disimilitudes sobre la cual aplicar los algoritmos de agrupamiento. De esta forma, se obtiene una segmentación del conjunto de votantes o alternativas donde se puede llegar a ver la estructura interna de las opiniones, facilitando el análisis de fenómenos como la polarización, el consenso parcial o la existencia de bloques de preferencias similares.

En esta sección se presentan algunos de los métodos de agrupamiento más relevantes en el contexto de las preferencias aprobatorias. En primer lugar, describiremos técnicas de partición como *K-means* y sus variantes, adaptadas al uso de distancias no euclídeas. A continuación, introduciremos los métodos jerárquicos, que permiten explorar la estructura de similitud a diferentes niveles. Posteriormente, presentaremos el algoritmo *Ranked K-Medoids (RKM)*, el cual está diseñado para trabajar directamente con las pseudodistancias asociadas a las preferencias aprobatorias. Finalmente, comentaremos brevemente otras alternativas que, aunque menos habituales en este ámbito, ofrecen perspectivas complementarias.

K-means

Uno de los métodos de análisis cluster más conocidos es el algoritmo *K-means*, una técnica de agrupamiento particional muy utilizada en el aprendizaje no supervisado. Su objetivo principal es dividir un conjunto de observaciones en k grupos (o clústeres), de modo que se minimice la variabilidad interna dentro de cada grupo y se maximice la diferencia entre grupos distintos. En términos generales, K-means busca que los elementos de un mismo cluster estén lo más próximos posible a un centro representativo, denominado *centroide*, mientras que los centroides de diferentes clústeres se encuentren lo más alejados posible entre sí.

Definición 4.1.1. Sea $X = \{x_1, \dots, x_n\}$ un conjunto de datos en un espacio vectorial euclídeo con una distancia $d(\cdot, \cdot)$ asociada (usualmente la distancia euclídea). El algoritmo *K-means* tiene como objetivo encontrar una partición $C = \{C_1, \dots, C_k\}$ de X que minimice la suma de las varianzas intra-clúster, es decir:

$$\min_C \sum_{j=1}^k \sum_{x_i \in C_j} d(x_i, \mu_j)^2,$$

donde μ_j es el centroide del cluster C_j , definido como el promedio de los elementos que lo componen.

El procedimiento estándar del algoritmo K-means se desarrolla en los siguientes pasos:

1. **Inicialización:** se seleccionan aleatoriamente k puntos del conjunto X que actuarán como centroides iniciales.
2. **Asignación:** cada observación x_i se asigna al cluster cuyo centroide esté más próximo, es decir, $x_i \in C_j$ si $d(x_i, \mu_j)$ es mínimo.
3. **Actualización:** se recalculan los centroides μ_j como el promedio de todos los elementos asignados al cluster C_j .
4. **Iteración:** los pasos de asignación y actualización se repiten hasta que los centroides no cambian significativamente, o bien hasta alcanzar un número máximo de iteraciones.

Proposición 4.1.1. *El algoritmo K-means converge en un número finito de pasos a una partición localmente óptima de los datos, en el sentido de que el criterio de varianza intra-cluster no puede reducirse más mediante reasignaciones locales de puntos a centroides.*

Aunque K-means es intuitivo y eficiente en espacios vectoriales euclídeos, tiene algunas limitaciones importantes en el contexto de preferencias aprobatorias. En este caso, las observaciones corresponden a preferencias codificadas que no son parte de un espacio vectorial euclídeo, sino de estructuras combinatorias donde el concepto de promedio de dos preferencias carece de sentido. Por ejemplo, calcular el “promedio” entre dos órdenes de preferencia no conduce a una nueva preferencia válida.

Para solucionar este problema, es útil considerar variantes del método que se apoyen únicamente en una función de distancia, sin necesidad de promediar elementos. En este sentido, el algoritmo *K-medoids* es una adaptación natural.

Definición 4.1.2. *Sea $X = \{x_1, \dots, x_n\}$ un conjunto de observaciones y $d(\cdot, \cdot)$ una función de distancia definida sobre X . El algoritmo K-medoids busca una partición $C = \{C_1, \dots, C_k\}$ y un conjunto de representantes $\{m_1, \dots, m_k\} \subset X$ tales que se minimice:*

$$\min_{C, \{m_j\}} \sum_{j=1}^k \sum_{x_i \in C_j} d(x_i, m_j).$$

Cada representante m_j (medoid) es un elemento real del conjunto X , lo que evita el problema de definir centroides abstractos en espacios no euclídeos.

En el marco de las preferencias aprobatorias, esta adaptación resulta especialmente adecuada. Las opiniones de los votantes pueden representarse mediante vectores binarios (aceptable o inaceptable) combinados con órdenes parciales o totales, y las distancias pueden definirse mediante métricas como la distancia de Hamming o distancias específicas como D_λ^r , que combinan aspectos de aprobación y de ordenación. De este modo, es posible aplicar K-medoids sobre la matriz de disimilitudes entre agentes o entre alternativas, respetando así datos discretos, [2].

En resumen, mientras que el algoritmo K-means sigue el principio general de agrupamiento mediante minimización de varianza intra-clúster, en el contexto de preferencias aprobatorias es necesario recurrir a variantes basadas en distancias, como K-medoids. Estas permiten utilizar funciones de disimilitud diseñadas específicamente para este tipo de preferencias, como D_λ^r , asegurando que los grupos obtenidos reflejen de manera fiel la estructura subyacente en los datos.

Clustering jerárquico

Otro método de análisis cluster es el *clustering jerárquico*, que tiene como objetivo construir una estructura anidada de particiones que describe las relaciones de similitud entre los elementos a diferentes niveles de granularidad. A diferencia de los métodos particionales como K-means o K-medoids, que necesitan fijar de antemano el número de grupos k , los métodos jerárquicos generan una *jerarquía de clústeres* que puede visualizarse mediante un diagrama denominado *dendrograma*. Esto es especialmente útil en situaciones exploratorias, donde no se conoce a priori cuántos grupos pueden existir o donde interesa analizar la estructura interna de los datos a múltiples escalas.

Definición 4.1.3. Sea $X = \{x_1, \dots, x_n\}$ un conjunto de elementos y $d(\cdot, \cdot)$ una función de distancia definida sobre X . Un procedimiento de clustering jerárquico construye una secuencia anidada de particiones

$$\{\{x_1\}, \dots, \{x_n\}\} = \mathcal{P}_n \succ \mathcal{P}_{n-1} \succ \dots \succ \mathcal{P}_1 = \{X\},$$

donde $\mathcal{P}_k$ denota una partición en k clústeres y la relación $\succ$ indica que cada cluster en $\mathcal{P}_k$ es subconjunto de algún cluster en $\mathcal{P}_{k-1}$.

Existen dos grandes familias de métodos jerárquicos:

- **Métodos aglomerativos:** parten de n clústeres iniciales (cada elemento por separado) y se van fusionando iterativamente los dos clústeres más cercanos según un criterio de enlace (*single-link*, *complete-link*, *average-link*, entre otros). El proceso termina cuando todos los elementos se han unido en un único clúster.
- **Métodos divisivos:** se empieza con un único cluster que contiene todos los elementos y, en cada paso, dividen recursivamente el cluster más heterogéneo hasta alcanzar un número deseado de particiones o un criterio de homogeneidad interna.

Ejemplo 4.1. Supongamos un conjunto de cuatro individuos con preferencias aprobatorias representadas en forma binaria, y consideremos la distancia de Hamming para medir sus discrepancias. Un algoritmo jerárquico aglomerativo comienza considerando cada individuo como un cluster separado. En el primer paso, se fusionan los dos individuos más similares (menor distancia). En los pasos siguientes, se repite la fusión según la menor distancia entre clústeres, hasta que finalmente todos los individuos pertenecen a un único grupo. El resultado puede representarse gráficamente mediante un dendrograma, donde la altura a la que se realiza cada unión refleja la distancia a la que se fusionaron los grupos.

Proposición 4.1.2. El clustering jerárquico produce una partición bien definida en cada nivel $k = 1, \dots, n$, y la estructura anidada puede representarse de forma única mediante un dendrograma. Además, cualquier corte horizontal en el dendrograma determina una partición de los datos en clústeres disjuntos.

En el contexto de preferencias aprobatorias, este tipo de métodos permite analizar tanto la similitud entre individuos (a partir de distancias como D_λ^x) como entre alternativas (mediante pseudodistancias como δ_λ). El clustering jerárquico es una buena opción para explorar la existencia de bloques de consenso y para identificar patrones de polarización en los perfiles de preferencias, [1]. La ventaja principal es que no es necesario fijar el número de grupos de antemano, lo cual es muy valioso en un análisis donde la estructura de las opiniones no siempre es clara.

Ejemplo 4.2. Veamos un ejemplo sencillo con la distancia de Hamming. Consideremos seis votantes, cuyas preferencias aprobatorias $(R_1, A_1), \dots, (R_6, A_6)$ se codifican binariamente sobre cinco alternativas $X = \{x_1, \dots, x_5\}$.

$$\begin{aligned} (R_1, A_1) : & \frac{x_1 \ x_2 \ x_5}{x_3 \ x_4} & (R_2, A_2) : & \frac{x_1 \ x_2 \ x_4 \ x_5}{x_3} & (R_3, A_3) : & \frac{x_1 \ x_5}{x_2 \ x_3 \ x_4} \\ (R_4, A_4) : & \frac{x_3}{x_1 \ x_2 \ x_4 \ x_5} & (R_5, A_5) : & \frac{x_3 \ x_5}{x_1 \ x_2 \ x_4} & (R_6, A_6) : & \frac{x_2 \ x_3}{x_1 \ x_4 \ x_5} \end{aligned}$$

Los vectores de aprobación correspondientes son

$$\begin{aligned} i_{A1} &= (1, 1, 0, 0, 1), & i_{A2} &= (1, 1, 0, 1, 1), & i_{A3} &= (1, 0, 0, 0, 1), \\ i_{A4} &= (0, 0, 1, 0, 0), & i_{A5} &= (0, 0, 1, 0, 1), & i_{A6} &= (0, 1, 1, 0, 0). \end{aligned}$$

Con la distancia de Hamming d_H sobre los vectores, la matriz de distancias es

$$D = (d_H(V_i, V_j))_{i,j=1}^6 = \begin{pmatrix} 0 & 1 & 1 & 4 & 3 & 3 \\ 1 & 0 & 2 & 5 & 4 & 4 \\ 1 & 2 & 0 & 3 & 2 & 4 \\ 4 & 5 & 3 & 0 & 1 & 1 \\ 3 & 4 & 2 & 1 & 0 & 2 \\ 3 & 4 & 4 & 1 & 2 & 0 \end{pmatrix}.$$

Apliquemos el clustering jerárquico aglomerativo. Partimos de 6 clústeres unitarios y, en cada paso, fusionamos los dos clústeres con menor distancia promedio entre sus elementos.

