

UAM-I
✓CBI

✓
CONDICIONES DE SUSTITUABILIDAD BAJO
LA LEY DE WALRAS EN LA TEORIA DEL
EQUILIBRIO GENERAL

TESIS QUE PRESENTA

✓ JAIME MUÑOZ FLORES

PARA LA OBTENCION DEL GRADO DE
✓ MAESTRO EN MATEMATICAS

JUNIO DE 1991 ✓
—

ASESOR DE TESIS:
DR. FELIPE PEREDO RODRIGUEZ

UNIVERSIDAD AUTONOMA METROPOLITANA
IZTAPALAPA



Casa abierta al tiempo

IZT 127988 A
BIBLIOTECA

✓ **CONDICIONES DE SUSTITUABILIDAD BAJO
LA LEY DE WALRAS EN LA TEORIA DEL
EQUILIBRIO GENERAL** ✓

TESIS DE MAESTRIA
MAT. JAIME MUÑOZ FLORES

INDICE.

INTRODUCCION.....	3
CAPITULO 1. CONCEPTOS TOPOLOGICOS.....	8
CAPITULO 2. TEOREMAS FUNDAMENTALES.....	26
CAPITULO 3. APLICACIONES A LA TEORIA DE MERCADOS COMPETITIVOS.....	46
CAPITULO 4. APLICACIONES A LA TEORIA DE CONSUMO.....	61
CAPITULO 5. TEOREMAS DE SCARF, BROUWER Y KAKUTANI.....	72
CAPITULO 6. EQUILIBRIO Y OPTIMOS DE PARETO.....	103
CAPITULO 7. ESTABILIDAD Y CONDICIONES DE SUSTITUABILIDAD.....	124
BIBLIOGRAFIA.....	138

INTRODUCCION.

El campo de aplicación de las matemáticas es cada día más amplio y diverso. Tanto en las Ciencias Naturales como en las Ciencias Sociales, el esfuerzo hacia el rigor sustituye con frecuencia resultados y razonamientos que no eran del todo correctos por los que lo son, ofreciendo también otras ventajas. Además de conducir a un entendimiento más profundo de los problemas de aplicación, puede derivar en un desarrollo de los instrumentos matemáticos involucrados.

Quizás la tarea más difícil del matemático aplicado radica precisamente en lograr situarse en la posibilidad de reunir elementos teóricos cualitativos de otras disciplinas del conocimiento, y traducirlos a un lenguaje formal y preciso, en un intento por reflejar la realidad de un problema de la mejor manera.

En este proceso, aparece casi en la totalidad de las formulaciones un importante número de hipótesis y de supuestos básicos. Como consecuencia, el alcance de los resultados es necesariamente restringido, lo que en muchos casos dificulta significativamente su experimentación y puesta en práctica.

El camino que se sigue para llegar a la formalización de proposiciones en Teoría Económica, por ejemplo, se ve a menudo interrumpido cuando el especialista descubre la necesidad de emplear con soltura conceptos matemáticos particulares, muchos de los cuales merecen cierta profundidad en su estudio.

En dirección contraria, el camino que sigue el matemático para extender y generalizar sus aplicaciones se interrumpe frecuentemente cuando encuentra cambios en las condiciones del problema original, mismos que deben ser examinados cuidadosamente por un economista para que se aseguren el sentido y la consistencia de los resultados.

Es en este contexto en donde se podrían ubicar los desarrollos expuestos en el presente trabajo. Los problemas centrales de la teoría que se discute son: La formulación analítica del papel de los precios en una economía competitiva (Capítulo 2), los teoremas fundamentales para el alcance de un equilibrio competitivo y de un óptimo de Pareto (Capítulo 2, Capítulo 5, Capítulo 6), la aplicación de propiedades de convexidad al análisis de actividades y a las teorías de producción y consumo (Capítulo 3 y Capítulo 4), la existencia de una función de utilidad continua (Capítulo 4), la estabilidad en el modelo y, por último, el papel de la sustituibilidad en el estudio de estabilidad (Capítulo 7).

Como preliminar se incluye un capítulo que comprende conceptos como adherencia, compacidad, numerabilidad, densidad, convexidad y conexidad, entre otros. También se enuncian algunos teoremas como son los de Bolzano y Weierstrass. La lectura de este capítulo podrá omitirla quien tenga frescos estos conceptos.

En el capítulo 2 se presenta la interpretación económica de la propiedad de homogeneidad de grado cero de una función. Se hace también una descripción de precios y mercancías -un tanto vaga e imprecisa-, pero necesaria para la discusión subsecuente. Por último, se demuestran distintas versiones del teorema de Minkowski (hiperplano separador) así como algunos lemas y corolarios relacionados.

El capítulo 3 trata la teoría correspondiente a mercados competitivos. Se introduce además el concepto de análisis de actividades y algunas hipótesis sobre los conjuntos de producción.

Se definen también los llamados *puntos eficientes* y se prueban algunos lemas y teoremas relacionados.

Las aplicaciones a la teoría de consumo aparecen en el capítulo 4. En este mismo capítulo se definen algunos supuestos sobre los conjuntos de consumo, y se establecen ciertas precisiones sobre órdenes y preórdenes, relaciones de indiferencia y equivalencia, contornos, preferencias y, finalmente, se desarrolla la importante demostración debida a G. Debreu [1] acerca de la existencia de una función de utilidad continua.

El teorema de Scarf (que además de servir para generar la prueba del teorema de Brouwer, sirve para establecer un teorema más general debido a Kakutani) se encuentra desarrollado en el Capítulo 5. El algoritmo que se discute ahí no sólo brinda los elementos necesarios para llevar a cabo la prueba; también proporciona un método para calcular un punto fijo hasta cualquier nivel de aproximación que se desee. El costo en cuanto a grado de complejidad de los argumentos que aparecen es probablemente elevado, si se compara con las pruebas tradicionales de teoremas de punto fijo basadas en razonamientos topológicos.

Sin embargo, la proximidad al cálculo computacional para ejercicios de aplicación, y la integridad en la que se prueban los teoremas de Scarf, Brouwer y Kakutani, justifican la discusión del Capítulo.

El capítulo 6 se dedica al desarrollo de resultados referentes a los llamados *óptimos de Pareto*, y a su relación con el equilibrio competitivo. Para el efecto, se incluyen definiciones muy particulares y se demuestran algunos lemas y corolarios intermedios.

El problema de estabilidad se encuentra tratado en el Capítulo 7, en donde se ilustra el análisis mediante la técnica de diagramas de fase. El capítulo termina con la discusión de las condiciones

de sustituibilidad bruta necesarias para las propiedades de estabilidad.

La relación que existe entre el material de cada uno de los capítulos merece algunos comentarios. La teoría del equilibrio general es basta y se encuentra aun en una etapa de desarrollo.

El contenido de las proposiciones comprendidas en éste trabajo, dista mucho de agotar el material existente sobre el tema. También es cierto que quedan fuera muchas otras ideas que influyen directa o indirectamente en la esencia o en la forma de las axiomatizaciones.

La transición de la discusión informal de las interpretaciones a la construcción de resultados es un tanto abrupta, y la parte teórica se desarrolla en algunos casos de manera restringida, incluso cuando se dispone de generalizaciones casi inmediatas.

A pesar de estas salvedades, cabe señalar que los esfuerzos que se han hecho por liberar al material aquí expuesto de sus tradiciones de cálculo han sido fecundos. El cambio hacia la utilización de propiedades topológicas y de convexidad ha resultado en un incremento notable de la generalidad de la teoría. La cantidad y la importancia de problemas que existen sin resolver, han abierto un panorama prometedor para las investigaciones en matemáticas aplicadas y el trabajo interdisciplinario.

Finalmente, quisiera expresar un agradecimiento muy especial para el Dr. Felipe Peredo Rodríguez, por su invaluable orientación para el desarrollo del presente.

CAPITULO I. CONCEPTOS TOPOLOGICOS.

Sea X un subconjunto de $\mathbb{R}^m$; un punto x_0 es adherente a X si hay una sucesión de puntos de X que tiende a x_0 . Puede decirse también, de forma menos precisa, que x_0 es adherente a X si hay puntos de X que esten arbitrariamente cerca de x_0 . Cualquier punto x de X es adherente a X ; basta tomar una sucesión cuyos puntos son todos iguales a x .

El conjunto de los puntos de $\mathbb{R}^m$ adherentes a X se denomina la adherencia (frecuentemente también la clausura) de X , y se representa por $\bar{X}$. Obsérvese que:

$$(1) \bar{X} \subset X$$

también resulta cierto que que:

$$X \subset Y \text{ implica que } \bar{X} \subset \bar{Y}.$$

Puede demostrarse que:

Si X es un subconjunto de $\mathbb{R}^m$, su adherencia ($\bar{X}$) es cerrada.

Es decir, $\bar{X} = \overline{\bar{X}}$. Como todo conjunto cerrado que contiene a X contiene evidentemente a $\bar{X}$, la adherencia de X puede caracterizarse como el conjunto cerrado más pequeño que contiene a X .

Se puede probar que:

Si Ξ es una colección de subconjuntos cerrados de $\mathbb{R}^m$, su intersección $\bigcap_{x \in \Xi} X$

es un subconjunto cerrado de $\mathbb{R}^m$.

También resulta cierto que:

Si $X_1, X_2, \dots, X_n$ son subconjuntos cerrados de $\mathbb{R}^m$, su unión $\bigcup_{i=1}^n X_i$, es un subconjunto cerrado de $\mathbb{R}^m$.

Ejemplo: Sea k un número real positivo y considérese en $\mathbb{R}^2$ el cuadrado (con perímetro excluido) $X = \{x \in \mathbb{R}^2 \mid |x_i| < k \text{ para } i=1,2\}$; entonces $\bar{X} = \{x \in \mathbb{R}^2 \mid |x_i| \leq k \text{ para } i=1,2\}$.

Utilizando lo anterior se obtiene que $\bar{Q} = \mathbb{R}$; la adherencia del conjunto de los números racionales es el conjunto de los números reales.

Un subconjunto X de $\mathbb{R}^m$ es cerrado si contiene sus puntos adherentes. Esto puede expresarse como $\bar{X} \subset X$. Así, en virtud de (1), X es cerrado si y sólo si es igual a su adherencia. También puede decirse que un subconjunto es cerrado si y sólo si todos los puntos a distancia cero del conjunto pertenecen al conjunto.

Ejemplo: En el ejemplo anterior, el subconjunto X de $\mathbb{R}^2$, y Q el subconjunto de $\mathbb{R}$, no son cerrados; $\bar{X}$ es cerrado.

Sea S un subconjunto de $\mathbb{R}^m$; para un subconjunto X de S se define la adherencia de X en S como el conjunto de puntos de S adherentes a X . De modo similar, X se dirá cerrado en S si posee los puntos de S adherentes a X .

Sea k un número real positivo, y tómesese $S = \{x \in \mathbb{R}^2 \mid |x_i| < k \text{ para } i=1,2\}$, y $X = \{x \in S \mid x_i \geq 0\}$. Un croquis mostrará claramente que el conjunto X es cerrado en S , pero no cerrado en $\mathbb{R}^2$. Nótese también que todo subconjunto T de $\mathbb{R}^m$ es cerrado en T .

Intuitivamente, un conjunto S es *numerable* si tiene a lo sumo tantos elementos como $\mathbb{N}$ (el conjunto de los naturales). De manera más precisa, un conjunto se define como numerable si puede ponerse en correspondencia biunívoca con un subconjunto de $\mathbb{N}$. Cuando el conjunto numerable S es no vacío, se puede siempre escoger el subconjunto correspondiente de $\mathbb{N}$ de manera que posea el 1 y, además, siempre que le pertenezcan dos números enteros positivos, le pertenezcan los enteros positivos comprendidos entre ambos. La imagen de $x \in S$ en la correspondencia se denominará entonces el *rango* de x . Así un conjunto es numerable si, y sólo si, todos sus elementos pueden ser ordenados sin que haya dos elementos con el mismo rango.

Puede demostrarse que:

Todo subconjunto S de $\mathbb{R}^m$ contiene un conjunto numerable X tal que $S \subset \bar{X}$.

Es decir, un subconjunto arbitrario S (por lo tanto quizás no numerable) de $\mathbb{R}^m$ contiene un conjunto numerable X que es *denso* en S , es decir, tal que para cualquier punto x de S hay puntos de X arbitrariamente cerca de x .

Ejemplo: Sea S el mismo $\mathbb{R}^m$. De $\bar{\mathbb{Q}} = \mathbb{R}$ se sigue que $\bar{\mathbb{Q}}^m = \mathbb{R}^m$. Así $\mathbb{Q}^m$, el producto cartesiano de m conjuntos iguales a $\mathbb{Q}$, es decir, el conjunto de puntos de $\mathbb{R}^m$ cuyas coordenadas son todas racionales, es denso en $\mathbb{R}^m$.

Además, $\mathbb{Q}^m$ es numerable.

Definición: Sea X un subconjunto de $\mathbb{R}^m$; un punto $x \in \mathbb{R}^m$ es *interior* a X si no es adherente al complemento de X .

Esto significa, intuitivamente, que x está completamente rodeado por puntos de X . Claramente cualquier punto x de $\mathbb{R}^m$ interior a X pertenece a X . El interior de X es el conjunto de los puntos interiores a X , a saber $C(\bar{C}X)$.

Ejemplo: Sea k un número real positivo, y considérese en $\mathbb{R}^2$ el

cuadrado cerrado $X = \{x \in \mathbb{R}^2 \mid |x_i| \leq k \text{ para } i = 1, 2\}$; su interior es $\{x \in \mathbb{R}^2 \mid |x_i| < k \text{ para } i = 1, 2\}$. El conjunto Q^2 y la recta $Y = \{y \in \mathbb{R}^2 \mid y_2 = 0\}$ tienen interiores vacíos en $\mathbb{R}^2$.

Como en el caso de los conjuntos anteriores, si X es un subconjunto de S (que a su vez es subconjunto de $\mathbb{R}^m$), se dice que un punto $x \in S$ es interior a X en S si no es adherente a $C_S X$, el complemento de X en S .

Considérese como ejemplo la recta en $\mathbb{R}^2$, $S = \{z \in \mathbb{R}^2 \mid z_1 = 0\}$ y su subconjunto $Z = \{z \in \mathbb{R}^2 \mid z_1 = 0, 0 \leq z_2 \leq 1\}$. El interior de Z en S es $\{z \in \mathbb{R}^2 \mid z_1 = 0, 0 < z_2 < 1\}$; su interior en $\mathbb{R}^2$ es vacío. Sea X un subconjunto de $\mathbb{R}^m$; diremos que un punto de $\mathbb{R}^m$ es un punto fronterizo si es adherente tanto a X como a su complemento CX . La frontera de X es el conjunto de sus puntos fronterizos, a saber $\bar{X} \cap C\bar{X}$.

Ejemplo: En el último grupo de ejemplos, la frontera del cuadrado cerrado X es su perímetro, la de Q^2 es $\mathbb{R}^2$.

Sea X un subconjunto de $\mathbb{R}^m$; el exterior de X es el complemento de su adherencia, a saber, $C\bar{X}$. Se puede demostrar que el interior, la frontera y el exterior de X forman una partición de $\mathbb{R}^m$, o, de forma equivalente, que el interior y la frontera de X forman una partición de la adherencia de X . La última consideración nos lleva a que el conjunto de X es cerrado si y sólo si contiene a su frontera.

Sea x un elemento de $\mathbb{R}^m$; su norma $|x|$ es el número real $\max\{|x_1|, |x_2|, \dots, |x_n|\}$, es decir, el máximo de los valores absolutos de sus coordenadas.

Sea k un número real no negativo; el conjunto $K = \{x \in \mathbb{R}^m \mid |x| \leq k\}$ es el cubo cerrado de $\mathbb{R}^m$ con centro en cero y lado $2k$.

Un subconjunto S de $\mathbb{R}^m$ se dirá acotado si está contenido en algún cubo cerrado K . Un conjunto no acotado es, pues, un conjunto que tiene puntos arbitrariamente lejos del origen.

Ejemplo: El conjunto J de los números enteros (un subconjunto de $\mathbb{R}$) no es acotado.

Se dice que un subconjunto S de $\mathbb{R}^m$ es compacto si es cerrado y acotado. Puede demostrarse que si $S_1, S_2, \dots, S_n$ son subconjuntos cerrados (o compactos) de $\mathbb{R}^{m_1}, \mathbb{R}^{m_2}, \dots, \mathbb{R}^{m_n}$, entonces $S_1 \times S_2 \times \dots \times S_n$ es un subconjunto cerrado (o compacto) de $\mathbb{R}^{m_1+m_2+\dots+m_n}$.

Por último, se dice que un subconjunto S de $\mathbb{R}^m$ es conexo si no es la unión de subconjuntos no vacíos, disjuntos y cerrados. En otros términos, un conjunto S es conexo si no puede hacerse una partición de S en dos subconjuntos cerrados en S .

Sea p un elemento de $\mathbb{R}^m$ distinto de 0, y c un número real; el conjunto $H = \{z \in \mathbb{R}^m \mid pz = c\}$ es un hiperplano con normal p . Si z' es un punto de H (también se dice que H pasa por z') se tiene que $pz' = c$, y la expresión anterior puede escribirse de nuevo como $H = \{z \in \mathbb{R}^m \mid p(z-z') = 0\}$. Así el hiperplano H es el conjunto de puntos z de $\mathbb{R}^m$ tales que $z-z'$ es ortogonal a p ; se dice que p y H son ortogonales. Si multiplicamos p y c por un mismo número real distinto de cero, el hiperplano H no cambia. Una intersección de hiperplanos se denomina una *variedad lineal*.

Dado un hiperplano H con normal p , se dice que el punto z de $\mathbb{R}^m$ está por encima de H si $pz > c$. El semiespacio cerrado encima de H es $\{z \in \mathbb{R}^m \mid pz \geq c\}$. Se obtienen definiciones parecidas sustituyendo encima, $>$, $\geq$ por debajo, $<$, $\leq$. Se puede ver que un semiespacio cerrado es cerrado y convexo. Lo mismo sucede para un hiperplano, ya que es la intersección de dos semiespacios cerrados, y por lo tanto para una variedad lineal.

Se dice que un hiperplano H es acotador para un subconjunto S de $\mathbb{R}^m$ si S está contenido en uno de los dos semiespacios cerrados determinados por H .

Teorema: Para todo vector distinto de cero n y todo punto P_0 , $n(P-P_0) = 0$, es una ecuación vectorial de un plano que pasa por P_0 y tiene como normal a n .

Demostración:

Hay que probar que $\mathcal{L} = \{P \mid n(P-P_0) = 0\}$ es un plano que pasa por P_0 y tiene n como normal. Necesitamos demostrar tan sólo que existen vectores linealmente independientes a y b ambos ortogonales a n . Entonces, n es una normal al plano

$$\mathcal{P} = \{P_0 + ua + vb \mid u, v \in \mathbb{R}\}$$

y en consecuencia $\mathcal{L} = \mathcal{P}$.

Considérese $n = (n_1, n_2, n_3)$. Como $n \neq 0$, al menos uno de sus componentes es distinto de cero. Supongamos que $n_1 \neq 0$. Entonces

$$a = (-n_2, n_1, 0)$$

es un vector distinto de cero ortogonal a n . Como a y n son vectores ortogonales distintos de cero, son linealmente independientes. Entonces $b = n \times a$ es un vector distinto de cero ortogonal a n y a a . Luego a y b son vectores linealmente independientes cada uno de los cuales es ortogonal a n , lo que completa la prueba.

Sea f una función de S a T , y considérese un punto $x_0 \in S$. La función f es continua en el punto x_0 si

$$x_q \rightarrow x_0, y_0 = f(x_0) \text{ y } y_q = f(x_q) \text{ implica que } y_q \rightarrow y_0$$

En otros términos, f es continua en x_0 si la imagen de cualquier sucesión que tiende a x_0 es una sucesión que tiende a la imagen de x_0 .

Ejemplo: Sea f la función de una variable real con valores reales (es decir, una función de $\mathbb{R}$ a $\mathbb{R}$) definida por $y=1/x$ si $x \neq 0$, $y = 0$ si $x = 0$. Como su gráfico sugiere inmediatamente, la función f es continua en cualquier punto de $\mathbb{R}$, exceptuando el cero.

La función f es continua en S si es continua en cada uno de los puntos de S .

Sean S_1, S_2, T subconjuntos de $\mathbb{R}^{m_1}, \mathbb{R}^{m_2}, \mathbb{R}^n$ respectivamente, y sea f una función de S_1 a S_2 , y g una función de S_2 a T . Defínase una función h de S_1 a T por $h(x) = g(f(x))$ para todo x en S_1 . Es inmediato que:

Si f es continua en el punto x de S_1 y si g es continua en el punto $f(x)$ de S_2 , entonces h es continua en x .

Para cada $k = 1, 2, \dots, p$ sea T_k un subconjunto de $\mathbb{R}^{n_k}$ y considérese el producto $T = \prod_{k=1}^p T_k$. Es inmediato que la proyección sobre T_k es continua en T .

Con las notaciones anteriores, sea f_k una función de S a T_k , y defínase una función f de S a T por $f(x) = f_k(x)$ para todo x en S , donde $(f_k(x))$ indica la f -upla $(f_1(x), f_2(x), \dots, f_k(x))$. Es inmediato que:

Si todas las f_k son continuas en el punto x de S , entonces f es continua en x .

Se puede demostrar que:

Una función de S a T es continua en S si y sólo si la imagen inversa de todo conjunto cerrado en T es cerrada en S .

En el caso particular importante en que f es una función de valor

real en S , es decir, si f es una función de S a $\mathbb{R}$, se puede demostrar que:

Una función de S a $\mathbb{R}$ es continua en S si sólo si la imagen inversa de todo intervalo de $\mathbb{R}$ de la forma $(-\alpha, y]$ ó $[y, \alpha)$ es cerrada en S . En el caso de $\mathbb{R}^+$, el conjunto de los números reales positivos, considérese la función de $S = \mathbb{R} \times \mathbb{R}$ a $\mathbb{R}$ definida por $y = f(u, v) = u/v$.

"La función f es continua en S " equivale a decir "la imagen inversa de $(-\alpha, y]$, es decir, el conjunto $\{(u, v) \in S \mid u - yv \leq 0\}$ es cerrado en S , y lo mismo para $[y, \alpha)$ "; la segunda afirmación se sugiere inmediatamente por un dibujo en $\mathbb{R}^2$, y se demuestra con relativa facilidad.

Sea f una función de S a T . Si f es continua en S , y si S es compacto, entonces $f(S)$ es compacto.

Aplicando esto último al caso particular en que f tiene valores reales, se obtiene inmediatamente el teorema de Weierstrass, que se enuncia como sigue:

Sea f una función de S a $\mathbb{R}$. Si f es continua en S , y si S es compacto y no vacío, entonces $f(S)$ tiene un máximo y un mínimo.

Sea f una función de S a T . Si f es continua en S y S es conexo, entonces $f(S)$ es conexo.

Aplicando este resultado al caso particular en que f tiene valores reales se obtiene inmediatamente el teorema de Bolzano, que se enuncia como sigue:

Sea f una función de S a $\mathbb{R}$. Si f es continua en S y si S es conexo, entonces $f(S)$ es conexo.

Así, si x_1 y x_2 son dos puntos de S y y es un número real tal que $f(x_1) \leq y \leq f(x_2)$, entonces existe un $x \in S$ tal que $f(x) = y$. En otros términos, f toma todos los valores en $f(x_1)$ y $f(x_2)$. Este resultado se aplica a menudo al caso en que S es un intervalo de $\mathbb{R}$.

Sea S un subconjunto de $\mathbb{R}^m$, T un subconjunto compacto de $\mathbb{R}^n$, (x_q) una sucesión de puntos en S , (y_q) una sucesión de puntos en T y φ una correspondencia de S a T .

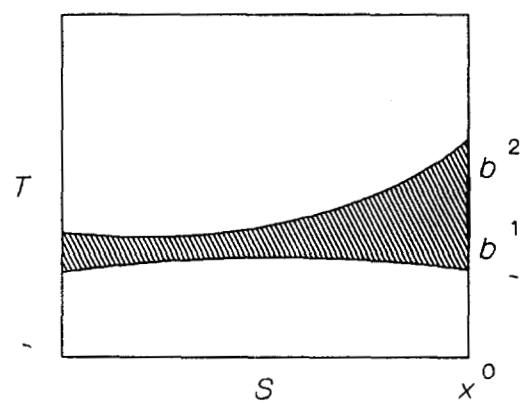
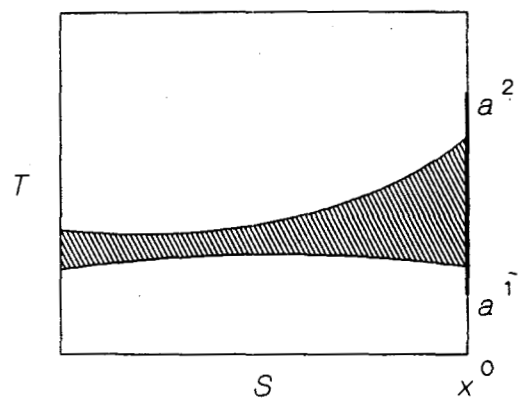
El concepto de continuidad de una correspondencia será introducido en tres etapas.

Sea x_0 un punto de S ; se define en primer lugar:

La correspondencia φ es semicontinua superiormente en el punto x_0 si

$$x_q \rightarrow x_0, y_q \in \varphi(x_q), y_q \rightarrow y_0 \text{ implica } y_0 \in \varphi(x_0).$$

En otros términos, si x_q tiende a x_0 , y si y_q tiende a y_0 al mismo tiempo que pertenece al conjunto imagen de x_q para todo q , se requiere que y_0 pertenezca al conjunto imagen de x_0 . Puede decirse también: Si x_q tiende a x_0 , y si la distancia de y_0 al conjunto imagen de x_q tiende a cero, se requiere que pertenezca al conjunto imagen de x_0 . Esta es una condición de continuidad natural aunque más débil, como se muestra a continuación:



El gráfico de φ es la región sombreada, incluidas las fronteras; $\varphi(x_0)$ es el intervalo $[a_1, a_2]$. La correspondencia φ es semicontinua superiormente en x_0 . Se define en segundo lugar:

La correspondencia φ es semicontinua inferiormente en el punto x_0 si

$x_q \rightarrow x_0$, $y_0 \in \varphi(x_0)$ implica que existe (x_q) tal que $y_q \in \varphi(x_q)$ y $y_q \rightarrow y_0$.

En otros términos, si x_q tiende a x_0 , y si y_0 pertenece al conjunto imagen de x_0 , se requiere que exista una sucesión infinita (y_q) tal que y_q tienda a y_0 al mismo tiempo que pertenezca al conjunto imagen de x_q para todo q .

Puede decirse también que si x_q tiende a x_0 , y si y_0 pertenece al conjunto imagen de x_0 , se requiere que exista una sucesión infinita (y_q) tal que y_q tienda a y_0 al mismo tiempo que pertenezca al conjunto imagen de x_q . Nótese cómo los papeles de " $y_0 \in \varphi(x_0)$ " e " $y_q \in \varphi(x_q)$, $y_q \rightarrow y_0$ " están permutados en las dos definiciones precedentes.

Finalmente se define:

La correspondencia φ es continua en el punto x_0 si es continua superior e inferiormente en x_0 .

Cuando para todo $x \in S$, $\varphi(x)$ consiste en un sólo elemento, las definiciones de semicontinuidad superior e inferior equivalen a la definición de continuidad para una función.

Las semicontinuidades y la continuidad en S se definen como semicontinuidades y continuidad en todos los puntos de S .

Puede verse que una correspondencia φ es semicontinua superiormente en S si y sólo si su gráfico es cerrado en $S \times T$.

Sean S_1 , S_2 , T subconjuntos de $\mathbb{R}^{m_1}$, $\mathbb{R}^{m_2}$ y $\mathbb{R}^n$ respectivamente, sea f una función de S_1 a S_2 , y φ una correspondencia de S_2 a T . Definamos una correspondencia ψ de S_1 a T por $\psi(x) = \varphi(f(x))$ para todo x en S_1 . Si f es continua en el punto x de S_1 , y si φ es semicontinua superiormente (o semicontinua inferiormente) en el punto $f(x)$ de S_2 , entonces ψ es semicontinua superiormente (o semicontinua inferiormente) en x .

Sea T_k un subconjunto de $\mathbb{R}^{n_k}$ para cada $k = 1, \dots, p$ y sea φ_k una correspondencia de S a T_k . Considérese el producto $T = \prod_{k=1}^p T_k$ y defínase una correspondencia φ de S a T por $\varphi(x) = \prod_{k=1}^p \varphi_k(x)$ para todo x de S . Si todos los T_k son compactos, T también lo es en virtud de lo anterior. En este caso puede verse que:
 Si cada φ_k es semicontinua superiormente (o semicontinua inferiormente) en el punto x de S , entonces φ es semicontinua superiormente (o semicontinua inferiormente) en x .

Como ejemplo, considérese a f como una función continua en $S \times T$ con valores reales, e interpretemos $f(x, y)$ como la ganancia de un agente económico cuando su ambiente es x y su acción es y .
 Dado x , hay interés por conocer los elementos $\varphi(x)$ que maximizan a f (ahora una función solamente de y) en $\varphi(x)$ forman un conjunto $\mu(x)$. Qué puede decirse de la continuidad de la correspondencia μ de S a T ? También hay interés en $g(x)$, el valor máximo de f en $\varphi(x)$ para un x dado. Qué puede decirse de la continuidad de la función con valores reales g en S ? El resultado siguiente da una respuesta a esas preguntas.

Si f es continua en $S \times T$, y si φ es continua en $x \in S$, entonces μ es semicontinua superiormente en x , y g es continua en x .

Ejemplo: Sean S y T el intervalo real $[0, 1]$. Defínase φ por $\varphi(x) = [0, 1]$ para todo $x \in S$, y f por $f(x, y) = xy$. Para $x \neq 0$, $\mu(x)$ consiste en el único elemento 1; para $x = 0$, $\mu(x) = [0, 1]$. Para todo x , $g(x) = x$.

Nótese que todo lo anterior supone que T es compacto; es esencial por varias razones. En las aplicaciones, cuando el conjunto T no es compacto, puede ser lícito sustituir T por algún conjunto

compacto sin modificar el problema, y, de este modo, utilizar los resultados anteriores

Sean x^1 y x^2 dos puntos de $\mathbb{R}^m$, y t^1 , t^2 dos números reales tales que $t^1 + t^2 = 1$. El punto $t^1 x^1 + t^2 x^2$ se denomina el *promedio ponderado* de x^1 y x^2 con *ponderaciones* t^1 y t^2 (respectivamente).

Sean x e y dos puntos distintos de $\mathbb{R}^m$. La *recta* xy es $\{z \in \mathbb{R}^m \mid t \in \mathbb{R}, z = (1-t)x + ty\}$. La *semirrecta cerrada* xy (el origen se escribe primero) es $\{z \in \mathbb{R}^m \mid t \in \mathbb{R}, 0 \leq t, z = (1-t)x + ty\}$. La *semirrecta abierta* xy es $\{z \in \mathbb{R}^m \mid t \in \mathbb{R}, 0 < t, z = (1-t)x + ty\}$.

Puede que sea más sugestivo volver a escribir la expresión de z como $z = t(y - x) + x$, es decir, se multiplica el vector fijo $(y - x)$ por el número variable t ; esto nos da la recta (o semirrecta) $0(y-x)$; entonces se le suma el vector fijo x (lo que equivale a una translación).

Sean x , y dos puntos de $\mathbb{R}^m$ (distintos o no).

El *segmento cerrado* xy , designado por $[x, y]$ es $\{z \in \mathbb{R}^m \mid t \in \mathbb{R}, 0 \leq t \leq 1, z = (1-t)x + ty\}$. Cuando $x = y$ se dice que el segmento $[x, y]$ es *degenerado*.

Un subconjunto de $\mathbb{R}^m$ es *convexo* si contiene el segmento cerrado $[x, y]$ toda vez que posea a los puntos x, y ; es decir, si toda vez que contenga a x , y , contenga su promedio ponderado con ponderaciones arbitrarias.

Se puede demostrar que:

la intersección de una colección de conjuntos convexos es convexa.

la suma y el producto de n conjuntos convexos es un conjunto convexo.

la adherencia de un conjunto convexo es un conjunto convexo, y que un conjunto convexo es conexo.

Un subconjunto de $\mathbb{R}$ es conexo si y sólo si es un intervalo.

Sean X e Y dos subconjuntos de $\mathbb{R}^m$. La suma de estos dos conjuntos la constituyen los elementos de la forma $x + y$ con $x \in X$, $y \in Y$. En otros términos, se toman todas las combinaciones posibles de un elemento X y de un elemento Y y se suman; el conjunto obtenido así es $X + Y$.