Paso 1. Las distancias mínimas son $d_H(V_1, V_2) = 1$ y $d_H(V_4, V_5) = 1$, por lo que esta será la primera fusión. El orden entre empates es irrelevante para el resultado final en términos de jerarquía.

$$\{V_1\} \cup \{V_2\} \xrightarrow{h=1} C_{12} = \{V_1, V_2\}.$$

Paso 2. Volvemos a tener mínimo $d_H(V_4, V_5) = 1$:

$$\{V_4\} \cup \{V_5\} \xrightarrow{h=1} C_{45} = \{V_4, V_5\}.$$

Paso 3. Observemos que la distancia promedio de V_3 al cluster C_{12} es

$$\bar{d}(V_3, C_{12}) = \frac{1}{2}(d_H(V_3, V_1) + d_H(V_3, V_2)) = \frac{1+2}{2} = 1,5,$$

que es la mínima en este momento. Fusionamos:

$$\{V_3\} \cup C_{12} \xrightarrow{h=1,5} C_{123} = \{V_1, V_2, V_3\}.$$

Paso 4. La distancia promedio de V_6 al cluster C_{45} es

$$\bar{d}(V_6, C_{45}) = \frac{1}{2}(d_H(V_6, V_4) + d_H(V_6, V_5)) = \frac{1+2}{2} = 1,5.$$

Fusionamos:

$$\{V_6\} \cup C_{45} \xrightarrow{h=1,5} C_{456} = \{V_4, V_5, V_6\}.$$

Paso 5 (final). Por último, la distancia promedio entre los clústeres C_{123} y C_{456} :

$$\bar{d}(C_{123}, C_{456}) = \frac{1}{3 \cdot 3} \sum_{i \in \{1,2,3\}} \sum_{j \in \{4,5,6\}} d_H(V_i, V_j) = \frac{32}{9} \approx 3,56.$$

Fusionamos:

$$C_{123} \cup C_{456} \xrightarrow{h \approx 3,56} \{V_1, \dots, V_6\}.$$

Dendrograma. El dendrograma resultante se representa en el gráfico 4.1.

Con ejemplo podemos ver cómo, bajo la distancia de Hamming, parecen diferenciarse dos bloques naturales $\{V_1, V_2, V_3\}$ y $\{V_4, V_5, V_6\}$ a alturas de fusión relativamente bajas (1 y 1.5). Cómo la unión entre ambos bloques ocurre a una altura mucho mayor ($\approx 3,56$), parece haber una estructura interna de similitud clara a dos niveles.

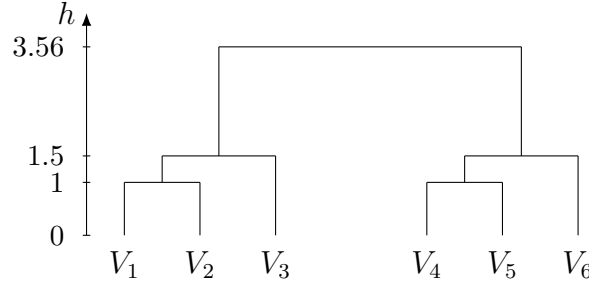


Gráfico 4.1: Dendrograma obtenido mediante clustering jerárquico aglomerativo.

Un aspecto importante en el uso de clustering jerárquico es la validación de la calidad del dendrograma obtenido. Para ello, se emplea con frecuencia el *coeficiente de correlación cophenética*, que evalúa hasta qué punto las distancias representadas en el dendrograma preservan las distancias originales entre los datos.

Definición 4.1.4. Sea $d(i, j)$ la distancia original entre dos elementos $x_i, x_j \in X$. La distancia cophenética $d_c(i, j)$ asociada a d se define como la altura del dendrograma en la que x_i y x_j se agrupan por primera vez en el mismo clúster.

Definición 4.1.5. Dadas las distancias originales $d(i, j)$ y las distancias cophenéticas $d_c(i, j)$, el coeficiente de correlación cophenética se define como

$$\text{Cophen}(X) = \frac{\sum_{i < j} (d(i, j) - \bar{d})(d_c(i, j) - \bar{d}_c)}{\sqrt{\sum_{i < j} (d(i, j) - \bar{d})^2 \sum_{i < j} (d_c(i, j) - \bar{d}_c)^2}},$$

donde $\bar{d}$ y $\bar{d}_c$ son las medias de las distancias originales y cophenéticas, respectivamente.

Un valor de este coeficiente cercano a 1 indica que el dendrograma es una buena representación de las distancias originales, mientras que valores bajos indican que la estructura jerárquica no captura bien la similitud entre los datos. Por esta razón, la correlación cophenética se emplea como criterio de validación en preferencias aprobatorias, ayudando a determinar qué medida de distancia genera particiones más fiables [1].

Ejemplo 4.3. Retomemos el dendrograma del ejemplo 4.2. Denotemos por $d(i, j) = d_H(V_i, V_j)$ las distancias de Hamming de la matriz D y por $d_c(i, j)$ la distancia cophenética asociada. Para $i \neq j$, las distancias cophenéticas correspondientes son

$$(d_c(i, j))_{i,j=1}^6 = \begin{pmatrix} 0 & 1 & 1,5 & 3,56 & 3,56 & 3,56 \\ 1 & 0 & 1,5 & 3,56 & 3,56 & 3,56 \\ 1,5 & 1,5 & 0 & 3,56 & 3,56 & 3,56 \\ 3,56 & 3,56 & 3,56 & 0 & 1 & 1,5 \\ 3,56 & 3,56 & 3,56 & 1 & 0 & 1,5 \\ 3,56 & 3,56 & 3,56 & 1,5 & 1,5 & 0 \end{pmatrix}.$$

Pasemos ahora a calcular el coeficiente de correlación cophenética presentado en la definición 4.1.5, donde las sumas recorren las $\binom{6}{2} = 15$ parejas, y $\bar{d}$, $\bar{d}_c$ son las medias muestrales de $\{d(i, j)\}$ y $\{d_c(i, j)\}$, respectivamente. En nuestro caso, la media de las distancias originales y la media de las distancias cophenéticas coinciden. Esto se debe a la simetría de los datos y al uso del criterio de enlace promedio. Sin embargo, en general no tiene por qué cumplirse que $\bar{d} = \bar{d}_c$. Tenemos, por tanto

$$\bar{d} = \bar{d}_c = \frac{1}{15} \sum_{i < j} d(i, j) = \frac{8}{3} \approx 2,6667,$$

y sustituyendo los valores se obtiene

$$\sum_{i < j} (d(i, j) - \bar{d})(d_c(i, j) - \bar{d}_c) = 18,1, \quad \sqrt{\sum_{i < j} (d(i, j) - \bar{d})^2 \sum_{i < j} (d_c(i, j) - \bar{d}_c)^2} = 21,420,$$

de donde

$$\text{Cophen}(X) \approx \frac{18,111}{21,420} = 0,846.$$

El valor $\text{Cophen} \approx 0,85$ indica que el dendrograma obtenido refleja bien las disimilitudes de Hamming, aunque se pierde algo de detalle en las distancias cruzadas. En este caso, confirma que existen dos bloques bien diferenciados.

En conclusión, el clustering jerárquico es una herramienta flexible y exploratoria que muy adecuada en el contexto de las preferencias aprobatorias. Su capacidad para representar múltiples niveles de agrupamiento lo convierte en un método especialmente útil para detectar tanto consensos amplios como divisiones internas significativas dentro de un colectivo.

4.1.1. Ranked K-Medoids (RKM)

El algoritmo *Ranked K-Medoids (RKM)* es una extensión del método clásico *K-medoids*, adaptada específicamente al contexto de las preferencias aprobatorias, [2]. El objetivo es agrupar *alternativas*, en lugar de votantes, de acuerdo con la forma en que estas son percibidas por el conjunto de individuos. De este modo, RKM permite identificar bloques de alternativas con perfiles de apoyo similares, revelando estructuras de sustitución o de polarización que dan más información que el análisis clásico centrado únicamente en los votantes.

Para aplicar RKM se necesita una función de disimilitud definida entre alternativas. Por ello, introducimos la pseudodistancia δ_λ , que combina dos dimensiones de información extraídas de las preferencias aprobatorias:

- (i) la *posición relativa* de las alternativas en los órdenes emitidos por los votantes, y
- (ii) la *aceptación o no* de dichas alternativas.

Notación 4.1. El símbolo $\oplus$ denota la operación lógica o exclusivo (*XOR*), que devuelve verdadero si exactamente una de las dos proposiciones se cumple, y falso en caso contrario.

Por otra parte, la notación $\mathbf{1}(\cdot)$ es la función indicatriz, que asigna el valor 1 cuando la condición dentro del paréntesis es verdadera, y 0 cuando es falsa.

Por tanto, se tiene que

$$\mathbf{1}(p \oplus q) = \begin{cases} 1 & \text{si solo uno es verdadero y el otro falso,} \\ 0 & \text{si } p = q. \end{cases}$$

Definición 4.1.6. Sea X el conjunto de alternativas y sean $(R_1, A_1), \dots, (R_n, A_n) \in R(X)$ las preferencias aprobatorias emitidas por n votantes respectivamente. Para dos alternativas $a, b \in X$ definimos:

- *La discordancia posicional:*

$$\rho_R(a, b) = \frac{1}{n} \sum_{i=1}^n \mathbf{1}(a \succ_i b \oplus b \succ_i a),$$

donde $\succ_i$ denota el orden inducido por $R(i)$. Esta cantidad mide la proporción de votantes que distinguen estrictamente entre a y b en su orden, es decir, aquellos que sitúan a a por encima de b o bien a b por encima de a , en contraposición a quienes los consideran indiferentes.