Ejemplo: Sea $X = \{x \in \mathbb{R}^2 \mid 0 \leq x_1 \leq 1, x_2 = 0\}$, $Y = \{y \in \mathbb{R}^2 \mid y_1 = 0, 0 \leq y_2 \leq 1\}$. Entonces $X + Y = \{z \in \mathbb{R}^2 \mid 0 \leq z_1 \leq 1, 0 \leq z_2 \leq 1\}$.

Se define $-X$ como el conjunto de elementos de $\mathbb{R}^m$ de la forma $-x$ donde $x \in X$. Se escribe $X - Y$ en lugar de $X + (-Y)$.

Si $X_1, X_2, \dots, X_n$ son subconjuntos de $\mathbb{R}^m$, la suma $\sum \bar{X}_j \subset \overline{\sum X_j}$. Es decir, la suma de sus adherencias está contenida en la adherencia de la suma.

Puede probarse también que la suma de n subconjuntos compactos de $\mathbb{R}^m$ es compacta.

Sea S un subconjunto de $\mathbb{R}^m$, y para cada $k = 1, \dots, p$ sea T_k un subconjunto de $\mathbb{R}^n$ y φ_k una correspondencia de S a T_k . Considérese la suma $T = \sum T_k$ y defínase una correspondencia φ de S a T por $\varphi(x) = \sum \varphi_k(x)$ para todo x de S . Si todos los T_k son compactos, también lo es T . En este caso puede demostrarse que si cada φ_k es semicontinua superiormente (o semicontinua inferiormente) en el punto x de S , φ es también semicontinua superiormente (o semicontinua inferiormente) en x .

Dado un subconjunto C de $\mathbb{R}^m$ y un punto x de C , se dice que C es un cono de vértice x si contiene la semirrecta cerrada (x, y) todas las veces que posea el punto y distinto de x .

Dados n conos $C_1, C_2, \dots, C_n$ de vértice cero, se dice que son positivamente semiindependientes si $x_j \in C_j$ para todo j , y $\sum x_j =$

0 implica que $x_j = 0$ para todo j ; es decir, si es imposible encontrar un vector en cada uno de los conos cuya suma sea nula a menos que todos ellos sean nulos. Dos conos C_1 y C_2 son positivamente semiindependientes si y sólo si $C_1 \cap C_2 = \{0\}$.

Considérese un subconjunto S de $\mathbb{R}^m$. Para describir aquellos puntos que se hallan infinitamente lejos de cero se introduce el concepto de cono asintótico como sigue.

Sea k un número real no negativo y désignese por S^k el conjunto $\{x \in S \mid |x| \geq k\}$ de vectores de S cuya norma es mayor o igual que k . Sea $\Gamma(S^k)$ el cono cerrado más pequeño de vértice 0 que contiene a S^k (es decir, la intersección de todos los conos cerrados de vértice cero que contienen a S^k). El cono asintótico de S , designado por AS , se define como la intersección de todos los $\Gamma(S^k)$, es decir, $AS = \bigcap \Gamma(S^k)$; es obviamente un cono cerrado de vértice cero.

Se puede demostrar que:

Un conjunto cerrado, convexo que posee a 0, contiene a su cono asintótico.

Sea S un subconjunto de $\mathbb{R}^m$. Su *cápsula convexa*, designada por $\bar{S}$, se define como la intersección de todos los conjuntos convexos que contienen a S . En virtud de que la adherencia de un conjunto convexo es convexa, $\bar{S}$ es en efecto convexa, y puede, por lo tanto, caracterizarse como el conjunto convexo más pequeño que contiene a S .

Se puede demostrar que:

Dados n subconjuntos de $\mathbb{R}^m$, la cápsula convexa de su suma es igual a la suma de sus cápsulas convexas.

La *cápsula convexa cerrada* de S es, por definición la adherencia $\bar{\bar{S}}$ de la cápsula convexa de S . Como la adherencia de todo subconjunto de $\mathbb{R}^n$ es cerrada, $\bar{\bar{S}}$ es, por cierto, cerrado. Puede verse que un conjunto cerrado y convexo que contiene a S , contiene

necesariamente a $\bar{S}$, que puede, por lo tanto, caracterizarse también como el conjunto convexo y cerrado más pequeño que contiene a S .

CAPITULO 2. TEOREMAS FUNDAMENTALES

Para ilustrar el sentido que se le dará al concepto de *homogeneidad de grado k* en funciones o correspondencias de variable vectorial, introduciremos de manera directa los razgos principales de una importante aplicación en el terreno de la teoría económica, misma que ocupará gran parte de nuestra atención en lo subsecuente.

Supóngase que se estudia una economía en la que participa un número definido de agentes que, dadas ciertas reglas de preferencia, deben elegir entre los distintos tipos de bienes disponibles para satisfacer sus necesidades. El sentido que se otorga a cada uno de estos términos es algo diferente al del uso corriente, pero se precisará más adelante; entre tanto, seguiremos hablando de una manera algo simplificada e imprecisa. Se considera a la economía en un momento dado llamado *momento presente*. Una mercancía se caracteriza por sus propiedades físicas, la fecha en que estará disponible y el lugar donde estará disponible.

El precio de una mercancía es la cantidad que debe pagarse (ahora) por la disponibilidad (futura) de una unidad de esa mercancía.

No se ofrece ninguna teoría del dinero y se supone que la economía funciona sin la ayuda de un bien que sirva como medio de cambio. Así, el papel de los precios es el siguiente: con cada mercancía se asocia un número real, su precio. Cuando un agente económico se compromete a aceptar la entrega de una cierta cantidad, el producto de esta cantidad por el precio de la mercancía es un número real anotado en el débito de su cuenta. Este número será llamado la suma pagada por el agente. Similarmente, un compromiso de entrega reusulta en un número real anotado en el activo de su

cuenta, y llamado la suma recibida por el agente.

El saldo de su cuenta, es decir, el valor neto de todos sus compromisos, guía sus decisiones de una forma que será especificada más adelante.

Para vincular el concepto precedente de precio con la noción habitual de una suma de dinero pagada cuándo y dónde la mercancía está disponible, se debe introducir el concepto de precio en determinada fecha, en determinado lugar. Se obtiene entonces, comparando precios en el mismo lugar, en distintas fechas, los tipos de interés y descuento; comparando precios en la misma fecha, en distintos lugares, los tipos de cambio.

El intervalo de tiempo en que la actividad económica tiene lugar se divide en un número finito de *intervalos elementales compactos* de igual longitud. Estos intervalos elementales pueden numerarse en orden cronológico. El origen del primero se llama el momento presente. Su longitud común, que puede ser un año, un mes, un día, etc., se elige lo suficientemente pequeña para que, desde el punto de vista del análisis, todos los momentos de un intervalo elemental sean indistinguibles. Un intervalo elemental será llamado una fecha, y, por lo tanto, la expresión "en la fecha t " será equivalente a "en algún momento del t -ésimo intervalo elemental". Similarmente, la región del espacio donde la actividad económica tiene lugar se divide en un número finito de regiones elementales compactas. Estas regiones elementales, que pueden numerarse arbitrariamente, se escojen suficientemente pequeñas para que, desde el punto de vista del análisis, todos los puntos de una de ellas sean indistinguibles. Una región elemental será llamada un *lugar* y por lo tanto la expresión "en el lugar S " es equivalente a "en algún punto de la s -ésima región elemental". Una mercancía se define, pues, por la especificación de todas sus características físicas, de su fecha de disponibilidad y de su lugar de disponibilidad. Tan pronto como uno de estos tres factores cambia, aparece una mercancía *distinta*.

Por lo que toca al tipo de bienes en el mercado, se harán dos distinciones fundamentales. Lo que se pone a disposición de un agente económico se llama *insumo* para él; lo que es hecho disponible por un agente económico es un *producto* para él. Para algunos agentes los insumos serán representados por números no negativos y los productos por números no positivos. Para otros agentes se hará la convención contraria. Una convención uniforme puede parecer deseable, pero una más flexible hará la interpretación más sencilla.

Una mercancía es un bien o servicio completamente especificado física, temporal y espacialmente. Se supone que hay solamente un número l de mercancías distinguibles entre sí, indicadas por un índice h que va de 1 a l . Se supone también que la cantidad de cada una de ellas puede ser cualquier número real. De ahora en adelante la gran generalidad del concepto de mercancía, tal como ha sido ilustrado anteriormente, debe tenerse siempre en cuenta.

Centrando la atención en los cambios de fecha se obtiene, como un caso particular de la teoría general de las mercancías, una teoría del ahorro, capital e interés.

Similarmente, centrando la atención en cambios de lugar se obtiene, como otro caso particular de la misma teoría general, una teoría de la localización, transporte, comercio internacional e intercambio. La interpretación de los resultados en esos términos, no será tratada en el presente trabajo.

El espacio $\mathbb{R}^l$ se denominará *espacio de mercancías*. Para cada agente económico, un plan de acción completo (confeccionado en este momento para todo el futuro), o, más brevemente, una acción, es una especificación, para cada mercancía, de la cantidad que él hará disponible o que será puesta a su disposición; es decir, una numeración completa de sus insumos y de sus productos. Cualquiera que sea la convención de signos que se adopte, una acción se representará, por lo tanto, por un punto de $\mathbb{R}^l$.

A cada mercancía, digamos la h -ésima, se asocia un número real, su precio, p_h . Este precio puede interpretarse como la cantidad pagada ahora por (o a) un agente por cada unidad de la h -ésima mercancía que será puesta a su disposición (o hecha disponible por él). El precio que se entenderá aquí está estrechamente ligado a lo que se entiende en un *mercado de futuros*. Allí, un contrato de venta se refiere a un bien definido con precisión, a ser entregado en una fecha especificada, en un lugar especificado. El precio a pagar también se especifica ahora, pero se entiende que será pagado en la fecha de entrega en el lugar de entrega. Esta diferencia con el concepto de precio que aquí se utiliza es insustancial. Debe aclararse que existe una diferencia de otro tipo. Mercados de futuros organizados corresponden solamente a un número reducido de bienes, lugares y fechas (no demasiado alejadas del futuro), mientras que aquí se supone implícitamente que existen mercados para *todas* las mercancías.

El signo del precio p_h de una mercancía no es una propiedad intrínseca de la misma; depende de la tecnología, los recursos, y puede decirse que de la economía en general.

El *sistema de precios* será la l -upla $p = (p_1, p_2, \dots, p_l)$; puede claramente representarse por un punto de $\mathbb{R}^l$. El valor de una acción a relativa al sistema de precios p es $\sum p_h a_h$, es decir, el producto interior $p \cdot a$. Imagínese que un cierto bien circula como dinero en el lugar s , en la fecha t , y sea k el índice de la mercancía así definida. Para obtener el precio en s , en t , de la mercancía h -ésima, es decir, el número de unidades de ese dinero que deben pagarse en s , en t para obtener la disponibilidad de una unidad de la mercancía h -ésima, se dividiría p_h por p_k . Haciendo esto para todos los precios en p se obtendría el sistema de precios en s , en t . En vez de referir todos los precios a una unidad monetaria en s , en t , se podría referirlos, por ejemplo, a algún bien o servicio en s , en t . Uno se ve llevado así al concepto general de *sistema de precios en el lugar s , en la fecha*

t , derivado de p por multiplicación por un cierto número real positivo $\lambda^{s,t}$ (determinado por la unidad de valor elegida en s , en t).

En $p^{s,t}$ el lugar s y la fecha t corresponden al pago, el lugar y la fecha implícitamente determinados por h corresponden a la entrega; el primer y segundo par no están relacionados, pudiendo en particular la fecha de pago ser más temprana, simultánea o posterior a la fecha de entrega. p aparece ahora como el sistema de precios en un lugar no especificado, en un momento no especificado (que es a menudo conveniente imaginar como el momento presente).

Sean t^1, t^2 dos fechas tales que $t^1 < t^2$. El número α_{t^1, t^2}^s definido por $p^{s, t^2} = p^{s, t^1} \alpha_{t^1, t^2}^s$ se llama el factor de acumulación en s de t^1 a t^2 .

Por lo pronto se supondrá que p es distinto de cero; en consecuencia α_{t^1, t^2}^s es un número positivo definido de manera única. Su significado es simple: entregando una unidad de valor en s , en t^1 , se reciben α_{t^1, t^2}^s unidades de valor en s , en t^2 .

Cuando en particular, $t^1 = t$ y $t^2 = t + 1$ se define el tipo de interés en s de t a $t+1$ por $i_{t, t+1}^s = \alpha_{t, t+1}^s - 1$. Es la diferencia entre el valor en s , en $t+1$, que se recibe y la unidad de valor en s , en t que se entrega. Los tipos de interés cotizados usualmente, por ejemplo 0.02 ó 2%, son tipos anuales; aquí todos los tipos de interés (y de descuento) son tipos por intervalo de tiempo elemental.

De $\alpha_{t^1, t^2}^s = i_{t^1, t^2}^s + 1$ se deriva que

$$\alpha_{t^1, t^2}^s = (i_{t^1, t^1+1}^s + 1) \dots (i_{t^2-1, t^2}^s + 1)$$

un producto de $t^2 - t^1$ términos. Esto sugiere la definición del tipo de interés en s de t^1 a t^2 , i_{t^1, t^2}^s como

$$\alpha_{t^1, t^2}^s = (1 + i_{t^1, t^2}^s)^{t^2 - t^1}$$

Similarmente, el número positivo β_{t^1, t^2}^s definido por

$$p^{s,t1} = p^{s,t2} \beta_{t2,t1}^s$$

se llama el *factor de descuento en s de t^2 a t^1* . Para recibir una unidad de valor en s, en t^2 , se dan $\beta_{t2,t1}^s$ unidades de valor en s, en t^1 . Podrá verse que

$$\beta_{t2,t1}^s = \frac{1}{\alpha_{t1,t2}^s} + \frac{1}{(1 + i_{t1,t2}^s)}.$$

Se define también el *tipo de descuento en s de t^2 a t^1* , $d_{t2,t1}^s$ por

$$\beta_{t2,t1}^s = (1 - d_{t2,t1}^s)^{t2-t1}$$

Sean s^1 , s^2 dos lugares. El número positivo $\epsilon_t^{s2,s1}$ definido por $p_{t1,t}^{s1} = p_{t2,t}^{s2} \epsilon_t^{s2,s1}$ se llama el *tipo de cambio en t, en s^1 sobre s^2* . Dando $\epsilon_t^{s2,s1}$ unidades de valor en t, en s^1 , se recibe una unidad de valor en t, en s^2 . Lo que se ha dicho acerca de la generalidad del concepto de mercancía podría repetirse ahora acerca del concepto de precio. Debe tenerse siempre presente que cuando el sistema de precios p es conocido y los números $\lambda^{s,t}$ están dados, todos los precios propiamente dichos, todos los factores de acumulación, los tipos de interés, y los tipos de cambio están determinados en cada fecha y en cada lugar.

Centrémonos ahora en una primera aproximación al significado que otorgaremos al término *homogeneidad de grado k*. Cuando las acciones de los agentes dependen de las tasas de cambio de los bienes entre sí, y no de su tasa de cambio frente a la unidad de cuenta, tendremos, expresado de otra manera, que

$$z(p) = z(kp) \quad \text{para todo } p > 0 \text{ y } k > 0$$

en donde la z representa la *función de demanda excedente*, que está

definida en términos de la diferencia entre producción y consumo¹. Algunos supuestos y resultados importantes en torno a esta demanda excedente aparecerán cuando se analice con detenimiento el comportamiento de los precios ante una oferta y una demanda dadas. Por el momento, basta decir que una consecuencia de la *homogeneidad de grado cero* es que se puede fijar arbitrariamente el nivel de p sin restringir de modo alguno el análisis. La forma más conveniente para nuestros fines, será considerar sólo los precios que pertenezcan al simplex S_n de n dimensiones, lo que se define como

$$S_n = \{p \mid \sum p_i = 1, p > 0\}.$$

Por otra parte, pasaremos a continuación a establecer algunos resultados que, además de intervenir de manera importante en los capítulos 6 y 7, representan por cuenta propia material de interés.

¹ Mas adelante se retomara este concepto.