- *La discordancia de aprobación:*

$$\rho_A(a, b) = \frac{1}{n} \sum_{i=1}^n \mathbf{1}((a \in A(i)) \oplus (b \in A(i))) = \frac{1}{n} \sum_{i=1}^n \mathbf{1}(i_{A(i)}(a) \neq i_{A(i)}(b)),$$

que refleja la proporción de votantes que consideran aceptables a exactamente una de las dos alternativas, pero no a ambas.

A partir de estas dos medidas se define, para cada $\lambda \in [0, 1]$, la pseudodistancia

$$\delta_\lambda(a, b) = \lambda \rho_R(a, b) + (1 - \lambda) \rho_A(a, b).$$

Proposición 4.1.3. Para todo $\lambda \in [0, 1]$, la función δ_λ es una pseudodistancia sobre el conjunto de alternativas, es decir:

1. $\delta_\lambda(a, b) \geq 0$, con igualdad si a y b son indistinguibles en el perfil $\mathcal{P}$,
2. $\delta_\lambda(a, b) = \delta_\lambda(b, a)$ (simetría),
3. $\delta_\lambda(a, c) \leq \delta_\lambda(a, b) + \delta_\lambda(b, c)$ (desigualdad triangular).

Demostración. (i) *No negatividad e identidad de los indistinguibles.* Por construcción, cada sumando en ρ_R y ρ_A es un indicador en $\{0, 1\}$, luego $0 \leq \rho_R, \rho_A \leq 1$ y, por tanto, $0 \leq \delta_\lambda(a, b) \leq 1$. Además, si a y b son indistinguibles en el perfil (ya que para todo i ni $a \succ_i b$ ni $b \succ_i a$ y $a \in A(i) \Leftrightarrow b \in A(i)$), entonces todos los indicadores valen 0, de modo que $\rho_R(a, b) = \rho_A(a, b) = 0$ y, en consecuencia, $\delta_\lambda(a, b) = 0$.

(ii) *Simetría.* Tanto ρ_R como ρ_A son simétricas, ya que al intercambiar a y b no cambia el valor de los indicadores (en ρ_R la condición “ $a \succ_i b$ xor $b \succ_i a$ ” es invariante por intercambio y en ρ_A la condición “ $a \in A(i)$ xor $b \in A(i)$ ” también lo es). Por tanto, $\delta_\lambda(a, b) = \delta_\lambda(b, a)$.

(iii) *Desigualdad triangular.* Demostraremos primero que ρ_R y ρ_A satisfacen la desigualdad triangular y después que cualquier combinación convexa hereda la propiedad.

Triangularidad de ρ_A . Para cada votante i , defínase $u_i(a) := \mathbf{1}\{a \in A(i)\} \in \{0, 1\}$. Entonces

$$\mathbf{1}((a \in A(i)) \oplus (b \in A(i))) = |u_i(a) - u_i(b)|.$$

Dado $a, b, c \in X$, se tiene para $u, v, w \in \{0, 1\}$ la desigualdad $|u - w| \leq |u - v| + |v - w|$. Aplicada coordenada a coordenada y sumando en i ,

$$\sum_{i=1}^n |u_i(a) - u_i(c)| \leq \sum_{i=1}^n |u_i(a) - u_i(b)| + \sum_{i=1}^n |u_i(b) - u_i(c)|,$$

y al dividir por n se obtiene $\rho_A(a, c) \leq \rho_A(a, b) + \rho_A(b, c)$.

Triangularidad de ρ_R . Para cada votante i , definimos

$$\sigma_i(a, b) := \mathbf{1}(a \succ_i b \oplus b \succ_i a) \in \{0, 1\}.$$

Observemos que, al ser $R(i)$ un orden débil (completo y transitivo), la relación de indiferencia $\sim_i$ es transitiva. Por tanto: si $\sigma_i(a, b) = \sigma_i(b, c) = 0$ (esto es, $a \sim_i b$ y $b \sim_i c$), entonces $a \sim_i c$ y $\sigma_i(a, c) = 0$; en caso contrario, al menos uno de $\sigma_i(a, b)$ o $\sigma_i(b, c)$ es 1, y entonces siempre se verifica $\sigma_i(a, c) \leq 1 \leq \sigma_i(a, b) + \sigma_i(b, c)$. En ambos casos,

$$\sigma_i(a, c) \leq \sigma_i(a, b) + \sigma_i(b, c).$$

Sumando en i y normalizando,

$$\rho_R(a, c) = \frac{1}{n} \sum_{i=1}^n \sigma_i(a, c) \leq \frac{1}{n} \sum_{i=1}^n \sigma_i(a, b) + \frac{1}{n} \sum_{i=1}^n \sigma_i(b, c) = \rho_R(a, b) + \rho_R(b, c).$$

Conclusión para δ_λ . Si d_1 y d_2 son pseudodistancias, entonces para $\alpha, \beta \geq 0$ con $\alpha + \beta = 1$ la combinación $\alpha d_1 + \beta d_2$ también lo es, pues

$$\begin{aligned} \alpha d_1(a, c) + \beta d_2(a, c) &\leq \alpha(d_1(a, b) + d_1(b, c)) + \beta(d_2(a, b) + d_2(b, c)) \\ &= (\alpha d_1 + \beta d_2)(a, b) + (\alpha d_1 + \beta d_2)(b, c). \end{aligned}$$

Aplicando esto con $d_1 = \rho_R$, $d_2 = \rho_A$, $\alpha = \lambda$ y $\beta = 1 - \lambda$, se obtiene la desigualdad triangular para δ_λ , luego queda probado que δ_λ cumple las propiedades enunciadas para todo $\lambda \in [0, 1]$. $\square$

El algoritmo RKM sigue la filosofía de K-medoids, pero operando directamente sobre la matriz de distancias δ_λ entre alternativas. A diferencia de K-means, no requiere centroides abstractos, sino que selecciona un conjunto de *medoides*, es decir, alternativas reales que representan a cada clúster. La particularidad de RKM es que emplea una *matriz de rangos* construida a partir de las distancias δ_λ , lo que garantiza que las asignaciones de alternativas a clústeres respetan la estructura de similitud percibida por los votantes.

Definición 4.1.7. Sea $X = \{x_1, \dots, x_m\}$ el conjunto de alternativas y sea δ_λ la pseudodistancia presentada en la definición 4.1.6. El objetivo del algoritmo es encontrar una partición en k clústeres

$$C = \{C_1, \dots, C_k\}, \quad \bigsqcup_{j=1}^k C_j = X,$$

y un conjunto de medoides $M = \{m_1, \dots, m_k\} \subset X$ tales que se minimice el coste total

$$\Phi(C, M) = \sum_{j=1}^k \sum_{x \in C_j} \delta_\lambda(x, m_j).$$

El procedimiento iterativo del algoritmo es el siguiente:

1. **Inicialización.** Se eligen aleatoriamente k alternativas distintas de X , que constituyen el conjunto inicial de medoides $M^{(0)} = \{m_1^{(0)}, \dots, m_k^{(0)}\}$.

2. **Asignación.** Dado un conjunto de medoides $M^{(t)}$, se define la partición $C^{(t)} = \{C_1^{(t)}, \dots, C_k^{(t)}\}$ asignando cada alternativa $x \in X$ al cluster del medoide más cercano:

$$C_j^{(t)} = \left\{ x \in X : \delta_\lambda(x, m_j^{(t)}) \leq \delta_\lambda(x, m_\ell^{(t)}) \quad \forall \ell \in \{1, \dots, k\} \right\},$$

con una regla fija de desempate en caso de igualdad.

3. **Actualización de medoides.** Para cada cluster $C_j^{(t)}$, se calcula para cada candidato $a \in C_j^{(t)}$ la suma de distancias internas

$$S_j^{(t)}(a) = \sum_{y \in C_j^{(t)}} \delta_\lambda(a, y).$$

El nuevo medoide $m_j^{(t+1)}$ se elige como el elemento que minimiza esta suma:

$$m_j^{(t+1)} \in \arg \min_{a \in C_j^{(t)}} S_j^{(t)}(a).$$

En caso de empate, se escoge según una regla de desempate prefijada (p. ej. el primero según un orden lexicográfico en X).

4. **Iteración.** Se repiten los pasos 2 y 3 hasta que los medoides no cambian entre dos iteraciones consecutivas, es decir, $M^{(t+1)} = M^{(t)}$, o hasta alcanzar un número máximo de iteraciones.

Ejemplo 4.4. Supongamos un conjunto de cuatro alternativas $X = \{a, b, c, d\}$ evaluadas por un grupo de votantes. Si las alternativas a y b son aceptables por un subconjunto muy similar de votantes y suelen ocupar posiciones cercanas en sus órdenes, mientras que c y d suelen ser inaceptables, la pseudodistancia δ_λ devolverá valores pequeños para $\delta_\lambda(a, b)$ y para $\delta_\lambda(c, d)$, y valores altos para $\delta_\lambda(a, c)$, $\delta_\lambda(b, d)$, etc. Aplicando RKM con $k = 2$, es muy probable que se formen los clústeres $\{a, b\}$ y $\{c, d\}$, revelando la existencia de dos bloques bien diferenciados de alternativas.

El método RKM resulta particularmente útil para:

- detectar *bloques de alternativas* con perfiles de apoyo semejantes,
- explorar *estructuras de sustitución*, es decir, grupos de alternativas que pueden ser vistas como equivalentes por los votantes,
- analizar la *polarización* en torno a conjuntos de alternativas que generan patrones de aprobación muy distintos.

Ejemplo 4.5. Veamos un ejemplo con para enseñar el funcionamiento de Ranked K-Medoids (RKM) sobre alternativas. Consideramos seis votantes $\{V_1, \dots, V_6\}$ y cinco alternativas $X = \{x_1, \dots, x_5\}$:

$$V_1 : \begin{array}{cc} x_1 & x_2 \\ x_5 & \\ x_3 & x_4 \end{array} \quad V_2 : \begin{array}{cc} x_1 & \\ x_4 & x_5 \\ x_2 & \\ x_3 & \end{array} \quad V_3 : \begin{array}{cc} x_5 & \\ x_1 & \\ x_2 & x_3 \\ x_4 & \end{array}$$

$$V_4 : \begin{array}{c} x_3 \\ x_1 \ x_2 \\ x_4 \ x_5 \end{array} \quad V_5 : \begin{array}{c} x_3 \ x_5 \\ x_1 \ x_2 \\ x_4 \end{array} \quad V_6 : \begin{array}{c} x_2 \ x_3 \\ x_4 \ x_5 \\ x_1 \end{array}$$

Aplicando las definiciones de ρ_A y ρ_R se obtienen las matrices:

$$\rho_A = \begin{pmatrix} 0 & \frac{1}{2} & 1 & \frac{1}{3} & \frac{1}{6} \\ \frac{1}{2} & 0 & \frac{1}{2} & \frac{1}{2} & \frac{2}{3} \\ 1 & \frac{1}{2} & 0 & \frac{2}{3} & \frac{5}{6} \\ \frac{1}{3} & \frac{1}{2} & \frac{2}{3} & 0 & \frac{1}{2} \\ \frac{1}{6} & \frac{2}{3} & \frac{5}{6} & \frac{1}{2} & 0 \end{pmatrix}, \quad \rho_R = \begin{pmatrix} 0 & \frac{1}{2} & 1 & 1 & 1 \\ \frac{1}{2} & 0 & \frac{2}{3} & 1 & 1 \\ 1 & \frac{2}{3} & 0 & \frac{5}{6} & \frac{5}{6} \\ 1 & 1 & \frac{5}{6} & 0 & \frac{1}{2} \\ 1 & 1 & \frac{5}{6} & \frac{1}{2} & 0 \end{pmatrix}.$$

Tomando $\lambda = \frac{1}{2}$, la matriz de disimilitudes es

$$\Delta = \delta_{1/2} = \frac{1}{2}(\rho_A + \rho_R) = \begin{pmatrix} 0 & \frac{1}{2} & 1 & \frac{2}{3} & \frac{7}{12} \\ \frac{1}{2} & 0 & \frac{7}{12} & \frac{3}{4} & \frac{5}{6} \\ 1 & \frac{7}{12} & 0 & \frac{3}{4} & \frac{5}{6} \\ \frac{2}{3} & \frac{3}{4} & \frac{3}{4} & 0 & \frac{1}{2} \\ \frac{7}{12} & \frac{5}{6} & \frac{5}{6} & \frac{1}{2} & 0 \end{pmatrix}.$$

Veamos ahora la ejecución de RKM con $k = 2$.

(1) *Inicialización.* Tomamos como medoides iniciales

$$M^{(0)} = \{x_1, x_3\}.$$

(2) *Asignación.* Cada alternativa $x \in X$ se asigna al medoide más cercano según Δ :

	x_1	x_2	x_3	x_4	x_5
$\delta_{1/2}(x, x_1)$	0	$\frac{1}{2}$	1	$\frac{2}{3}$	$\frac{7}{12}$
$\delta_{1/2}(x, x_3)$	1	$\frac{7}{12}$	0	$\frac{3}{4}$	$\frac{5}{6}$

De aquí resulta la partición

$$C_1^{(0)} = \{x_1, x_2, x_4, x_5\}, \quad C_2^{(0)} = \{x_3\}.$$

(3) *Actualización de medoides.* Para cada candidato $a \in C_1^{(0)}$, calculamos la suma de distancias internas

$$S_1^{(0)}(a) = \sum_{y \in C_1^{(0)}} \delta_{1/2}(a, y).$$

Explícitamente:

$$\begin{aligned} S_1^{(0)}(x_1) &= \frac{1}{2} + \frac{2}{3} + \frac{7}{12} = \frac{55}{36} \approx 1,528, \\ S_1^{(0)}(x_2) &= \frac{1}{2} + \frac{3}{4} + \frac{5}{6} = \frac{25}{12} \approx 2,083, \\ S_1^{(0)}(x_4) &= \frac{2}{3} + \frac{3}{4} + \frac{1}{2} = \frac{25}{12} \approx 2,083, \\ S_1^{(0)}(x_5) &= \frac{7}{12} + \frac{5}{6} + \frac{1}{2} = \frac{31}{12} \approx 2,583. \end{aligned}$$

El mínimo se alcanza en x_1 , que permanece como medoide de $C_1^{(0)}$. El cluster $C_2^{(0)}$ es unitario, por lo que su medoide es x_3 . Así,

$$M^{(1)} = \{x_1, x_3\}.$$

(4) *Iteración.* Como $M^{(1)} = M^{(0)}$, la partición se mantiene y el algoritmo converge.

El resultado final es el siguiente:

$$\boxed{C_1 = \{x_1, x_2, x_4, x_5\} \text{ (medoide } x_1), \quad C_2 = \{x_3\} \text{ (medoide } x_3)}.$$

Las indiferencias hacen más pequeño ρ_R en aquellos pares que algunos votantes colocan en la misma clase (como $x_4 \sim x_5$ para V_2 y V_6), mientras que la aprobación sigue marcando separaciones fuertes (por ejemplo, x_3 frente a $\{x_1, x_2, x_5\}$). El resultado final nos muestra un bloque mayoritario $\{x_1, x_2, x_4, x_5\}$ y una alternativa singular x_3 , lo que es coherente con una estructura de sustitución dentro del bloque y de polarización respecto a x_3 .

Veamos ahora la ejecución de RKM con $k = 3$.

(1) *Inicialización.* Tomamos como medoides iniciales

$$M^{(0)} = \{x_1, x_3, x_5\}.$$

(2) *Asignación.* Cada alternativa se asigna al medoide más cercano. A partir de la matriz Δ se obtiene:

	x_1	x_2	x_3	x_4	x_5
$\delta_{1/2}(x, x_1)$	0	$\frac{1}{2}$	1	$\frac{2}{3}$	$\frac{7}{12}$
$\delta_{1/2}(x, x_3)$	1	$\frac{7}{12}$	0	$\frac{3}{4}$	$\frac{5}{6}$
$\delta_{1/2}(x, x_5)$	$\frac{7}{12}$	$\frac{5}{6}$	$\frac{5}{6}$	$\frac{1}{2}$	0

La asignación inicial es entonces:

$$C_1^{(0)} = \{x_1, x_2\}, \quad C_2^{(0)} = \{x_3\}, \quad C_3^{(0)} = \{x_4, x_5\}.$$

(3) *Actualización de medoides.*

- En $C_1^{(0)} = \{x_1, x_2\}$,

$$S_1^{(0)}(x_1) = \delta(x_1, x_2) = \frac{1}{2}, \quad S_1^{(0)}(x_2) = \delta(x_2, x_1) = \frac{1}{2}.$$

Ambos empatan, así que se mantiene $m_1^{(1)} = x_1$ por la regla de desempate.

- En $C_2^{(0)} = \{x_3\}$, el medoide es $m_2^{(1)} = x_3$.
- En $C_3^{(0)} = \{x_4, x_5\}$,

$$S_3^{(0)}(x_4) = \delta(x_4, x_5) = \frac{1}{2}, \quad S_3^{(0)}(x_5) = \delta(x_5, x_4) = \frac{1}{2}.$$

También hay empate, luego se mantiene $m_3^{(1)} = x_4$ por desempate.

Por tanto,

$$M^{(1)} = \{x_1, x_3, x_4\}.$$

(4) *Iteración.* Con los nuevos medoides $M^{(1)}$ se repite la asignación, pero la partición no cambia:

$$C_1^{(1)} = \{x_1, x_2\}, \quad C_2^{(1)} = \{x_3\}, \quad C_3^{(1)} = \{x_4, x_5\}.$$

Así, el algoritmo converge.

El resultado final es el siguiente:

$$\boxed{C_1 = \{x_1, x_2\} \text{ (medoide } x_1), \quad C_2 = \{x_3\} \text{ (medoide } x_3), \quad C_3 = \{x_4, x_5\} \text{ (medoide } x_4)}.$$

Con $k = 3$, el algoritmo identifica una estructura más fina: un subbloque $\{x_1, x_2\}$ con gran similitud, un bloque $\{x_4, x_5\}$ unido por sus indiferencias frecuentes, y la alternativa x_3 que sigue apareciendo como un outlier. Esto nos muestra cómo la elección del número de clústeres permite detectar distintos niveles de granularidad en la organización de las alternativas.

En conclusión, RKM proporciona una herramienta metodológica sólida para el estudio de alternativas en el marco de las preferencias aprobatorias. Su fundamento en pseudodistancias específicas como δ_λ permite capturar bien la percepción colectiva de las alternativas, y su carácter particional lo hace compatible con análisis empíricos donde se buscan estructuras internas en los datos. En el Capítulo 5 se verá cómo la combinación de δ_λ con RKM puede permitir descubrir patrones en el espacio de alternativas que no serían evidentes mediante el análisis exclusivo de las distancias entre votantes.

Otras técnicas

Además de los métodos descritos anteriormente, existen otros tipos de análisis cluster que, si bien no emplearemos en este trabajo, merecen ser mencionados por su relevancia general y su posible aplicabilidad al contexto de las preferencias aprobatorias:

- **DBSCAN** (*Density-Based Spatial Clustering of Applications with Noise*): se trata de un método basado en densidad que permite identificar clústeres de forma arbitraria sin necesidad de especificar de antemano el número de grupos. Resulta útil cuando los datos presentan regiones de alta densidad separadas por zonas de baja densidad. También es buena opción para detectar posibles observaciones atípicas.
- **Mean-Shift**: es un algoritmo no paramétrico que busca modos o máximos locales en la distribución de los datos. No requiere fijar el número de clústeres y se adapta bien a estructuras complejas. Sin embargo, es limitado aplicado a datos categóricos como las preferencias aprobatorias, salvo que se realicen las transformaciones adecuadas del espacio de representación.
- **Clustering espectral**: se basa en la teoría de grafos y utiliza una matriz de afinidad o similitud entre los elementos para realizar una partición mediante descomposición espectral del Laplaciano asociado. Es especialmente potente cuando los clústeres no son convexos y puede adaptarse al contexto de preferencias aprobatorias siempre que se construya una matriz de afinidad a partir de distancias como D_λ^* .

En resumen, estos métodos son perspectivas alternativas que pueden resultar útiles en ciertos contextos específicos. En este trabajo no son tratados ya que no se adecúan tanto a la naturaleza de los datos que analizaremos.

4.2. Elección de la distancia y del método según el objetivo

En esta sección formalizaremos tres fenómenos fundamentales que pueden observarse al analizar preferencias aprobatorias. Estos son la *polarización*, el *consenso parcial* y los *bloques homogéneos de alternativas*.