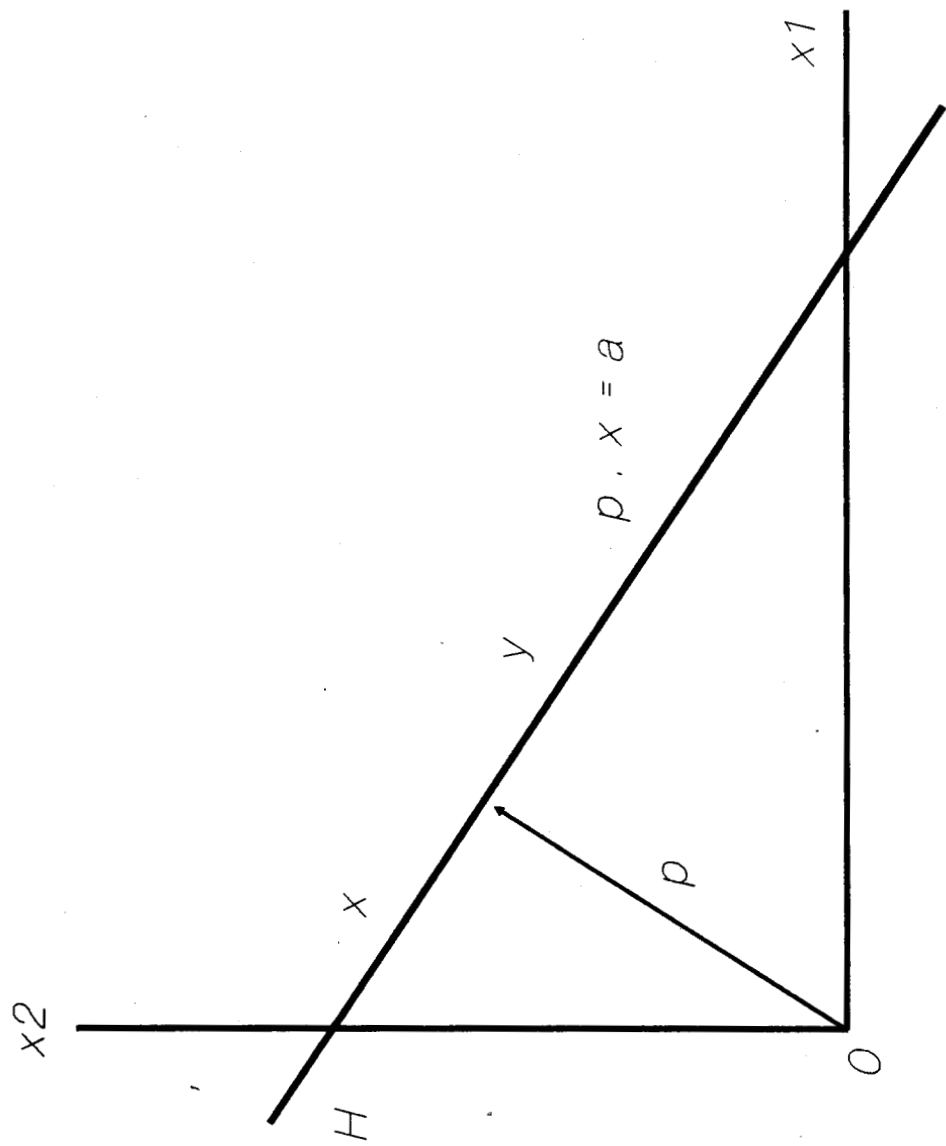
El teorema conocido como teorema de separación, que establece la existencia de un hiperplano que separa dos conjuntos convexos ajenos, es probablemente el teorema más importante en la teoría de optimización. Veremos ahora la esencia de este resultado a través de un caso simple, pero de gran trascendencia; consideraremos el espacio completo $\mathbb{R}^n$, mismo que será considerado simultáneamente como espacio topológico (con la métrica usual) y como espacio vectorial.

Definición: Dados dos subconjuntos no vacíos X y Y en $\mathbb{R}^n$ y un hiperplano $H = \{ x \in \mathbb{R}^n \mid p \cdot x = \alpha \}$, decimos que X y Y están separados por H (o bien, que H separa a X y a Y) si

$$p \cdot x \leq \alpha \quad \text{para todo } x \in X$$

y

$$p \cdot x \geq \alpha \quad \text{para todo } x \in Y$$



Si tenemos desigualdades estrictas en las relaciones anteriores, entonces decimos que X e Y están estrictamente separados por H .

Definición: Dado un conjunto no vacío X en $\mathbb{R}^n$, un hiperplano H se llama *acotador de X* , si X está contenido en alguno de los dos semiespacios determinados por H . Si H es un hiperplano acotador para X y tiene algún punto en común con la frontera de X , entonces H se conoce como el *hiperplano que sustenta a X* .

Probaremos a continuación el teorema de Minkowski en su versión más sencilla.

Teorema 1: Sea $\underline{X}$ un conjunto no vacío, cerrado y convexo de $\mathbb{R}^n$. Sea $x_0 \notin \underline{X}$. Entonces se cumple lo siguiente:

(i) Existe un punto $a \in \underline{X}$ tal que $d(x_0, a) \leq d(x_0, x)$ para todo $x \in \underline{X}$, y $d(x_0, a) > 0$.

(ii) Existen un $p \in \mathbb{R}^n$, $p \neq 0$, $|p| < \infty$, y un $\alpha \in \mathbb{R}$ tales que

$$p \cdot x \geq \alpha \quad \text{para todo } x \in \underline{X}$$

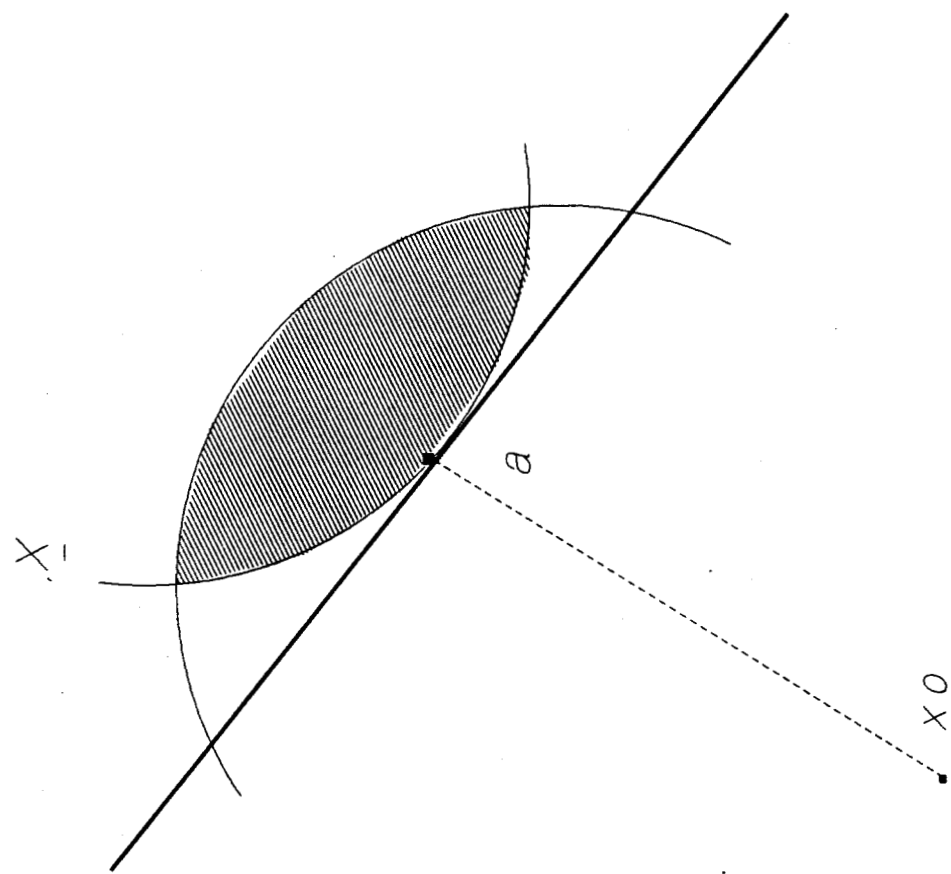
y

$$p \cdot x_0 < \alpha$$

En otras palabras, $\underline{X}$ y x_0 están separados por el hiperplano $H = \{x \in \mathbb{R}^n \mid p \cdot x = \alpha\}$

Prueba:

(i) Sea $\underline{B}(x_0)$ una bola cerrada con centro en x_0 , que comparta elementos con $\underline{X}$ (es decir, $\underline{B}(x_0) \cap \underline{X}$). Escribese $A = \underline{B}(x_0) \cap \underline{X}$. El conjunto A es no vacío, cerrado y acotado (por lo tanto, compacto). Como A es compacto y la función distancia es continua, $d(x_0, x)$ alcanza su mínimo en A como consecuencia del teorema de Weierstrass. Esto es, existe un $a \in A$ tal que $d(x_0, a) \leq d(x_0, x)$



para todo $x \in \underline{X}$. Ya que $x_0 \notin \underline{X}$, $d(x_0, a) > 0$.

(ii) Sea $p = a - x_0$ y $\alpha = p \cdot a$. Obsérvese primero que $p \cdot x_0 = (a - x_0) \cdot (x_0 - a) + (a - x_0) \cdot a = -(a - x_0) \cdot (a - x_0) + (a - x_0) \cdot a = -|p|^2 + \alpha < \alpha$, donde $0 < |p| < \infty$. Sea $x \in \underline{X}$ un punto arbitrario. Como $\underline{X}$ es convexo y $a \in \underline{X}$, $x(t) \in \underline{X}$ donde $x(t) = (1 - t)a + tx$, $0 < t \leq 1$. Entonces $d(x_0, a) \leq d(x_0, x(t))$ por (i). En otros términos: $|a - x_0|^2 \leq |x(t) - x_0|^2 = |(1 - t)a + tx - x_0|^2 = |(1 - t)(a - x_0) + t(x - x_0)|^2 = (1 - t)^2|a - x_0|^2 + 2t(1 - t)(a - x_0) \cdot (x - x_0) + t^2|x - x_0|^2$. De donde obtenemos que $0 \leq t(t - 2)|a - x_0|^2 + 2t(1 - t)(a - x_0) \cdot (x - x_0) + t^2|x - x_0|^2$. Dividiendo ambos lados por t ($t > 0$) obtenemos $0 \leq (t - 2)|a - x_0|^2 + 2(1 - t)(a - x_0) \cdot (x - x_0) + t|x - x_0|^2$. Tomando el límite cuando $t \rightarrow 0$ tenemos que

$$0 \leq -2|a - x_0|^2 + 2(a - x_0) \cdot (x - x_0)$$

o bien

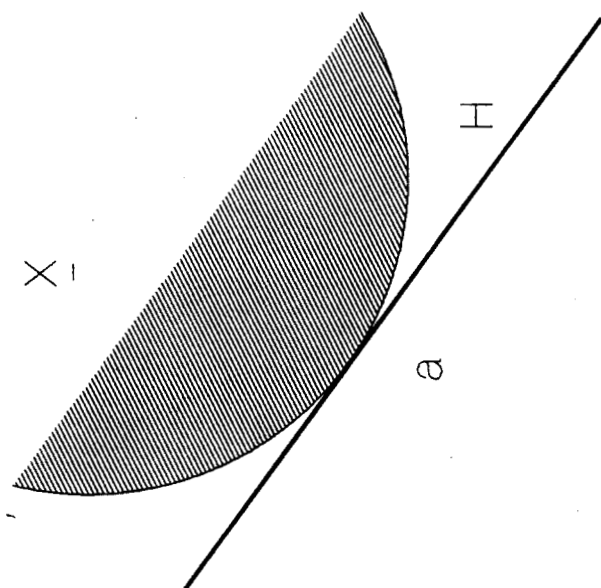
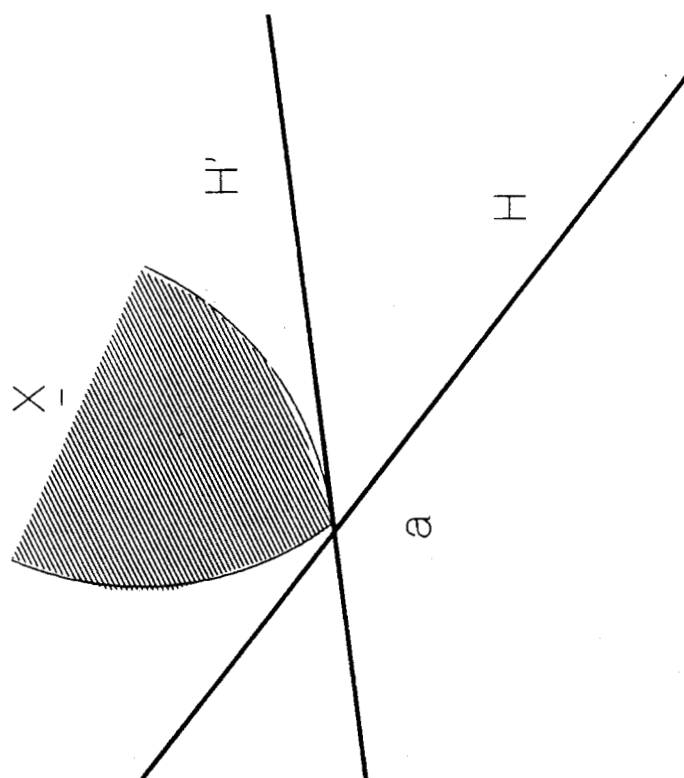
$$0 \geq (a - x_0) \cdot (a - x_0) - (a - x_0) \cdot (x - x_0) = (a - x_0) \cdot a - (a - x_0) \cdot x.$$

Como $p = (a - x_0)$ tenemos que $p \cdot x \geq p \cdot a = \alpha$ para todo $x \in \underline{X}$.

Q. E. D.

Cabe hacer notar lo siguiente: (1) La convexidad de $\underline{X}$ sólo se usa en la parte (ii) de la demostración. (2) Como $0 < |p| < \infty$, podemos encontrar p tal que $|p| = 1$. (3) El hiperplano $H = \{x \in \mathbb{R}^n \mid p \cdot x = \alpha\}$ es un hiperplano que sustenta a $\underline{X}$. Este hiperplano separa el conjunto $\underline{X}$ de el punto dado x_0 (que es, de hecho, un conjunto convexo). Por lo tanto, la existencia de dicho hiperplano es crucial en la prueba. Dado un conjunto no vacío, convexo y cerrado (digamos $\underline{X}$) y un punto (digamos a) en la frontera de $\underline{X}$, podemos asegurar la existencia de al menos un hiperplano que sustenta a $\underline{X}$ y que pasa por el punto a . De hecho, dicho hiperplano juega un importante papel en otras versiones del teorema de separación que probaremos en seguida.

Si la (hiper-) curva que define la frontera de X es suave en el punto a de la frontera, entonces el (hiper-) plano tangente a ese punto resulta ser el que sustenta a X . En este caso, sólo hay un hiperplano que sustenta a X y que pasa por el punto a . Sin embargo, si la (hiper-) curva no es suave, puede haber más de un hiperplano que pase por un punto dado. Estos dos casos se ilustran en la siguiente figura. Es importante notar que los comentarios relativos al hiperplano tangente a un punto, vinculan los teoremas de separación con el cálculo. El poder de los teoremas de separación radica en que la (hiper-) curva frontera no tiene que ser suave en a .



El teorema anterior puede ser visto también de la siguiente forma:

Corolario: Sea $\underline{X}$ un conjunto no vacío, cerrado y convexo en $\mathbb{R}^n$, que no contiene al origen. Entonces existe un $p \in \mathbb{R}^n$, $p \neq 0$, $|p| < \infty$, y un $\alpha > 0$ tales que

$$p \cdot x \geq \alpha \quad \text{para todo } x \in \underline{X}$$

en donde la desigualdad puede ser estricta.

Prueba: En el teorema 1, sea x_0 el origen. Entonces, como resultado del teorema, existen un $p \neq 0$ y un α tales que $p \cdot x \geq \alpha$ para todo $x \in \underline{X}$, donde p y α están definidos como $p = a - x_0 = a$ y $\alpha = p \cdot a$. Entonces, $\alpha > 0$.

Q.E.D.

La desigualdad en el corolario puede ser estricta eligiendo un punto entre a y el origen.