Definición 4.2.1. Sean n votantes con preferencias aprobatorias $(R_1, A_1), \dots, (R_n, A_n) \in R(X)$, y sea $d : R(X) \times R(X) \rightarrow [0, \infty)$ una función de distancia definida sobre sus preferencias. Diremos que existe polarización si la partición $C = \{C_1, \dots, C_k\}$ obtenida mediante un algoritmo de clustering satisface

$$\min_{1 \leq j \leq k} \frac{1}{|C_j|^2} \sum_{x, y \in C_j} d(x, y) \ll \min_{p \neq q} \frac{1}{|C_p||C_q|} \sum_{x \in C_p, y \in C_q} d(x, y),$$

es decir, los votantes son internamente homogéneos dentro de cada bloque, mientras que la distancia promedio entre bloques es grande.

Ejemplo 4.6 (Polarización política). Un ejemplo clásico de polarización se suele dar en elecciones presidenciales con sistemas bipartidistas. En tales contextos, la mayoría de los votantes se concentra en dos clústeres claramente separados, los que apoyan al candidato de un partido y los que apoyan al del partido contrario. Las distancias dentro de cada grupo son pequeñas (alta cohesión interna), pero las distancias entre grupos son máximas (fuerte oposición). Este fenómeno se suele dar de manera recurrente en democracias consolidadas con dos grandes fuerzas políticas enfrentadas.

Definición 4.2.2. Con las mismas notaciones anteriores, decimos que existe consenso parcial si la partición $C = \{C_1, \dots, C_k\}$ contiene al menos un cluster C_j de tamaño significativo (por ejemplo $|C_j| \gg |C_\ell|$ para $\ell \neq j$) tal que

$$\frac{1}{|C_j|^2} \sum_{x, y \in C_j} d(x, y) \text{ es bajo,}$$

mientras que los clústeres restantes son más pequeños y tienen una mayor heterogeneidad. El consenso no es total (pues existen discrepancias), pero sí se observa un núcleo mayoritario de proximidad en las preferencias.

Ejemplo 4.7. Un caso típico de consenso parcial puede encontrarse en encuestas de opinión sobre medidas sociales ampliamente aceptadas, como la obligatoriedad del cinturón de seguridad o la prohibición del tabaco en espacios cerrados. La mayoría de la población comparte posiciones similares (gran cluster de consenso), aunque persisten minorías que disienten (clústeres pequeños con mayor dispersión).

Definición 4.2.3. Sea $X = \{x_1, \dots, x_m\}$ el conjunto de alternativas y sea $\delta : X \times X \rightarrow \mathbb{R}_{\geq 0}$ una pseudodistancia definida a partir de las preferencias aprobatorias (por ejemplo, δ_λ). Diremos que existen bloques homogéneos de alternativas si la partición $C = \{C_1, \dots, C_r\}$ de A cumple que, para cada cluster C_j ,

$$\max_{x, y \in C_j} \delta(x, y) \ll \min_{x \in C_j, y \in C_\ell, j \neq \ell} \delta(x, y),$$

es decir, las alternativas dentro de cada bloque son percibidas como muy similares por los votantes, mientras que los bloques están bien separados entre sí.

Ejemplo 4.8. En ciertos sistemas electorales multipartidistas, los partidos con programas ideológicos parecidos son vistos por los votantes como opciones prácticamente intercambiables. Por ejemplo, en algunos procesos electorales europeos se han formado bloques de partidos de izquierda moderada que comparten gran parte de sus apoyos, mientras que partidos más extremos o de signo contrario aparecen en bloques bien diferenciados. En este caso, las alternativas dentro de cada bloque son homogéneas, pero las distancias entre bloques son grandes.

Estos fenómenos se estudian de manera más efectiva mediante combinaciones diferentes de distancias y algoritmos de agrupamiento:

- Para medir la *polarización*, se recomiendan métodos particionales como K-medoids aplicados sobre distancias entre votantes (por ejemplo D_λ^r o combinaciones de Kemeny-Hamming).
- Para medir el *consenso parcial*, los métodos jerárquicos permiten explorar la existencia de un núcleo amplio de acuerdo y su relación con minorías.
- Para medir los *bloques homogéneos de alternativas*, la combinación de la pseudodistancia δ_λ con el algoritmo RKM proporciona una aproximación específica al análisis de similitudes entre alternativas.

En el capítulo 5 aplicaremos estos métodos a un conjunto de datos. Según sus características decidiremos cuales son los métodos y distancias mas convenientes para estudiarlo.

Capítulo 5

Análisis empírico

Para ilustrar los métodos desarrollados en los capítulos anteriores, se analizará un conjunto de datos de carácter político y social. En concreto, se utiliza el *Democracy Fund Voter Study Group* (VSG) de 2019, [13], una encuesta representativa de la población estadounidense que recoge actitudes, percepciones y preferencias hacia partidos, líderes políticos, grupos sociales y políticas públicas.

Se ha elegido esta ola de 2019 por varias razones. En primer lugar, ofrece indicadores continuos y ordinales (como los “termómetros de favorabilidad” de 0-100) que pueden transformarse en preferencias aprobatorias fijando umbrales y ordenando las alternativas, lo que permite construir preferencias aprobatorias. En segundo lugar, se trata de un momento preelectoral, justo antes de las presidenciales de 2020, en el que cabe esperar tanto situaciones de consenso parcial como fenómenos de fuerte polarización. Por último, es un conjunto de alta calidad técnica y muy bien documentado, disponible en formatos estándar (CSV/Stata/SPSS), lo que facilita su preparación y análisis.

Para este trabajo no se han utilizado todas las variables del dataset, sino algunas de las que permiten construir preferencias aprobatorias y aplicar métodos de clustering.

En este contexto, el objetivo es estudiar los patrones de distancia hacia distintos grupos sociales de la población estadounidense. Para ello se utilizarán los *termómetros de favorabilidad* incluidos en la encuesta, que miden en una escala de 0 a 100 el grado de simpatía o antipatía hacia distintos grupos sociales, como se puede ver en la tabla 5.1. Según la descripción oficial del conjunto de datos [13], una calificación entre 50 y 100 significa que el votante siente afecto y simpatía por el grupo. Una calificación entre 0 y 49 significa que no siente afecto por el grupo y que no le importa demasiado.

Nombre de la variable	Descripción
termometro_afro	Opinión hacia afroamericanos (0–100).
termometro_blancos	Opinión hacia blancos (0–100).
termometro_asia	Opinión hacia asiáticos (0–100).
termometro_latino	Opinión hacia latinos (0–100).
termometro_judio	Opinión hacia judíos (0–100).
termometro_muslim	Opinión hacia musulmanes (0–100).
termometro_gay	Opinión hacia personas homosexuales (0–100).
termometro_inmig	Opinión hacia inmigrantes (0–100).

Tabla 5.1: Variables de termómetros de favorabilidad seleccionadas.

Estas variables se transformarán en *preferencias aprobatorias* mediante umbrales que permiten distinguir las categorías “aceptable”/“inaceptable” y, a la vez, ordenarlas de mayor a menor calidez. En este caso, el umbral será de 50 puntos. De esta manera es posible aplicar los métodos vistos en los capítulos anteriores sobre un conjunto real de opiniones sociales.

La estrategia consiste en identificar clusters o perfiles de votantes en función de sus simpatías hacia estos grupos, para después compararlos con otras variables sustantivas como las que aparecen en la tabla 5.2.

Nombre de la variable	Descripción
anio_nac	Año de nacimiento.
genero	Género.
raza	Raza/etnicidad.
educacion	Nivel educativo.
ingreso_fam	Ingresos familiares.
religion	Afiliación religiosa.
importancia_relig	Importancia de la religión.
dueno_armas	Propiedad de armas en el hogar.
interes_noticias	Interés por noticias/política.
partido_id_3	Identificación partidista (3 categorías).
partido_id_7	Identificación partidista (7 categorías).
ideologia_5pt	Autoubicación ideológica (5+NS).
voto_2016	Recuerdo de voto en 2016.
voto_2020	Voto/intención presidencial 2020 (pre-elección).

Tabla 5.2: Variables sociodemográficas y político-electorales utilizadas como comparativas en el análisis.

En resumen, el análisis servirá para ver de manera aplicada cómo las preferencias aprobatorias pueden utilizarse para descubrir patrones de la opinión pública en los Estados Unidos previa a las elecciones presidenciales de 2020.

Avancemos a la preparación de los datos. El primer paso consiste en transformar las variables originales de la encuesta en estructuras de preferencias aprobatorias. Para ello, se realiza lo siguiente:

- **Selección y renombrado de variables.** Por una parte, se utilizarán las variables del cuestionario original que permiten construir preferencias aprobatorias o caracterizar los clústeres. Estas son los *termómetros de favorabilidad* (0-100) de la tabla 5.1. En el segundo grupo de variables utilizadas, correspondiente a las de la tabla 5.2, se encuentran variables sociodemográficas, ideológicas y de comportamiento electoral, que no intervienen en la construcción de las distancias pero sirven para interpretar los clústeres.
- **Codificación de preferencias aprobatorias.** Cada individuo queda representado por un par (R, A) :
 - El subconjunto A recoge las alternativas *aceptables*, que en nuestro caso serán los que superen el umbral de 50 puntos. Es decir, si la valoración de un individuo hacia una alternativa es mayor o igual que 50, ésta se considera aceptable. En caso contrario, inaceptable.

- El orden R se obtiene a partir de la puntuación de apoyo que han dado los votantes. En los termómetros, las alternativas se ordenan de mayor a menor puntuación.

Este proceso convierte las respuestas de la encuesta en objetos matemáticos, a los que ya se pueden aplicar las técnicas de análisis cluster presentadas en el capítulo 4. Los parámetros λ y r controlan, respectivamente, la importancia relativa de la parte aprobatoria frente a la ordinal y la sensibilidad de la distancia, y serán explorados en el análisis de sensibilidad posterior.

5.1. Elección del método

Para el análisis cluster de las preferencias aprobatorias es necesario primero seleccionar las distancias y el método de agrupamiento que se usará. Esto se selecciona en función de la naturaleza de los datos y del desarrollo teórico expuesto en los capítulos 3 y 4. El objetivo es integrar de forma equilibrada la información sobre la aceptabilidad de las alternativas y el orden relativo que cada votante establece.

Para combinar tanto la dimensión ordinal como la binaria, se usará la *distancia* D_λ^r . Como se vio en la sección 3.5, esta familia de distancias permite ponderar la importancia relativa de las discrepancias ordinales y de aprobación mediante dos parámetros:

- $\lambda \in (0, 1)$, que regula el peso relativo de las dos componentes.
- $r \geq 1$, que controla la sensibilidad frente a desacuerdos intensos, penalizando más fuertemente las discrepancias grandes cuando se quiere resaltar los casos de polarización fuerte.

Para este estudio se han elegido los parámetros

- $\lambda = 0,5$,
- $r = 1$.

Se ha tomado esta elección porque ofrece una buena capacidad de ajuste a un conjunto de datos heterogéneo como el de [13], en el que hay variables continuas transformadas y binarias y donde pueden aparecer tanto áreas de consenso parcial como focos polarización marcados.

Finalmente, para la etapa de agrupamiento se adopta el método *Ranked K-Medoids (RKM)*. Como se vio en la subsección 4.1.1, este algoritmo particional está diseñado para trabajar directamente con matrices de distancias definidas sobre preferencias aprobatorias. Tiene una gran capacidad para identificar medoides representativos y manejar de forma eficiente grandes poblaciones, por lo que es una herramienta idónea para el análisis de las opiniones sociopolíticas con el que se realizará el estudio.

5.2. Clustering de votantes

En esta sección se aplica el procedimiento de clustering sobre los votantes empleando exclusivamente la distancia combinada D_λ^r definida en la sección 3.5, con los parámetros $\lambda = 0,5$ y $r = 1$.

Partiendo de las variables de cluster de la tabla 5.1, se han realizado las transformaciones indicadas anteriormente:

- 1) se han eliminado a los votantes con códigos de no respuesta o que se han saltado la pregunta.

- 2) las respuestas de 0 a 100 se han binarizado en alternativas aceptables (≥ 50) e inaceptables (< 50).
- 3) se ha construido, para cada votante, la posición media $PR(x_j)$ de cada alternativa ordenando de mayor a menor puntuación.

A partir de las matrices de aprobaciones IA y posiciones PR se han calculado todas las discordancias posicionales y de aprobación para cada par de alternativas. Con ellas se ha formado la matriz de distancias D_λ^r entre todos los votantes. Esta matriz, simétrica y con diagonal nula, es la base del análisis de clúster.

Para determinar el número óptimo de clústeres k se ha aplicado el algoritmo *Partitioning Around Medoids* (PAM), una versión robusta de *K-medoids*, sobre la matriz de distancias D_λ^r . El criterio de validación interna que se ha empleado ha sido la *anchura media de silhouette*, que evalúa la cohesión interna a la vez que la separación entre grupos.

En el gráfico 5.1 se puede ver la silhouette media en función de $k \in [2, 10]$. El máximo se alcanza en $k = 2$ con un valor de aproximadamente 0,271. Este resultado podría sugerir que la estructura de los datos se describe mejor mediante dos grandes grupos de votantes.

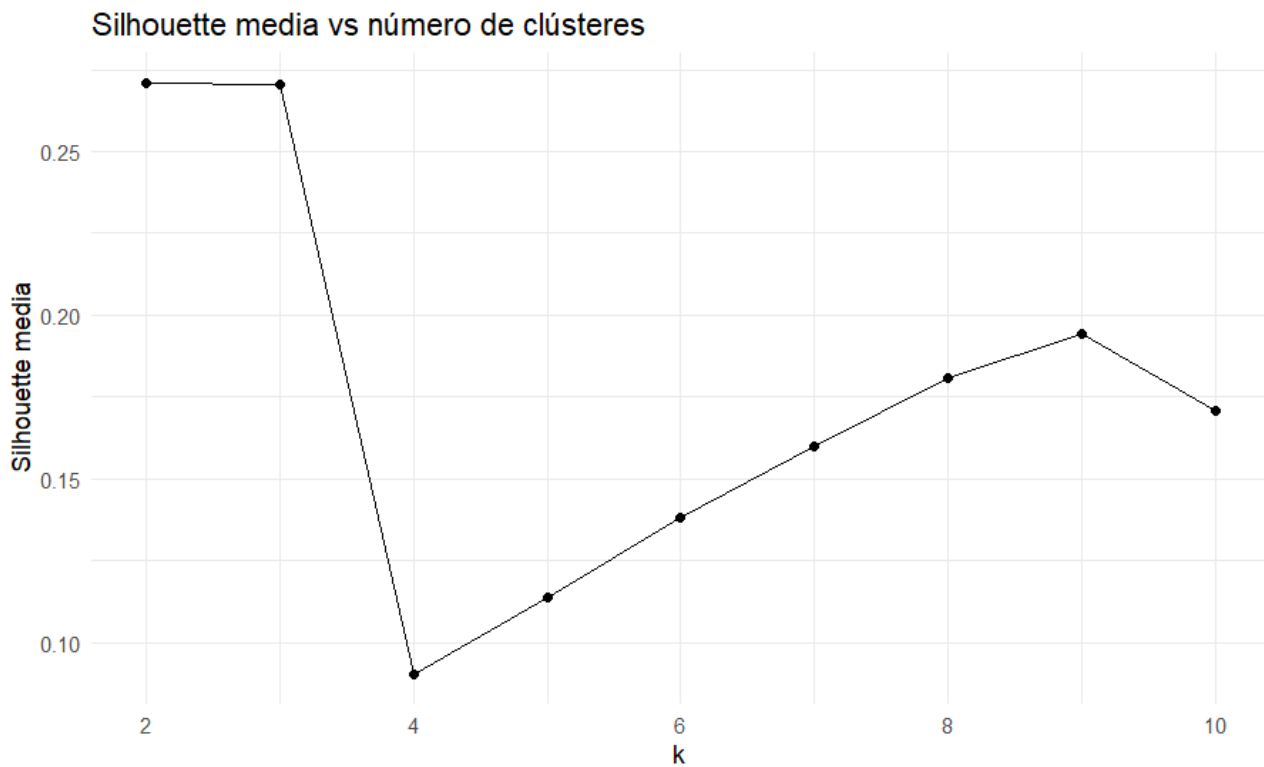


Gráfico 5.1: Silhouette media en función del número de clústeres k .

El valor de 0,271, aunque indica la presencia de dos bloques, es relativamente bajo en la escala habitual del índice (donde valores superiores a 0.5 se interpretan como separación fuerte). Por tanto, existe una polarización pero *moderada*. Es decir, dos bloques principales, pero con solapamiento y una fracción apreciable de individuos en posiciones intermedias.

Para facilitar la interpretación se ha realizado un *Multidimensional Scaling* (MDS) clásico que proyecta las distancias D_λ^r en un plano bidimensional (gráfico 5.2). Cada punto representa a un votante, coloreado según el cluster asignado.

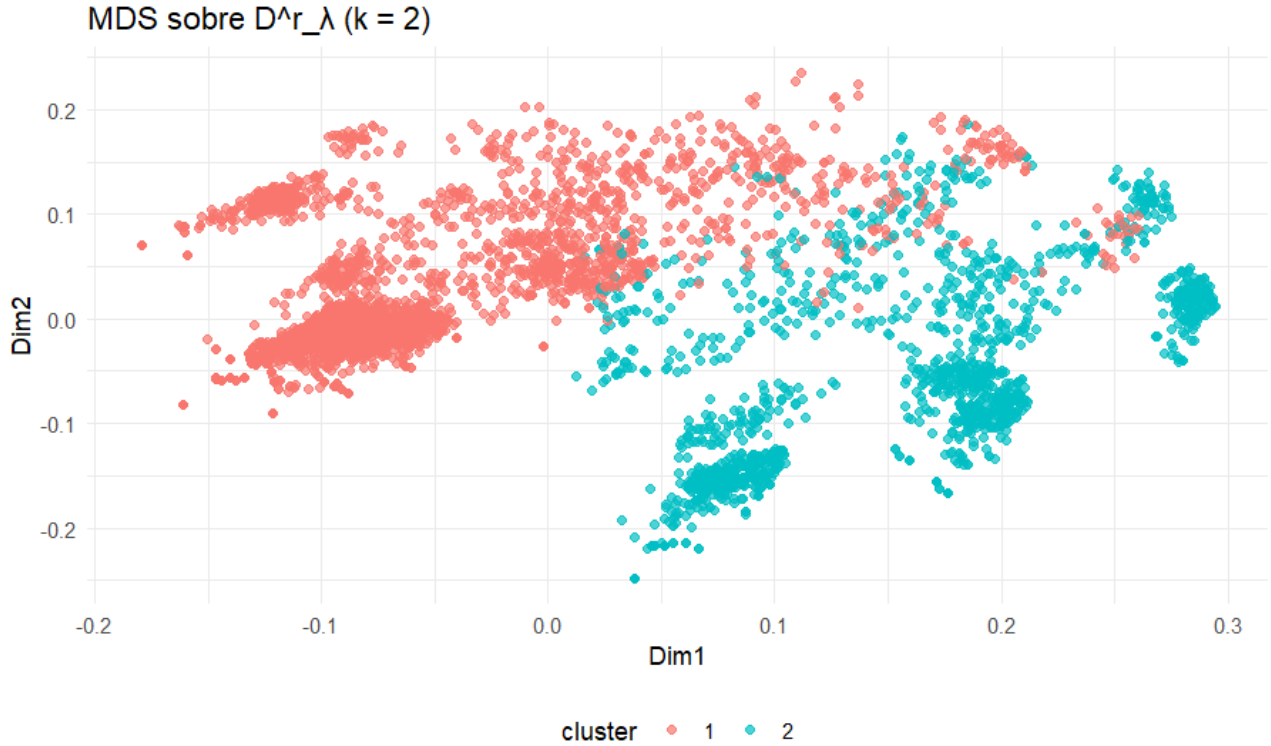


Gráfico 5.2: Representación en dos dimensiones (MDS) de los clústeres de votantes obtenidos con $k = 2$ y distancia D_λ^r .