Por otra parte, observemos que la hipótesis de cerradura de $\underline{X}$ puede ser relajada para obtener el siguiente

Teorema 2: Sea X un conjunto no vacío y convexo de $\mathbb{R}^n$. Sea x_0 un punto en $\mathbb{R}^n$ que no está en X . Entonces existe un $p \in \mathbb{R}^n$, $p \neq 0$, $|p| < \infty$, tal que

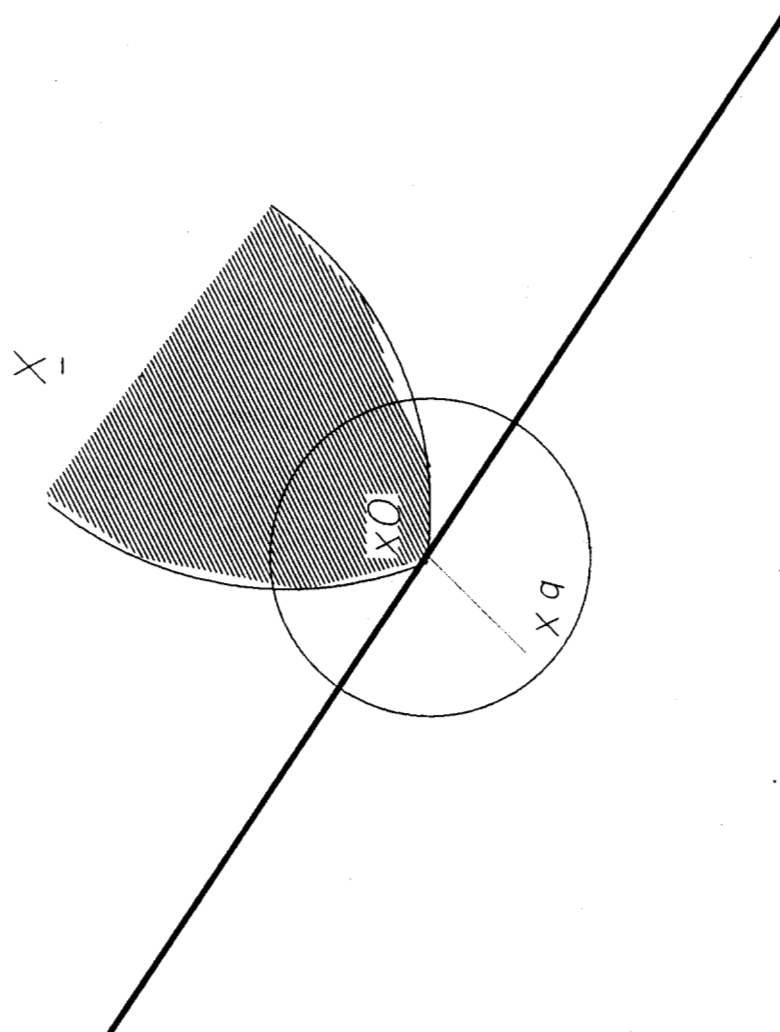
$$p \cdot x \geq p \cdot x_0 \quad \text{para todo } x \in X.$$

Prueba:

(i) Supóngase que $x_0 \in \underline{X}$, en donde $\underline{X}$ es la cerradura de X . Entonces, por el teorema 1, existe un $p \in \mathbb{R}^n$, $p \neq 0$, y un $\alpha \in \mathbb{R}$, tales que $p \cdot x \geq \alpha$ para todo $x \in \underline{X}$ y $p \cdot x_0 < \alpha$. Entonces $p \cdot x > p \cdot x_0$ para todo $x \in \underline{X}$. Luego $p \cdot x > p \cdot x_0$ para todo $x \in X$.

(ii) Supóngase que $x_0 \in \underline{X}$. Ya que $x_0 \notin X$ por hipótesis, x_0 es un

punto de la frontera de $\underline{X}$. Entonces, para toda bola cerrada que contenga a x_0 , existe un punto que no se encuentra en $\underline{X}$. Es decir, existe una sucesión $\{x_q\}$ tal que $x_q \notin \underline{X}$ y $x_q \rightarrow x_0$. Como $x_q \notin \underline{X}$ y $\underline{X}$ es no vacío, cerrado y convexo, existe, por el teorema 1, un $p_q \in \mathbb{R}^n$, $p_q \neq 0$ tal que $p_q \cdot x > p_q \cdot x_q$ para todo $x \in \underline{X}$. Esto se ilustra en la siguiente figura



Ahora, sin pérdida de generalidad, podemos elegir p_q tal que $|p_q| = 1$. Entonces la sucesión $\{p_q\}$ yace en la esfera unitaria de $\mathbb{R}^n$. Ya que la esfera unitaria es compacta, existe una subsucesión convergente en la esfera; es decir, existe una subsucesión tal que $p_{q_s} \rightarrow p$ con $|p| = 1$, donde $\{p_{q_s}\}$ corresponde a $\{x_{q_s}\}$. Tómese el límite de $p_{q_s} \cdot x > p_{q_s} \cdot x_{q_s}$ cuando $q_s \rightarrow \infty$. Como el producto interior es una función continua, tenemos que $p \cdot x \geq p \cdot x_0$ para todo $x \in \underline{X}$. De aquí que $p \cdot x > p \cdot x_0$ para todo $x \in \underline{X}$.

Q.E.D.

Teorema 3 (Minkowski): Sean X y Y dos conjuntos convexos no vacíos en $\mathbb{R}^n$ (no necesariamente cerrados) tales que $X \cap Y = \emptyset$. Entonces existen $p \in \mathbb{R}^n$, $p \neq 0$, $|p| < \infty$, y $\alpha \in \mathbb{R}$ tales que $p \cdot x \geq \alpha$ para todo $x \in X$ y $p \cdot y \leq \alpha$ para todo $y \in Y$.

Prueba: Considérese $S = X + (-Y)$ (el conjunto que se obtiene por la suma de vectores). Como X y Y son convexos, S es convexo. Se puede ver que $0 \notin S$, ya que de lo contrario, existen $x' \in X$ y $y' \in Y$ tales que $x' + (-y') = 0$, o bien $x' = y'$ (lo que contradice que $X \cap Y = \emptyset$). Luego, debido al teorema 2, existe $p \neq 0$ tal que $p \cdot z \geq p \cdot 0 = 0$ para todo $z \in S$. Escribimos $z = x - y$, $x \in X$, $y \in Y$. Entonces tenemos que $p \cdot x \geq p \cdot y$ para todo $x \in X$, $y \in Y$. En otros términos, $\inf_{x \in X} p \cdot x \geq \sup_{y \in Y} p \cdot y$. Podemos luego escoger un α de modo tal que la conclusión del teorema se cumpla.

Q.E.D.

Hagamos ahora las siguientes observaciones: (1) El teorema 2 es un caso particular incluido en el teorema 3, en donde uno de los dos conjuntos consiste solamente en un punto (que es claramente convexo). (2) Hemos utilizado expresiones como

$$p \cdot x \geq \alpha \quad \text{para todo } x \in X \quad (\text{o bien, } p \cdot x > \alpha \text{ para todo } x \in X)$$

y

$p \cdot y \leq \alpha$ para todo $y \in Y$ (o bien, $p \cdot y < \alpha$ para todo $y \in Y$)

Los sentidos de las desigualdades son en realidad irrelevantes para el establecimiento del teorema; podemos cambiarlos definiendo $p = -p$ y $\alpha' = -\alpha$. Entonces $p \cdot x \leq \alpha'$ para todo $x \in X$ (o bien, $p \cdot x < \alpha'$ para todo $x \in X$), y $p \cdot y \geq \alpha'$ para todo $y \in Y$ (o bien, $p \cdot y > \alpha'$ para todo $y \in Y$). (3) Nótese también que al establecer los teoremas anteriores, nada se especificó acerca de los signos de α y de las componentes de p .

Terminaremos la presentación de los teoremas de separación demostrando una de sus aplicaciones importantes: probaremos el teorema de Minkowski - Farkas aplicando el teorema 1. Para hacer esto, necesitaremos primero establecer el siguiente

Lema: Sea K un cono en $\mathbb{R}^n$ con vértice en el origen, y sea p un punto dado en $\mathbb{R}^n$. Si $p \cdot x$ está acotado por abajo para todo $x \in K$, entonces $p \cdot x \geq 0$ para todo $x \in K$.

Prueba:

Por hipótesis, existe un $\alpha \in \mathbb{R}$ tal que $p \cdot x \geq \alpha$ para todo $x \in K$. Ya que K es un cono con vértice en el origen, $x \in K$ implica que $\theta x \in K$ para todo $\theta \geq 0$. De ahí que $p \cdot (\theta x) \geq \alpha$, o bien, $p \cdot x \geq \alpha/\theta$ para todo $x \in K$ y $\theta > 0$. Tomando el límite cuando $\theta \rightarrow \infty$ resulta que $p \cdot x \geq 0$.

Q.E.D.

Ahora podemos demostrar el siguiente

Teorema 4 (Minkowski - Farkas): Sean $a_1, a_2, \dots, a_m$ y $b \neq 0$ puntos de $\mathbb{R}^n$. Supóngase que $b \cdot x \geq 0$ para todo x tal que $a_i \cdot x \geq 0$, $i = 1, 2, \dots, m$. Entonces, existen coeficientes $\lambda_1, \lambda_2, \dots, \lambda_m$, todos ≥ 0 , pero no iguales a cero simultáneamente, tales que

$$b = \sum_{i=1}^m \lambda_i a_i.$$

Prueba: Sea K un cono polihédrico convexo generado por $a_1, a_2, \dots, a_m$. Entonces K es un conjunto cerrado. Queremos mostrar que $b \in K$. Supóngase que $b \notin K$. Entonces K es un conjunto convexo, cerrado y no vacío, que además es ajeno a b . Luego, por el teorema 1, existen un $p \in \mathbb{R}^n$, $p \neq 0$, y un $\alpha \in \mathbb{R}$ tales que

$$p \cdot x \geq \alpha \quad \text{para todo } x \in K$$

y

$$p \cdot b < \alpha$$

Se sigue que $p \cdot x$ está acotado por abajo para todo $x \in K$. Aplicando el lema anterior, tenemos que $p \cdot x \geq 0$ para todo $x \in K$. Nótese que también $0 \in K$ significa que $p \cdot 0 \geq \alpha$, o bien, $\alpha \leq 0$. De modo que $p \cdot b < 0$. Ya que $a_i \in K$ para todo i , $p \cdot a_i \geq 0$ para todo i . Se puede ver entonces que para este p tenemos que $p \cdot b < 0$ con $a_i \cdot p \geq 0$, $i = 1, 2, \dots, m$. Esto contradice la hipótesis del teorema. Por lo tanto, $b \in K$. En consecuencia, existen $\lambda_1, \lambda_2, \dots, \lambda_m$, todos ≥ 0 , tales que $b = \sum_{i=1}^m \lambda_i a_i$. Ya que $b \neq 0$, las λ 's no son cero simultáneamente.

Q.E.D.

Observemos ahora lo siguiente: (1) Si $b = 0$, entonces es posible que $\lambda_i = 0$ para toda $i = 1, 2, \dots, m$. (2) El teorema se cumple también en el sentido inverso. Si $a_1, a_2, \dots, a_m$ y $b \neq 0$ son puntos en $\mathbb{R}^n$, y si existen coeficientes $\lambda_1, \lambda_2, \dots, \lambda_m$, todos ≥ 0 (pero no iguales a cero simultáneamente), tales que $b = \sum_{i=1}^m \lambda_i a_i$, entonces $b \cdot x \geq 0$ para x tales que $a_i \cdot x \geq 0$, $i = 1, 2, \dots, m$. (*Prueba:* Supóngase que $\lambda_1, \lambda_2, \dots, \lambda_m$, todos ≥ 0 , son coeficientes tales que

$$b = \sum_{i=1}^m \lambda_i a_i; \text{ entonces } b \cdot x = \left(\sum_{i=1}^m \lambda_i a_i \right) \cdot x = \sum_{i=1}^m \lambda_i (a_i \cdot x) \geq 0).$$

De acuerdo a las observaciones anteriores, el teorema de Minkowski - Farkas se puede establecer de la siguiente manera:

Teorema 5: Dados los puntos a_i , $i = 1, 2, \dots, m$ y $b \neq 0$ en $\mathbb{R}^n$, solamente una de las dos afirmaciones siguientes se cumple:

(i) Existen λ_i , $i = 1, 2, \dots, m$ todas ≥ 0 (pero no iguales a cero simultáneamente) tales que

$$b = \sum_{i=1}^m \lambda_i a_i$$

(ii) Existe un $x \in \mathbb{R}^n$ tal que

$$b \cdot x < 0 \quad \text{y} \quad a_i \cdot x \geq 0, \quad i = 1, 2, \dots, m.$$

(El teorema se puede enunciar como "si (ii) no se cumple, entonces (i) no se cumple").

Definiendo una matriz $m \times n$ $A = \begin{bmatrix} a_1 \\ a_2 \\ \vdots \\ a_m \end{bmatrix}$ y un vector $x = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix}$ además de

vector fila $\lambda = (\lambda_1, \lambda_2, \dots, \lambda_m)$, podemos establecer las afirmaciones (i) e (ii) de la siguiente manera:

(i) La ecuación $b = \lambda A$ tiene una solución no negativa $\lambda \geq 0$, $\lambda \neq 0$.

(ii) Las desigualdades $b \cdot x < 0$ y $Ax \geq 0$ tienen solución en x .

El teorema 4 puede ser reescrito como " $b \cdot x \geq 0$ para todo x tal que $Ax \geq 0$ implica que existe un $\lambda \geq 0$ con $\lambda \neq 0$, que cumple con $b = \lambda A$ ".

Una interpretación geométrica del teorema cuatro se ilustra en la siguiente figura, en donde se puede ver que

(i) La desigualdad $a_i \cdot x \geq 0$ para toda $i = 1, 2, \dots, m$, significa

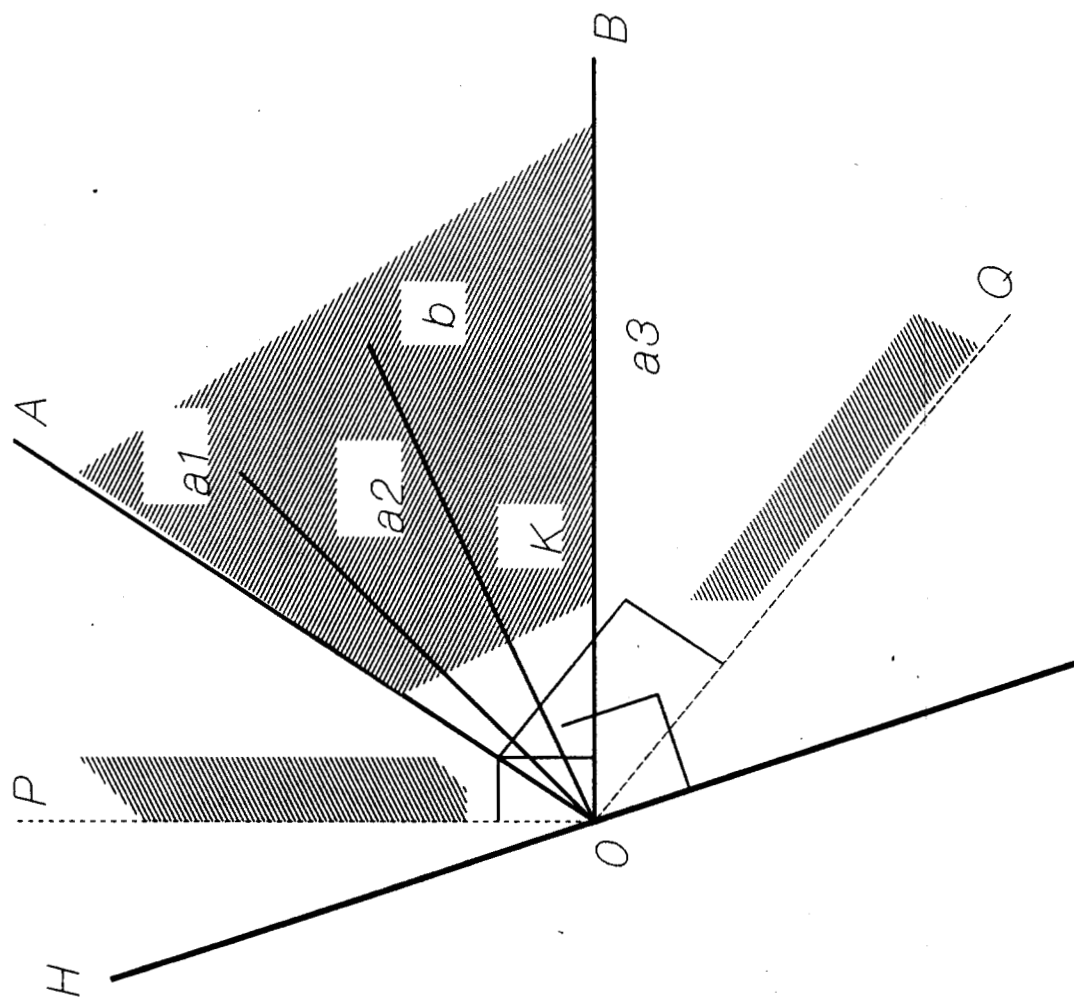
que x esta en el cono POQ (región sombreada).

(ii) La desigualdad $b \cdot x$ significa que x está en el semiespacio determinado por el hiperplano $H = \{x \in \mathbb{R}^n \mid b \cdot x = 0\}$, que contiene al punto b .

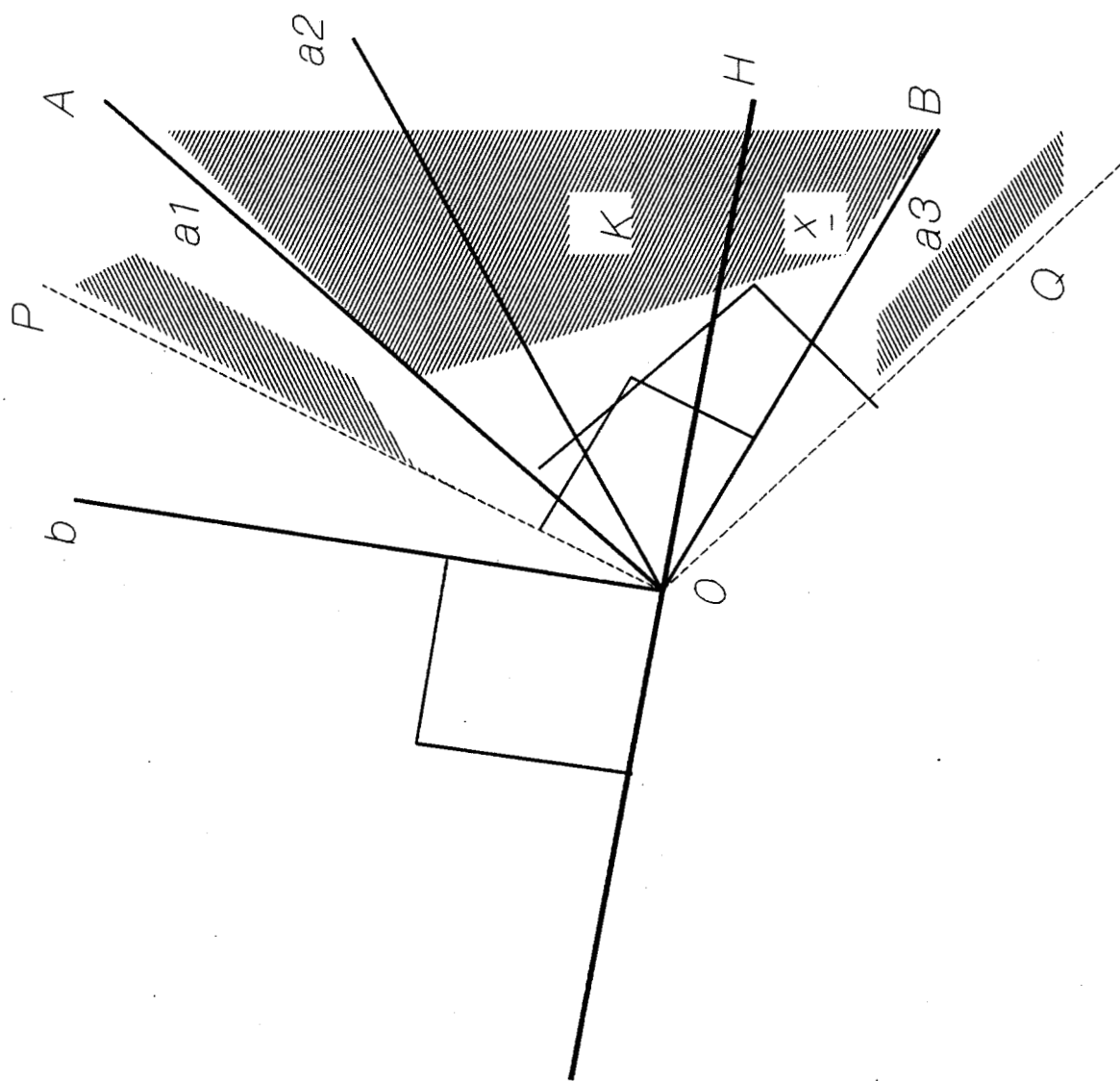
(iii) La conclusión del teorema es " $b \in K$ ", donde $K = \{y \mid y = \lambda A, \lambda \geq 0\}$.

(iv) Si $b \in K$ (caso a de la figura), entonces tenemos que $b \cdot x \geq 0$ para todo x tal que $a_i \cdot x \geq 0, i = 1, 2, \dots, m$ (es decir, el cono AOB está contenido en el semiespacio determinado por H que contiene a b).

(v) Si $b \notin K$ (caso b de la figura), entonces no podemos tener que $b \cdot x \geq 0, i = 1, 2, \dots, m$. En otros términos, existe un punto (digamos x') en el cono AOB , pero no en el semiespacio que contiene a b .



caso a: b pertenece a K



caso b: b no pertence a K

Como una última observación, señalaremos que el teorema de Minkowski - Farkas juega un papel muy importante en la teorías de programación lineal (teorema de dualidad), teoría de juegos (juego de dos personas, suma cero), y de programación no lineal (teorema de Kuhn - Tucker).

CAPITULO 3. APLICACIONES A LA TEORIA DE MERCADOS COMPETITIVOS.

En este capítulo estudiaremos la teoría correspondiente a mercados competitivos. Consideraremos una economía consistente en dos tipos de "agentes económicos", uno llamado *productores* y el otro *consumidores*. Existen m consumidores y k productores en la economía. Estos agentes económicos están relacionados con los *bienes*. Un paquete de bienes estará representado por un elemento de $\mathbb{R}^n$. En la teoría de mercados competitivos se supone que cada agente económico es tan pequeño en relación con el tamaño de la economía en su conjunto, que no puede afectar los niveles de precios prevalecientes en la misma (o bien, el impacto de sus acciones ya sean como consumidor o como productor, no se refleja sensiblemente en los precios de mercado). Esta suposición implica necesariamente que existe un gran número tanto de productores como de consumidores. También se supone la existencia de ciertas reglas de comportamiento que deben seguir los agentes. Se supone que cada consumidor maximiza su satisfacción dentro de el conjunto de planes de consumo que quedan a su alcance, de acuerdo con su presupuesto. Por su parte, los productores maximizan su beneficio eligiendo entre sus planes de producción posibles. El análisis típico de actividades se desarrolla en la mayoría de los casos mediante la caracterización de un productor cuyo conjunto de posibilidades de producción es un cono polihédrico convexo. En la descripción de la teoría correspondiente a la elección que hace el consumidor, asumiremos que sólo los planes de consumo deseados por éste intervienen en su toma de decisiones, y además que los precios de los bienes no afectan sus preferencias. En realidad, es posible y quizá frecuente que se desee consumir un bien sólo por

que es caro. Esto, que llamaremos "comportamiento esnobista", se descarta de la discusión¹. A diferencia de la teoría del productor, en la del consumidor no contamos con un criterio cuantificable que mida el comportamiento (como es el beneficio para el productor). Sin embargo, se discutirán importantes resultados relativos a la elección del consumidor y a la teoría de mercados competitivos sin referirnos a medidas de la satisfacción individual.

Adicionalmente al estudio del comportamiento de cada tipo de agente económico, la teoría de mercados competitivos introduce elementos que tienen que ver con la interacción de los agentes en una economía tendiente a alcanzar el *equilibrio competitivo*. Esto es, esencialmente, un estado en el que cada consumidor maximiza su satisfacción de acuerdo a sus restricciones presupuestales definidas por el vector de precios prevalecientes; por su parte, el productor maximiza su beneficio conforme a el mismo vector de precios dado, y la demanda total iguala a la oferta total de bienes.

La existencia de un equilibrio competitivo depende de la existencia de un vector de precios tal que permita el estado de la economía arriba descrito. Este vector de precios reflejará la consistencia entre las acciones de un gran número de consumidores y de productores, así como la compatibilidad derivada de la interacción mutua.

En el desarrollo de la teoría de mercados competitivos, se hace generalmente una suposición que juega un papel crucial para la extensión de resultados: la ausencia de factores externos. En otros términos, se supone que la interdependencia entre agentes económicos puede soslayarse. Hablaremos de independencia entre consumidores o productores, cuando cada uno de ellos se ocupe sólo de sus planes de consumo o de producción, haciendo caso omiso de los de otros agentes.

Para precisar la notación apuntaremos lo siguiente. El proceso de

¹Por el momento, su consideración complica innecesariamente el desarrollo teórico.

producción y_j para el j -ésimo productor es un vector n -dimensional cuyos elementos negativos denotan los insumos, mientras que los positivos se reservan para los productos. El conjunto de posibilidades de consumo para el i -ésimo consumidor se denota por X_i , y se llamará conjunto de consumo. Dado un vector de precios p y un ingreso M_i , el *conjunto presupuestal* es $\{x_i | x_i \in X_i, p \cdot x_i \leq M_i\}$. Cada individuo recibe un ingreso por la venta de sus recursos. Sus recursos están constituidos por sus bienes físicos y por los servicios que pueda prestar (como trabajo). Por el momento no se considera particularmente la circulación de dinero. Si existe, aceptar dinero significa exclusivamente aceptar un "pago". Podemos considerar simplemente que el dinero es inherente a un sistema institucional económico particular.

Las relaciones entre el presupuesto de un consumidor y su plan de consumo merecen algunos comentarios. Primero, la palabra *riqueza* se usará para referirse al valor total de las pertenencias materiales del consumidor (factibles de ser vendidas en el mercado), mas el valor total de los servicios que éste pueda prestar. Como se apuntó anteriormente, el valor que representa un bien está fechado y, en consecuencia, la riqueza de un consumidor. Otra implicación de fechar los bienes es que se supone que cada consumidor conoce todas las posibilidades de consumo que encontrará en el futuro y los precios que prevalecerán; en otros términos, el consumidor puede predecir situaciones. Esta suposición puede justificarse claramente para el caso de estados estacionarios de la economía; sin embargo, resulta de escasos sustentos en otras condiciones cualesquiera. Las extensiones de la teoría para economías de incertidumbre se han comenzado a estudiar recientemente.

Por el momento, regresemos al vector de consumo, y llamémosle x_i . Si el i -ésimo consumidor tiene x_i recursos y los pone a la venta, su ingreso será $p \cdot x_i$ en donde p es el vector de precios prevaleciente en el mercado. Su restricción presupuestal es, en consecuencia

$$p \cdot x_1 \leq p \cdot x_1$$

Advertimos ahora la posibilidad de que el consumidor i reserve una parte de sus recursos para destinarla a satisfacer sus necesidades elementales. Esto puede pensarse como una venta total de sus posesiones seguida de una compra de parte de lo que vendió, que considera indispensable (el efecto es el mismo que el de vender sólo parte de sus bienes). Para clarificar esto a través de un ejemplo, supóngase que el consumidor i tiene sólo del bien A . Suponga además que su dotación inicial de este bien es X_a y que de ahí vende una parte, digamos X'_a . Su consumo de A será entonces $(X_a - X'_a)$. Pensemos ahora que quiere consumir del bien B , del cual nunca tuvo dotación inicial. Si p_a y p_b son los precios respectivos, su restricción presupuestal se escribe como

$$p_a(X_a - X'_a) + p_b X_b \leq p_a X_a$$

o bien como

$$p_b X_b \leq p_a X'_a$$

Y mientras los montos totales de recursos estén establecidos, no habrá diferencia alguna derivada de cómo se escriba la restricción presupuestal. De hecho, puede escribirse nuevamente como

$$p_a(-X'_a) + p_b X_b \leq 0$$

Así, podemos considerar su plan de consumo como $(-X'_a, X_b)$. Sus preferencias estarán definidas entonces dentro de los valores posibles de $(-X'_a, X_b)$, en lugar de $(X_a - X'_a, X_b)$. Obsérvese que estos supuestos nos han llevado a un vector con componentes negativas.

En general, podemos escribir $p \cdot x_1 \leq p \cdot x_1$ como $p \cdot z_1 \leq 0$, en donde $z_1 = x_1 - x_1$, siendo x_1 el vector que corresponde al plan de consumo y x_1 el vector que corresponde a las dotaciones iniciales

de bienes. Los elementos negativos de z_i representarán entonces cantidades entregadas, mientras que los positivos representarán cantidades recibidas.

Pueden surgir algunas complicaciones al considerar ahora a los productores. Primero, hay que tener en cuenta que los productores pueden reservarse para sí mismos cierta cantidad de recursos. La pregunta sería: Qué sucede con los ingresos que esos recursos representan? Una manera sencilla (aunque no única) de evitar esta cuestión se encuentra al suponer que todos los recursos pertenecen inicialmente a los consumidores, que en un momento dado, deciden vender parte a los productores. En consecuencia, los consumidores reciben un ingreso. En segundo lugar, cabe preguntar quién detendrá los beneficios que obtengan los productores. Esto se resuelve suponiendo que los consumidores son los dueños de las empresas. Sea $y_j \in \mathbb{R}^n$ el plan de producción (en donde se combinan insumos y productos) elegido por el j -ésimo productor para un vector de precios p dado. Los elementos negativos de y_j denotan insumos y los elementos positivos denotan productos. El beneficio para el productor es entonces $p \cdot y_j$. Sea θ_{ji} la fracción de la empresa j que pertenece al consumidor i . Advertimos que

$$\sum_{i=1}^m \theta_{ji} = 1 \text{ para todo } j \quad \text{y} \quad \theta_{ji} \geq 0 \text{ para todo } j \text{ y para todo } i$$

Entonces, el monto de los dividendos que obtiene el consumidor i del productor j es $\theta_{ji}(p \cdot y_j)$. Conservando x_i para denotar las dotaciones iniciales del consumidor i , su riqueza total (previa a cualquier consumo) puede escribirse ahora como

$$p \cdot x_i + \sum_{j=1}^k \theta_{ji}(p \cdot y_j)$$

Nótese que esta formulación no presupone la posibilidad de que un individuo juegue el papel de productor y de consumidor simultáneamente.

Al conjunto de posibles combinaciones de insumo-producto para el

j -ésimo productor se le llamará *conjunto de producción*, y se denotará por Y_j . En adelante será claro que $y_j \in Y_j$, y asumiremos que Y_j es un subconjunto de $\mathbb{R}^n$. Ante la ausencia de factores externos, tendremos que el conjunto de producción agregada de la economía, denotado por Y , se puede definir como

$$Y = \sum_{j=1}^k Y_j$$

Finalmente, hay que hacer notar que la definición de equilibrio competitivo (que hasta el momento no se ha precisado formalmente) es independiente de quién detente los recursos, ya que se puede describir el comportamiento de los consumidores sin especificar a cuánto asciende su ingreso total; simplemente se indica un $x_i \in X_i$ que el consumidor elige entre todos los x_i iguales o menores en valor.

Hechos estos comentarios, podemos entrar de lleno a la teoría de mercados competitivos. Debido a la simplicidad con que se harán los planteamientos, se podrá observar con cierta claridad la cantidad de problemas que surgen del proceso de abstracción. Estudios de este tipo son de fundamental importancia para el entendimiento de versiones más elaboradas (o mas "realistas") que, hasta el momento, no se encuentran completamente desarrolladas en la literatura existente.

El concepto central en el análisis neoclásico de la producción es el de *función de producción*. Una función de producción es aquella que describe la relación tecnológica entre insumos y productos. Si denotamos el vector de productos por $x = (x_1, x_2, \dots, x_k)$ y el vector de insumos por $v = (v_1, v_2, \dots, v_r)$, entonces la función de producción tiene la forma $F(x, v) = 0$. Esta relación define usualmente una superficie única en el plano x para un valor dado de v . Consideremos el caso en donde no existe producción conjunta, de modo que $x = x_1$. En este caso, la función de producción se escribiría como $x = f(v)$. Es común suponer que la función es inyectiva y diferenciable.

El *análisis de actividades* ha revolucionado el análisis tradicional de producción al descartar el concepto de función de producción. En su lugar, postula la existencia de un conjunto de procesos de producción disponibles en una economía dada.

A cada elemento del conjunto de producción se le conoce como *proceso* o *actividad*. Se supone que hay n bienes en la economía y que cada uno de ellos es cualitativamente homogéneo.