En la proyección se pueden ver dos agrupamientos grandes de votantes. El clúster 1 contiene a 3 687 votantes (aproximadamente un 71 % de la muestra), mientras que el clúster 2 a 1 518 votantes (alrededor de un 29 %). Sin embargo, al tener la zona central cierta mezcla de colores, se confirma que la separación, si bien existe, no es extrema. Esta observación concuerda con la anchura media moderada del índice de silhouette y con la disparidad entre las anchuras medias de los clústeres (0.32 frente a 0.16).

En conjunto, la distancia combinada D_λ^r ha permitido medir la divergencia entre votantes considerando simultáneamente el ordenamiento de las alternativas y su aprobación. El análisis de silhouette identifica dos clústeres principales, apoyando la hipótesis de polarización, pero el valor moderado del índice (0,271) y la menor cohesión del segundo cluster parecen sugerir que la separación no es muy grande. Analicemos más profundamente a qué puede deberse esta separación.

5.2.1. Diferencias internas entre los dos clústeres

Para comprender mejor el porqué de esta separación, se ha examinado cómo difieren los dos clústeres en las *aprobaciones* y en la *posición media* de cada alternativa.

En el gráfico 5.3 se observa que donde hay mayores discrepancias de aprobación es en el *termómetro musulmán*, *gay* e *inmigrantes*, todos ellos con una aprobación significativamente mayor en el clúster 1. Por otro lado, las diferencias a favor del clúster 2 son mas pequeñas, destacando de manera leve el colectivo *blanco*.

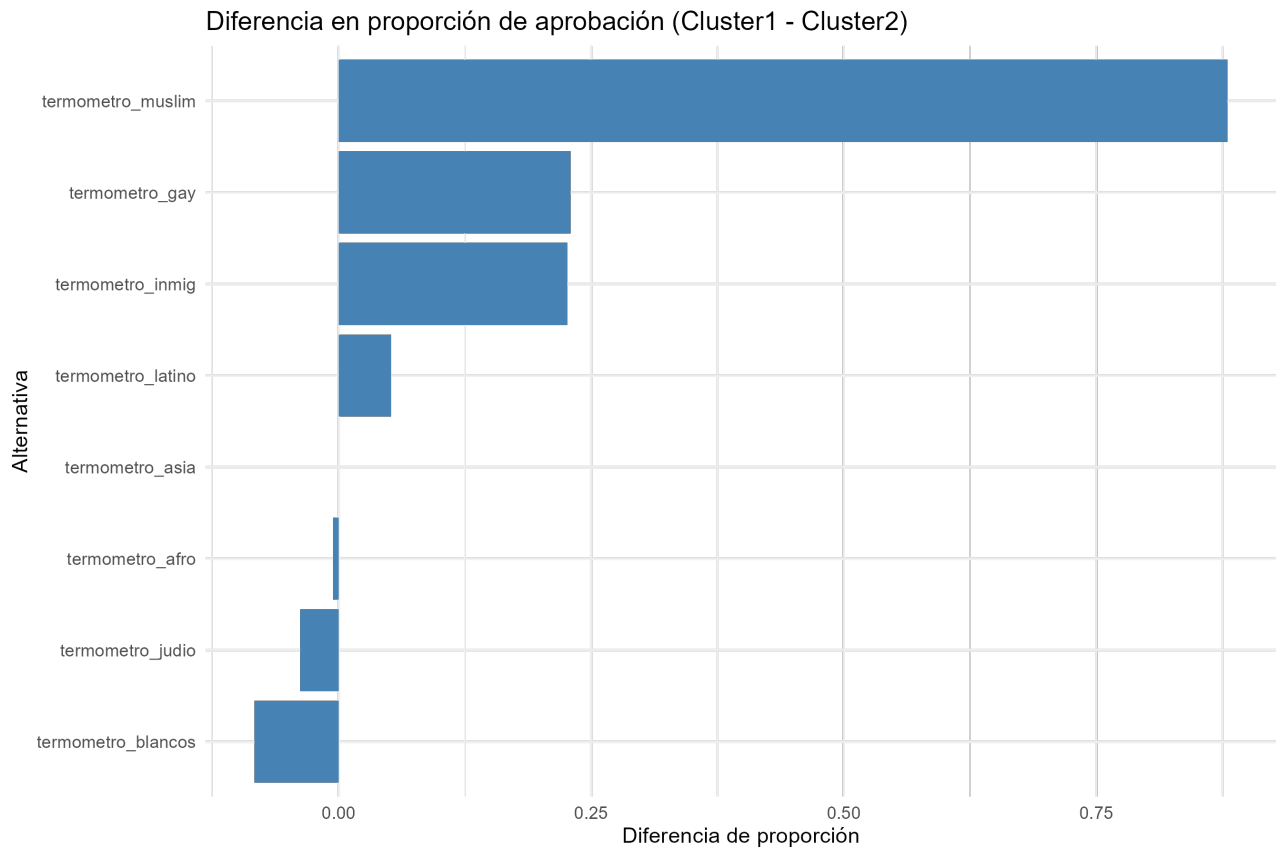


Gráfico 5.3: Diez alternativas con mayor diferencia en la proporción de aprobación (*Cluster 1* – *Cluster 2*). Las barras positivas indican mayor aprobación en el clúster 1 y las negativas, en el clúster 2.

En el gráfico 5.4 se puede observar que el clúster 2 tiene en una posición más destacada a *blancos*, y en menor medida a *judíos*, *asiáticos* y *afroamericanos*, mientras que el clúster 1 sitúa en posiciones más altas a *musulmanes*, *personas LGTBQ* e *inmigrantes*.

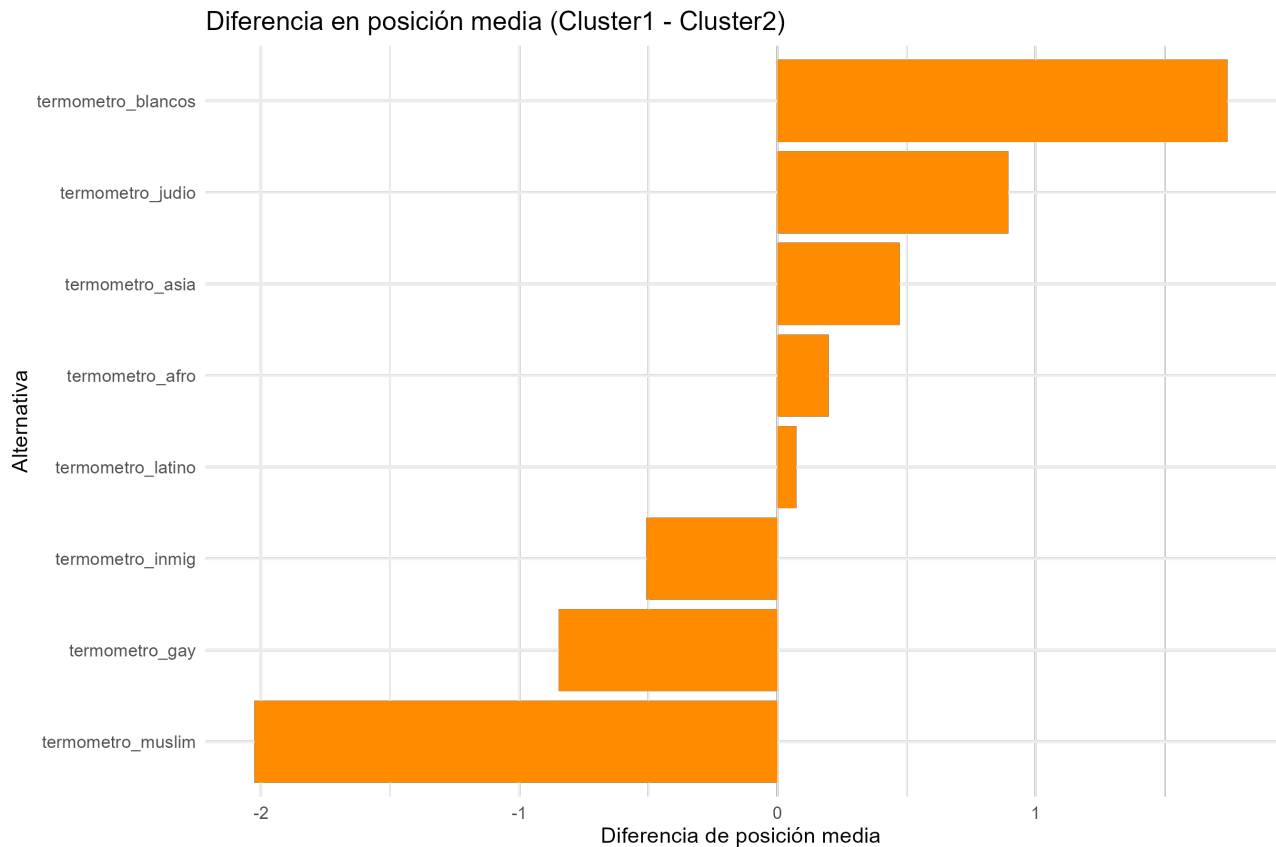


Gráfico 5.4: Diez alternativas con mayor diferencia en la posición media (*Cluster 1* – *Cluster 2*). Los valores negativos indican que la alternativa se sitúa en posiciones más altas en el clúster 2 y los positivos en el clúster 1.

Con esto podemos concluir que estos resultados indican que la principal línea de separación entre los clústeres se debe a la actitud hacia minorías religiosas, étnicas y sexuales, ya que el clúster 1 muestra una aprobación y prioridad significativamente mayores hacia dichos grupos, mientras que el clúster 2 da una importancia relativa algo mayor a los colectivos mayoritarios.

5.2.2. Relación con variables sustantivas

Para analizar si esta división se relaciona con características sociodemográficas o político-ideológicas, se han considerado las variables sustantivas descritas en la tabla 5.2. Como ejemplo representativo, se presenta en el gráfico 5.5 la distribución de la variable *ideología* por clúster. Cada barra corresponde a un cluster y está dividida en función de las categorías de la variable.