La convención de fechar los bienes, extiende el enfoque del análisis de actividades hacia procesos dinámicos y teoría del capital, en donde se involucra al tiempo de manera fundamental.

La forma de operar en el análisis de actividades se basa en la teoría de conjuntos y algunas otras ramas de las matemáticas. El análisis de actividades es axiomático, más riguroso y con fundamentos más formales que el tradicional análisis neoclásico. Los teoremas de separación juegan un papel esencial.

Veamos ahora algunos elementos que intervendrán de manera importante en la discusión subsecuente.

Sea Y el conjunto de todos los procesos de producción tecnológicamente posibles en una economía dada. Hemos dicho que $Y \in \mathbb{R}^n$, y que $y \in Y$ denota un proceso de producción. Introduzcamos las siguientes hipótesis:

(a-1) (Aditividad) $y, y' \in Y$ implica que $y + y' \in Y$.

(a-2) (Proporcionalidad) $y \in Y$ implica que $\alpha y \in Y$ para todo $\alpha \geq 0$, $\alpha \in \mathbb{R}$.

Si (a-1) y (a-2) se cumplen, Y es un cono convexo. Debido a la proporcionalidad, si $a_j \in Y$, entonces

$$\lambda_j a_j \in Y \text{ para toda } \lambda_j \geq 0, \text{ con } \lambda_j \in \mathbb{R} \text{ donde } a_j = \begin{bmatrix} a_{1j} \\ \vdots \\ a_{nj} \end{bmatrix}$$

El vector a_j se refiere entonces a la j -ésima actividad o proceso de Y . Por su parte, a_{ij} denota la cantidad del bien i necesaria para la producción de una unidad del bien j ; λ_j representa el nivel de actividad para j . El tercer supuesto es el siguiente:

(a-3) Existe un número finito de actividades básicas (a_j) de manera que Y es un cono polihédrico convexo (generado por las a_j 's)

En otros términos, un elemento típico y en Y puede ser expresado por una combinación lineal de $a_1, a_2, \dots, a_m$, donde m es un entero positivo. Como consecuencia de los supuestos anteriores, el conjunto de producción Y en el análisis de actividades puede ser escrito como $Y = \{y \mid y = A \cdot \lambda, \lambda \geq 0\}$, donde A es una matriz $n \times m$ (con entradas reales) formada por $[a_1, a_2, \dots, a_m]$ y λ es un vector m -dimensional cuyo j -ésimo elemento es λ_j .

Debe apreciarse que el supuesto de aditividad significa la acción independiente de cada actividad.

El hecho de que Y sea un cono polihédrico convexo implica algunas otras propiedades.

(i) $0 \in Y$ (posibilidad de inacción). Esto es, es posible para un productor permanecer inactivo.

(ii) Y es un conjunto cerrado (Lo cual resulta importante en el sentido matemático, y en el sentido económico significa que cualquier proceso de producción que pueda ser aproximado por una sucesión de procesos en Y , se encuentra también en Y)

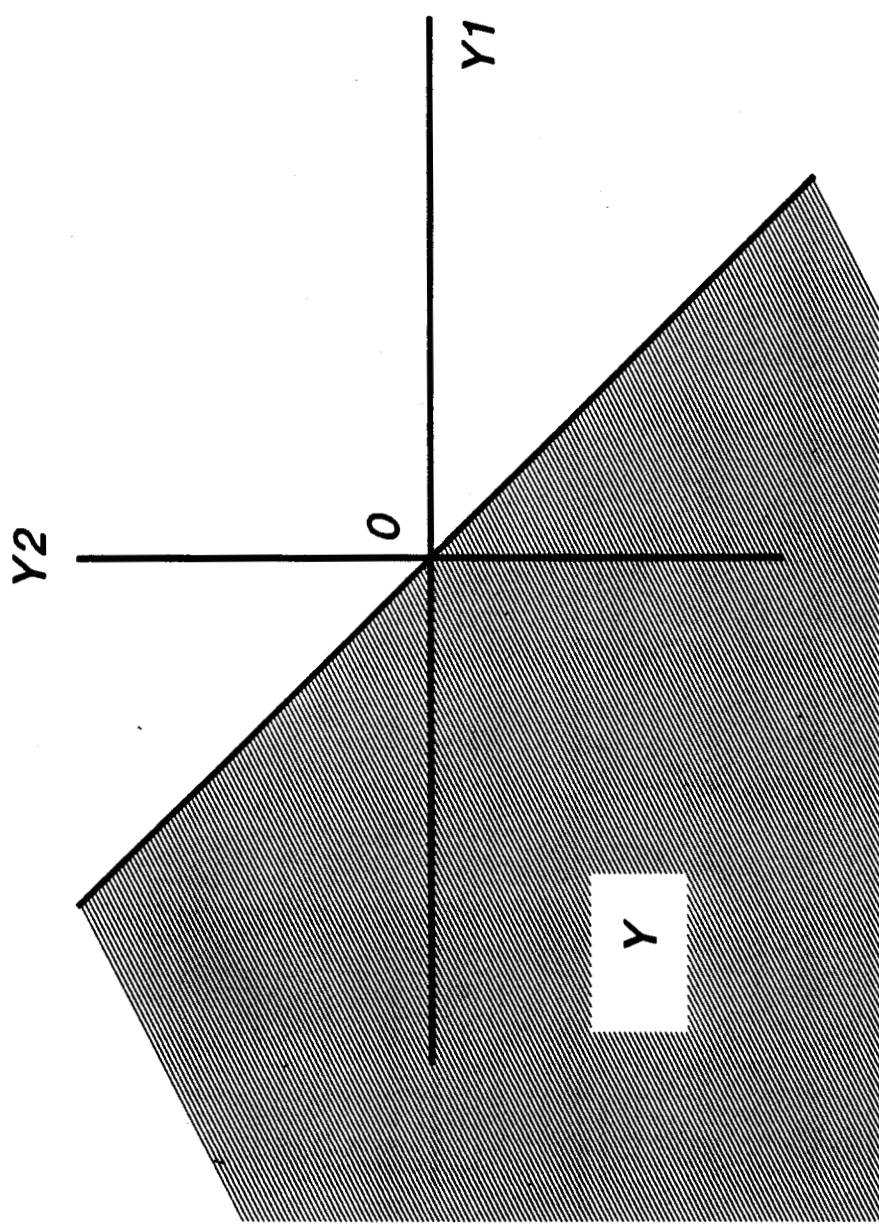
Koopmans sugiere los siguientes tres supuestos de importante significado económico:

(a-4) (Productividad) Existe al menos un elemento positivo para algún y en Y .

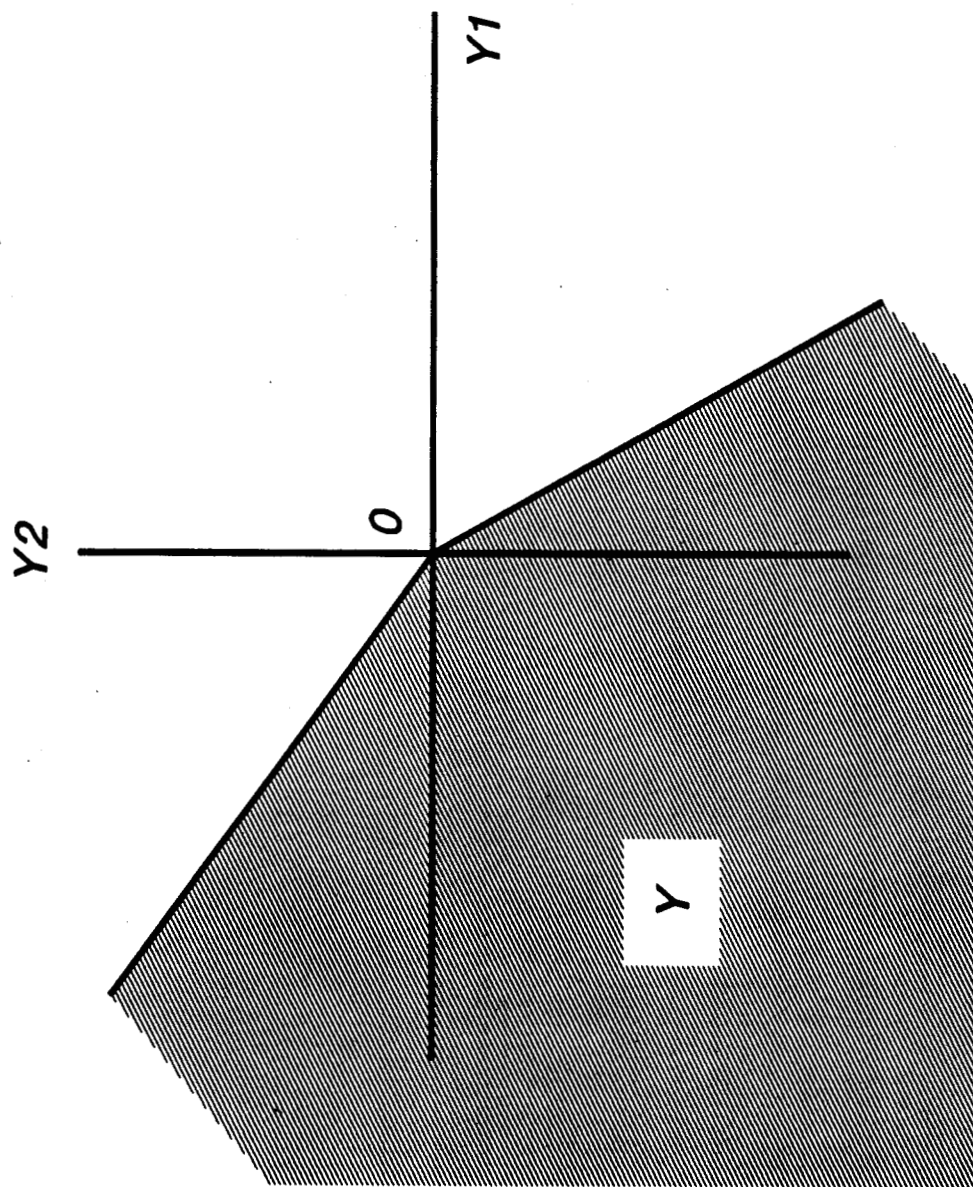
(a-5) (Imposibilidad de Producción libre) $y \geq 0$ implica que $y \notin Y$, o bien $Y \cap \Omega = \{0\}$.

(a-6) (Irreversibilidad) $y \in Y$, $y \neq 0$ implica que $-y \notin Y$. O bien $Y \cap (-Y) = \{0\}$.

Los siguientes diagramas ilustran el significado de algunos de los postulados anteriores. En el caso a, (a-4) y (a-5) se cumplen pero no (a-6). En el caso b, (a-4), (a-5) y (a-6) se cumplen.



caso a



caso b

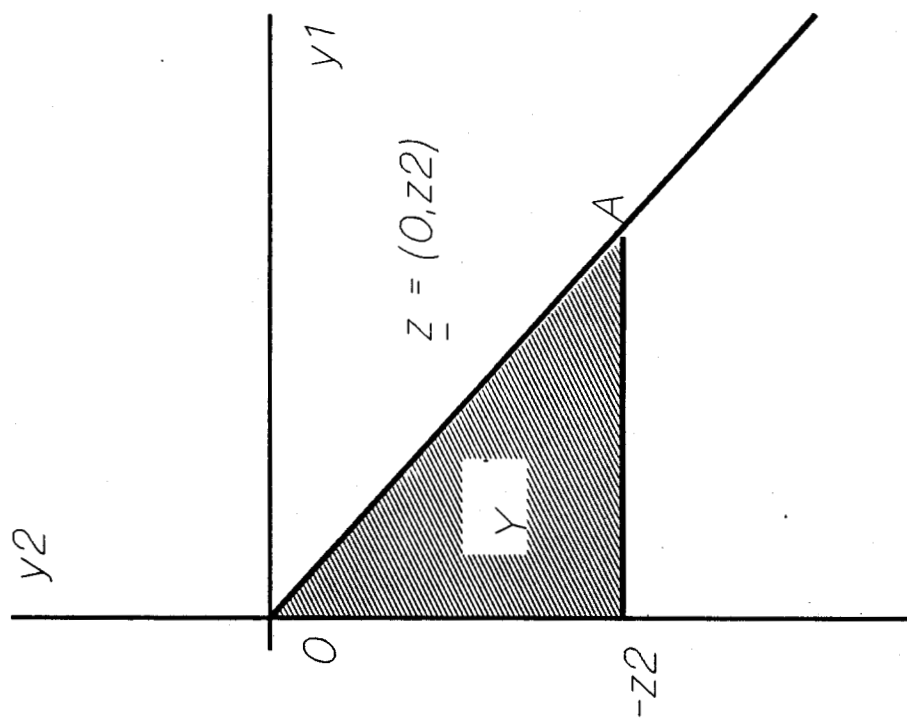
Teorema 1-3: Sea $Y = \{y \mid y = A \cdot \lambda, \lambda \geq 0\}$. Y satisface el postulado (a-5) si y sólo si existe un $p > 0$, $p \in \mathbb{R}^n$, $|p| < \infty$, tal que $p \cdot y \leq 0$ para todo $y \in Y$.

Si interpretamos a p como un vector de precios, entonces $p \cdot y$ representa el beneficio derivado de y . Luego, por ejemplo, $p \cdot y \leq 0$ para todo $y \in Y$ significa que el beneficio máximo es a lo más 0. Aquí establecimos los supuestos (a-4), (a-5) y (a-6) junto con el hecho de que el conjunto Y es un cono polihédrico convexo. En general estos postulados pueden ser establecidos aun cuando Y no sea un cono polihédrico convexo. Cuando Y representa la colección de combinaciones insumo-producto técnicamente factibles en una economía, y cuando no se requiere que Y sea un cono polihédrico convexo, llamamos a Y *conjunto de producción general*. Podemos enlistar algunas de las hipótesis más importantes que desearíamos se cumplieran para el conjunto de producción general.

- (i) El conjunto Y es cerrado
- (ii) Posibilidad de inacción ($0 \in Y$)
- (iii) Productividad
- (iv) Producción libre
- (v) Irreversibilidad ($Y \cap (-Y) \subset \{0\}$)
- (vi) Eliminación Libre [$y \in -\Omega$ implica que $y \in Y$, o bien $-\Omega \subset Y$]
- (vii-a) El conjunto Y es un cono polihédrico convexo
- (vii-b) El conjunto Y es un cono convexo
- (vii-c) El conjunto Y es convexo

Obsérvese que (vii-a) implica (i) e (ii) y que (vii-b) implica (ii). Los enunciados (vii-c) e (ii) juntos implican que si $y \in Y$, entonces $\alpha y \in Y$ para todo $0 \leq \alpha \leq 1$; en otras palabras, prevalecen rendimientos no decrecientes a escala. Nótese también que la convexidad de Y presupone la divisibilidad perfecta de todos los bienes.

El conjunto de producción Y como se acaba de describir indica las posibilidades tecnológicas en una economía dada. Se aprecia que $y \in Y$ indica cuánto producto puede ser obtenido después de especificar las cantidades de insumos, pero nunca se especifica en dónde o qué cantidad de estos insumos se encuentra disponible en la economía. Sin embargo, se puede introducir la existencia de limitantes en cuanto a la disponibilidad de recursos. Por ejemplo, Y puede ser un cono polihédrico convexo truncado, es decir, $Y = \{y \mid y = A \cdot \lambda, \lambda \geq 0, \lambda \in \mathbb{R}^m, y + z \geq 0\}$, donde $z \geq 0$ denota la limitante de recursos en la economía. Con el postulado de producción libre, el conjunto deja de ser un cono polihédrico convexo, pero sigue siendo un conjunto convexo. Nótese que ahora el conjunto es compacto. En la siguiente figura se ilustra el truncamiento.



Estrictamente hablando, el término *análisis de actividades* no comprende solamente el estudio del proceso de producción cuando el número de actividades básicas es finito. Por su parte, el conjunto Y puede ser un cono polihédrico convexo, o un cono polihédrico convexo truncado. Esto no modifica el sustancialmente el enfoque que daremos a los resultados subsecuentes. Precisaremos ahora uno de los conceptos fundamentales

Definición: Sea $Y \subset \mathbb{R}^n$ un conjunto de producción general. Un punto $\underline{y}$ de Y es un *punto eficiente* si no existe ningún otro $y \in Y$ tal que $y \geq \underline{y}$.