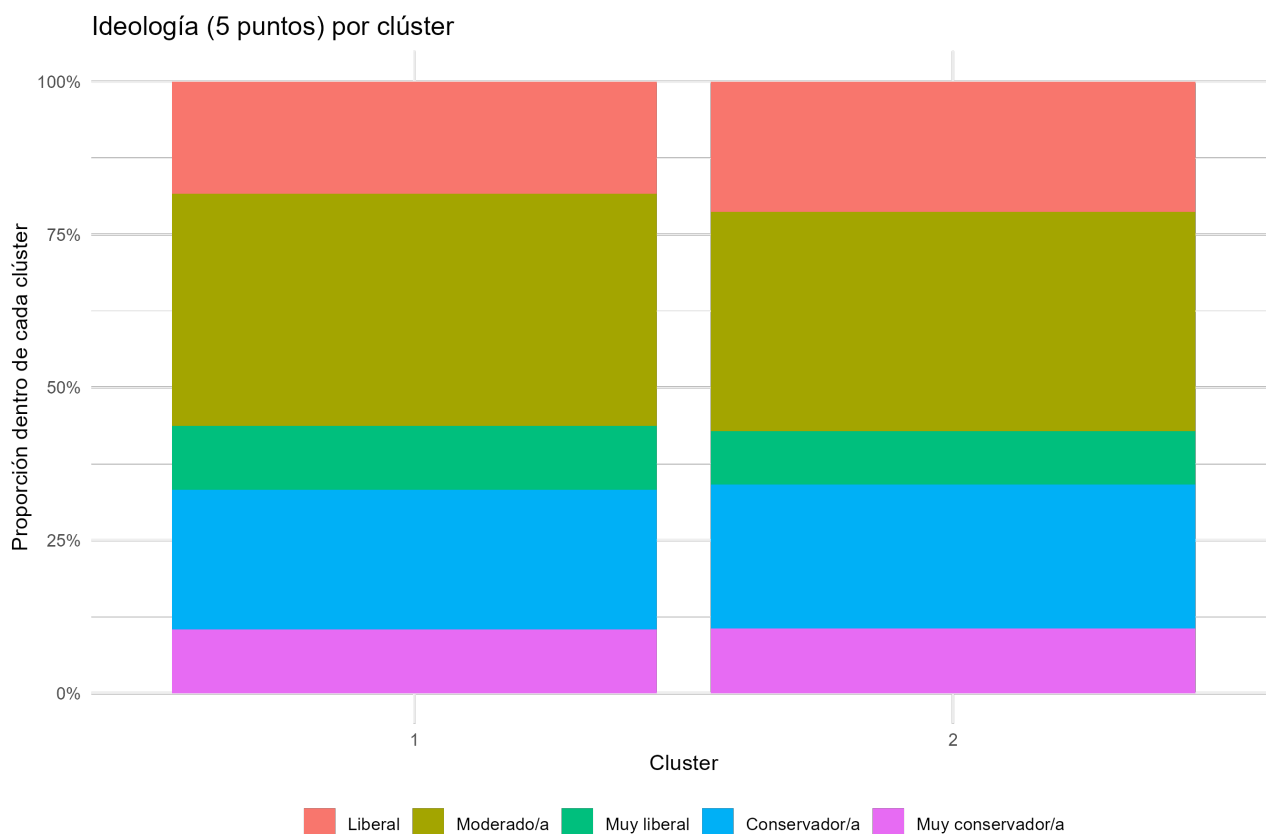


Gráfico 5.5: Distribución de la ideología política (escala de cinco puntos) en cada clúster.

La figura muestra que las proporciones de las distintas categorías ideológicas son muy similares en ambos clústeres. Este patrón se repite prácticamente sin variaciones en todas las demás variables sustantivas analizadas (edad, género, raza, nivel educativo, ingresos familiares, religión, importancia de la religión, posesión de armas, interés por las noticias, identificación partidista y recuerdo/intención de voto en 2016 y 2020), lo que confirma que ninguna de ellas explica la separación detectada.

Este resultado refuerza la conclusión de que la polarización hallada a partir de las preferencias aprobatorias no responde a divisiones sociodemográficas ni a las variables político-ideológicas tradicionales. La existencia de dos grandes clústeres parece deberse, más bien, a diferencias de actitud específicas frente a ciertos colectivos minoritarios, como se ha comentado en la subsección 5.2.1.

Capítulo 6

Conclusiones y otras líneas de investigación

En este trabajo se han estudiado las preferencias aprobatorias, una estructura que combina simultáneamente la información ordinal (orden de las alternativas) con la evaluativa (umbrales de aceptabilidad), y se han aplicado métodos de análisis cluster datos de votaciones reales.

Desde el punto de vista empírico, el estudio del conjunto de datos del *Democracy Fund Voter Study Group* (ola de 2019) ha permitido:

- Codificar las valoraciones de los votantes en forma de preferencias aprobatorias, fijando un umbral de aprobación del 50 %.
- Calcular distancias entre votantes que integran la parte binaria y la ordinal mediante la métrica D_{λ}^r con $\lambda = 0,5$ y $r = 1$.
- Agrupar a los individuos mediante *Ranked K-Medoids*, detectando dos clústeres principales de votantes.
- Verificar que la *separación es moderada*: la anchura media de silhouette (0,271) indica que los grupos son reales pero con solapamiento.
- Identificar la *diferencia clave*: el clúster 1 muestra mayor aprobación y prioridad hacia minorías religiosas, étnicas y sexuales (musulmanes, inmigrantes, colectivo LGTBQ), mientras que el clúster 2 concede algo más de importancia a colectivos mayoritarios.
- Comprobar que las principales variables sociodemográficas y político-ideológicas (edad, género, raza, educación, ingresos, religión, posesión de armas, interés por la política, identificación partidista e ideología) muestran distribuciones muy parecidas en ambos clústeres.

En conjunto, los resultados sugieren que la división detectada no responde a variables como edad, género, clase social o ideología, sino a diferencias de actitud específicas hacia determinados colectivos.

Los hallazgos empíricos confirman varias ideas expuestas en los capítulos teóricos:

- Las preferencias aprobatorias capturan simultáneamente el orden relativo y la aceptabilidad, lo que permite distinguir no sólo qué alternativas son preferidas, sino también cuáles se consideran mínimamente aceptables.

- Las distancias sobre este tipo de preferencias, y en particular la combinación D_{λ}^r , resultan adecuadas para medir la *divergencia real* entre votantes, más allá de la mera coincidencia de primeros puestos en un ranking.
- El uso de *Ranked K-Medoids* demuestra que los algoritmos de clustering basados en distancias son herramientas eficaces para detectar patrones de consenso y polarización en espacios de decisión complejos.

De este modo, el trabajo muestra la utilidad práctica de las herramientas matemáticas presentadas en los capítulos 2 a 4 para analizar datos sociales reales.

El estudio se podría ampliar de ciertas maneras:

- **Identificación de bloques de alternativas.** En lugar de agrupar votantes, podría buscarse la formación de *clústeres de alternativas* con perfiles de aprobación y ordenación semejantes, para descubrir subconjuntos de grupos sociales que tienden a ser aceptables o inaceptables conjuntamente.
- **Análisis de sensibilidad de λ y r .** Sería interesante evaluar cómo varían los resultados al modificar el peso relativo de la parte binaria y la ordinal (λ) o el exponente r en la distancia D_{λ}^r , comprobando la robustez de las conclusiones.
- **Aplicaciones a otros contextos.** La metodología podría aplicarse a encuestas en otros países o a ámbitos como comités de expertos, plataformas de recomendación o procesos de negociación, para estudiar similitudes y diferencias en los patrones de aprobación.
- **Modelos dinámicos.** Analizar varias olas temporales permitiría estudiar la evolución de las preferencias aprobatorias y detectar procesos de convergencia o polarización a lo largo del tiempo.

En definitiva, este Trabajo de Fin de Máster muestra que las preferencias aprobatorias constituyen un marco potente para estudiar la interacción entre consenso y polarización. El análisis empírico revela que, incluso en un contexto preelectoral como el de Estados Unidos en 2019, las divisiones no siempre coinciden con las tradicionales líneas sociodemográficas o ideológicas, sino que pueden surgir de diferencias más específicas en la actitud hacia determinados colectivos.

Bibliografía

- [1] A. Albano, J. L. García-Lapresta, A. Plaia, and M. Sciandra, A family of distances for preference–approvals, *Annals of Operations Research*, vol. 323, pp. 1–29, 2023.
- [2] A. Albano, J. L. García-Lapresta, A. Plaia, and M. Sciandra, Clustering alternatives in preference-approvals via novel pseudometrics, *Statistical Methods & Applications*, 2024.
- [3] M. Balinski and R. Laraki, A theory of measuring, electing and ranking, *Proceedings of the National Academy of Sciences of the United States of America*, vol. 104, no. 21, pp. 8720–8725, 2007.
- [4] M. Balinski and R. Laraki, Majority Judgment: Measuring, Ranking, and Electing, The MIT Press, Cambridge, MA, 2011.
- [5] D. Black, *The Theory of Committees and Elections*, Cambridge University Press, 1958.
- [6] S. J. Brams and P. C. Fishburn, Approval voting, *American Political Science Review*, vol. 72, no. 3, pp. 831–847, 1978.
- [7] S. J. Brams, Mathematics and democracy: Designing better voting and fair-division procedures, *Mathematical and Computer Modelling*, vol. 48, no. 9–10, pp. 1666–1670, 2008.
- [8] S. J. Brams and M. R. Sanver, Voting systems that combine approval and preference, in S. J. Brams, W. V. Gehrlein, and F. S. Roberts (eds.), *The Mathematics of Preference, Choice and Order: Essays in Honor of Peter C. Fishburn*, pp. 215–237. Springer, New York, 2009.
- [9] B. Erdamar, J. L. García-Lapresta, D. Pérez-Román, and M. R. Sanver, Measuring consensus in a preference-approval context, *Information Fusion*, vol. 17, pp. 14–21, 2014.
- [10] M. C. Herron and J. B. Lewis, Did Ralph Nader spoil Al Gore’s Presidential bid? A ballot-level study of Green and Reform Party voters in the 2000 Presidential election, *Quarterly Journal of Political Science*, vol. 2, no. 3, pp. 205–226, 2007.
- [11] G. Santos-García and J. C. R. Alcantud, A characterization of pairwise dictatorial approbatory social welfare functions, *Economics Letters*, vol. 248, 112217, 2025.
- [12] M. R. Sanver, Approval as an intrinsic part of preference, in J. F. Laslier and M. R. Sanver (eds.), *Handbook on Approval Voting*, Studies in Choice and Welfare, pp. 469–481. Springer, New York, 2010.
- [13] Democracy Fund Voter Study Group, *Voter Survey 2019*, Democracy Fund, Washington D.C., 2019.