Un punto eficiente representa un punto frontera, y es una combinación de insumos y productos tal, que ningún producto puede incrementarse sin que se disminuyan otros, o se aumenten insumos. En términos del último diagrama, OA es el conjunto de puntos eficientes en el cono truncado de producción. Si el conjunto de producción es todo el cono, la semirrecta que parte de O y pasa por A es el conjunto de puntos eficientes.

Un punto eficiente se define a menudo con la indicación explícita de las restricciones de recursos. Sin embargo, nuestra interpretación de Y es flexible, en el sentido de que considera el caso en donde no se toman en cuenta restricciones de recursos. En consecuencia, la definición de punto eficiente aquí presentada es también flexible.

Estableceremos y demostraremos ahora dos teoremas básicos para el análisis de actividades.

Teorema 2-3: Sea Y un conjunto de producción general en $\mathbb{R}^n$. Un punto $\underline{y}$ en Y es un punto eficiente si existe un $p > 0$, $p \in \mathbb{R}^n$, con $|p| < \infty$ tal que $p \cdot \underline{y} \geq p \cdot y$ para todo $y \in Y$.

Prueba: Supongámonos lo contrario. Entonces existe $y \in Y$ tal que $y \geq \underline{y}$. Se sigue que $p \cdot y > p \cdot \underline{y}$, ya que $p > 0$, lo que representa una

contradicción.

Nota: En este teorema no se especifica nada acerca de las hipótesis sobre Y .

Para probar el siguiente teorema (3-0), necesitaremos primero un resultado.

Lema: Sea Y un conjunto de producción general en $\mathbb{R}^n$. Si $\underline{y}$ es un punto eficiente de Y , entonces $(Y - \underline{y}) \cap \Omega_0 = \emptyset$, donde Ω_0 es el ortante positivo de $\mathbb{R}^n$ (es decir, el interior del ortante no negativo de $\mathbb{R}^n$).

Prueba: Supóngase lo contrario. Entonces existe un $z > 0$ (esto es, $z \in \Omega_0$) tal que $z \in (Y - \underline{y})$. Entonces $z + \underline{y} \in Y$. Se sigue que $\underline{y}$ no es un punto eficiente de Y , lo que constituye una contradicción.

Teorema 3-3: Sea Y un conjunto de producción convexo en $\mathbb{R}^n$. Si $\underline{y}$ es un punto eficiente de Y , entonces existe $p \geq 0$, $p \in \mathbb{R}^n$, $|p| < \infty$ tal que $p \cdot \underline{y} \geq p \cdot y$ para todo $y \in Y$.

Prueba: Sea $X = Y - \underline{y}$. Entonces X es convexo (como suma de dos conjuntos convexos). Ya que $\underline{y}$ es un punto eficiente de Y , tenemos que $X \cap \Omega_0 = \emptyset$ por el lema anterior. Tenemos entonces dos conjuntos convexos ajenos y no vacíos. De el teorema de separación de Minkowski se sigue que existen un $p \in \mathbb{R}^n$, $p \neq 0$, $|p| < \infty$, y un $\alpha \in \mathbb{R}$ tal que

$$p \cdot z \geq \alpha \quad \text{para todo } z \in \Omega_0$$

y

$$p \cdot x \leq \alpha \quad \text{para todo } x \in X$$

Nótese que $0 \in X$ para $\underline{y} \in Y$. Luego $\alpha \geq 0$. Tenemos entonces que $p \geq$

0. Si esto no es cierto, existe un elemento de p (digamos p_1) que es negativo (dado que $p \neq 0$). Eligiendo el elemento correspondiente en z (digamos z_1) lo suficientemente grande, podemos obtener $p \cdot z < 0$, lo que constituye una contradicción. También $\alpha \leq 0$, ya que de lo contrario, $p \cdot z \geq \alpha > 0$ para todo $z \in \Omega_0$. Esto es imposible, ya que eligiendo $|z|$ suficientemente pequeña podemos tener $p \cdot z < \alpha$. Ya que $\alpha \geq 0$ y $\alpha \leq 0$ tiene que ser $\alpha = 0$. Por lo tanto $p \cdot x \leq 0$ para todo $x \in X$ con $p \geq 0$, o $p \cdot (y - \underline{y}) \leq 0$; esto es, $p \cdot \underline{y} \geq p \cdot y$ para toda $y \in Y$ con $p \geq 0$. Q.E.D.

Conforme al teorema 3-3, el concepto de punto eficiente puede ser ahora caracterizado por la maximización del beneficio; es decir, maximización de $p \cdot y$ con respecto a $y \in Y$. La existencia de una solución para este problema de maximización la garantiza el hecho de que Y sea un conjunto compacto (teorema de Weierstrass), ya que el producto punto es una función continua.

Supóngase que Y puede ser caracterizado por las desigualdades

$$\sum_{j=1}^m a_{ij} \lambda_j \leq r_i, \quad i = 1, 2, \dots, n, \quad \lambda_j \geq 0, \quad j = 1, 2, \dots, m$$

Entonces el problema de encontrar una $\lambda = (\lambda_1, \lambda_2, \dots, \lambda_m)$ que maximiza $p \cdot y$ donde $y = \sum_{j=1}^m a_{ij} \lambda_j$, sujeto a las restricciones que aparecen arriba, es un problema típico de programación lineal, en donde existen métodos computacionales bien conocidos y usados en la práctica. En consecuencia, puede verse que el análisis de actividades tiene también un importante significado en cuanto a aplicaciones.

CAPITULO 4. APLICACIONES A LA TEORIA DE CONSUMO.

El concepto básico en la teoría de preferencias es el *conjunto de consumo*, que es análogo al conjunto de producción. Es necesario anotar los siguientes supuestos sobre este conjunto de consumo:

- (a-1) X es convexo.
- (a-2) Un individuo puede consumir cualquier cantidad de un bien.
- (a-3) Un individuo sobrevive mientras tenga cantidades positivas de al menos un bien.
- (a-4) X es un subconjunto de un espacio vectorial de dimensión finita.

El supuesto (a-2) no debe considerarse como demasiado restrictivo, ya que se entiende que el consumo representa una actividad en la que el bien pasa a manos de un individuo sin implicar que éste deba disponer de el bien al instante. Asumir (a-1) resulta algo muy conveniente para nuestros fines. Sin embargo, es una hipótesis relativamente fuerte; nos lleva, entre otras cosas, a considerar la *divisibilidad perfecta* de los bienes (es decir, que a sabiendas de que existen algunos para los cuales no tiene sentido hablar de cantidades fraccionarias, se considera que es posible consumir cualquier cantidad de cualquier bien). En este contexto, se podría cuestionar la necesidad de que X sea un subconjunto de $\mathbb{R}^n$. Siendo subconjunto de un espacio lineal, la multiplicación por escalar es permisible. Aun asumiendo divisibilidad perfecta, X podría no ser necesariamente un subconjunto de un espacio vectorial de dimensión finita. Considérese como ejemplo, el caso en que, al fechar los bienes a tiempo continuo, el vector de consumo $x(t) \in X$ nos lleva

a un subconjunto de un espacio vectorial de dimensión infinita.

Sea X el conjunto de consumo del i -ésimo individuo. Dados dos elementos $x, y \in X$, supongamos que el consumidor puede decidir, basado en el nivel de satisfacción que represente, (a) preferir x , (b) considerar x al menos tan satisfactorio como y , o bien (c) considerar ambos planes de consumo como equivalentes. Esto induce cierto tipo de relación (que llamaremos *relación de preferencia*) en los elementos de X . Podemos considerar esta relación como un conjunto de pares (x, y) con $x, y \in X$. En la notación, podemos expresar x es preferido a y como $x (>) y$, o bien xRy . En genreal, una relación (binaria) involucra siempre a un conjunto de pares ordenados.

Definición: Una relación R en X es llamada *preorden parcial* o *quasi-orden parcial* si satisface:

- i) xRx para todo $x \in X$ (reflexividad)
- ii) xRy y yRz implica que xRz cuando $x, y, z \in X$ (transitividad).

Adicionalmente, si xRy y yRx implica que $x = y$, entonces la relación es llamada *orden parcial* (o simplemente orden). Más aun, si en una relación tenemos necesariamente que xRy o que yRx para elementos arbitrarios de X ($x \neq y$), hablaremos de un *quasi-orden total* o *completo*. Un orden total o completo se puede definir análogamente.

Definición: Una *relación de equivalencia* se tiene cuando se cumplen la reflexividad, transitividad y

- iii) xRx implica yRx (simetría).

En otras palabras, una relación de equivalencia es un *quasi-orden simétrico* (denotado por $x \sim y$). La relación de preferencia es *indiferente* a es un ejemplo de relación de equivalencia. Si $X = \mathbb{R}^n$

y definimos para $x, y \in X$ $x \geq y$ si $x_i \geq y_i$ para todo i , la relación $\geq$ es un ejemplo de orden parcial.

Usaremos la notación siguiente para los órdenes de preferencias:

- (a) x es preferido a y : $x(>)y$
- (b) x no es preferido a y : $x(\leq)y$
- (c) x es indiferente a y : $x(=)y$.

En la teoría del consumidor, se supone que el orden de preferencias está definido en el conjunto de consumo de cada individuo. Un agente no examina ni compara los planes posibles para otros consumidores al tomar sus decisiones.

Definición: Sea X un conjunto de consumo y sea $\alpha \in X$. Entonces

- i) El conjunto $\{x | x \in X, x(\geq)\alpha\}$ se llamará *contorno superior de α*
- ii) El conjunto $\{x | x \in X, x(\leq)\alpha\}$ se llamará *contorno inferior de α*
- iii) El conjunto $\{x | x \in X, x(>)\alpha\}$ se llamará *conjunto preferido a α* .
- iv) El conjunto $\{x | x \in X, x(<)\alpha\}$ se llamará *conjunto no preferido a α* .
- v) El conjunto $\{x | x \in X, x(=)\alpha\}$ se llamará *conjunto indiferente a α* .

Sea X el conjunto de consumo de un individuo en particular, y supóngase que su satisfacción derivada de un plan determinado de consumo, puede expresarse por un índice real positivo (que llamaremos índice de utilidad). Este índice representa una función de X a $\mathbb{R}$, que denotaremos por $u(x)$ y llamaremos *función de utilidad*. Dado que los reales poseen un orden natural, se puede sugerir inmediatamente que este orden represente al orden de las preferencias. Es decir,

- i) $x(\geq)y$ si y sólo si $u(x) \geq u(y)$

- ii) $x(>)y$ si y sólo si $u(x) > u(y)$
- iii) $x(=)y$ si y sólo si $u(x) = u(y)$

Una característica fundamental de la función de utilidad, es que puede ser sustituida por cualquier función monótona creciente de $u(y)$ sin alterar el orden. En otras palabras, si ϕ es una función real tal que

$$\phi(u^1) \begin{matrix} \geq \\ \leq \end{matrix} \phi(u^2) \text{ según sea } u^1 \begin{matrix} \geq \\ \leq \end{matrix} u^2$$

entonces

- i) $x(\geq)y$ si y sólo si $\phi[u(x)] \geq \phi[u(y)]$
- ii) $x(>)y$ si y sólo si $\phi[u(x)] > \phi[u(y)]$
- iii) $x(=)y$ si y sólo si $\phi[u(x)] = \phi[u(y)]$

En el análisis gráfico, el concepto de *curvas de indiferencia* se introduce frecuentemente. Una curva de indiferencia es el lugar geométrico de $x = (x_1, x_2)$ de $\mathbb{R}^2$ tal que $u(x) = \text{constante}$. Generalmente se dibujan las curvas de indiferencia convexas con respecto al origen, para reflejar una tasa marginal de sustitución entre dos bienes. También es común suponer que $u(x)$ es estrictamente creciente (aunque podría suponerse lo contrario sin pérdida de generalidad).

Una cuestión importante referente a la existencia de curvas de indiferencia constituidas por "bandas", surge al preguntarse por la posibilidad de representar el orden de preferencias de un consumidor a través de una función de utilidad (es decir, a través de números reales). Para tratar ese aspecto, examinemos la siguiente

Definición: Un orden de preferencias ($\geq$) en X se llamará *continuo* o *cerrado* si para todo $\alpha \in X$, los conjuntos $\{x \mid x \in X, x (\geq) \alpha\}$ y $\{x \mid x \in X, x (\leq) \alpha\}$ son cerrados. En otros términos, si $\{x_q\}$ es una sucesión en X con $x_q (\geq) \alpha$ para todo q (respectivamente $x_q (\leq)$

α para todo q), entonces $x_q \rightarrow x_0$ implica que $x_0 (\geq) \alpha$ (respectivamente $x_0 (\leq) \alpha$).

Teorema 1 - 4: Sea X un subconjunto conexo de $\mathbb{R}^n$ y sea $(\geq)$ un preorden total continuo definido en X . Entonces existe una función de utilidad continua en X para $(\geq)$.

Prueba: El caso en el que todos los puntos de X son indiferentes puede ser resuelto inmediatamente. Cualquier función de valor real constante es una función de utilidad continua. Se excluirá este caso.

La demostración se basa en la existencia de un subconjunto numerable D de X que es denso en X . En la parte 2 de la demostración se define en D una función de valor real creciente convenientemente escogida. En la parte 3 esta función se extiende de D a X . En la parte 4 se demuestra que la función así definida es continua. La parte 1 nos proporciona un útil resultado preliminar.

1. Resultado preliminar.

(1) Si x' y x'' satisfacen $x' (<) x''$, existe un x en D tal que $x' (<) x (<) x''$.

Para demostrarlo considérense los dos conjuntos

$$X_{x'} = \{x \in X \mid x (\leq) x'\} \quad \text{y} \quad X^{x''} = \{x \in X \mid x'' (\leq) x\}$$

Son disjuntos, no vacíos y cerrados por hipótesis de continuidad. Puesto que X es conexo no puede ser su unión, de modo que:

$$X_{x'} \cup X^{x''} \neq X.$$

Supóngase ahora que no hubiera x alguno en D con la propiedad deseada; esto significaría que $D \subset X_{x'} \cup X^{x''}$. En virtud de que $X \subset Y$ implica que $\bar{X} \subset \bar{Y}$, $\bar{D}$, la adherencia de D en X , estaría contenida en la adherencia en X del conjunto $X_{x'} \cup X^{x''}$. Sin

embargo, este último es cerrado puesto que es la unión de dos conjuntos cerrados. Por lo tanto se tendría que $\bar{D} \subset X_x \cup X^{x''}$, o, puesto que $\bar{D} = X$,

$$X = X_x \cup X^{x''}.$$

Se obtendría, por consiguiente, una contradicción.

2. Una función de utilidad en D .

La función de utilidad a ser definida en D se designará por u' .

Elíjanse dos números reales a, b tales que $a < b$.

Si D tiene un elemento mínimo $x_{\min}$, tómese $u'(x_{\min}) = a$.

Si D tiene un elemento máximo $x_{\max}$, tómese $u'(x_{\max}) = b$.

Extraíganse de D todos los elementos indiferentes a $x_{\max}$ y a $x_{\min}$ y llámese D' al conjunto restante. Por (1),

(2) D' no tiene elemento mínimo ni máximo.

Defínase una función creciente de D' sobre el conjunto Q' de números racionales en el intervalo (a, b) como sigue. Puesto que D' es numerable sus elementos pueden ser ordenados $(x_1, x_2, \dots, x_p, \dots)$; este ordenamiento no está relacionado con el preorden $(\leq)$. Similarmente, Q' es numerable y sus elementos pueden ser ordenados $(r_1, r_2, \dots, r_q, \dots)$; este ordenamiento no está relacionado con el orden $\leq$. Los elementos de D' serán considerados sucesivamente; con x_q se asociará un elemento r_{qp} de Q' de tal forma que el preorden se conserve y que todo elemento de Q' acabe por ser tomado.

Considérese x_1 ; tómese $q_1 = 1$ y $u'(x_1) = r_{q1}$.

Considérese x_2 ; efectúese la partición de D' en los siguientes conjuntos: la clase de indiferencia de x_1 , los intervalos $\{x \in D' \mid x (\leq) x_1\}$ y $\{x \in D' \mid x (\geq) x_1\}$. Pueden ocurrir dos casos:

si x_2 es indiferente a x_1 (se escribe $x_2 \sim x_1$), tómese $q_2 = q_1$ y $u'(x_2) = r_{q2}$;

si x_2 está en uno de los dos intervalos, digamos $\{x \in D' \mid x (\leq) x_1\}$ (que escribiremos como $] \leftarrow, x_1[$), considérese el intervalo correspondiente $]a, r_1[$, de Q' y selecciónese en él el número

racional de menor rango r_{q2} ; tómesese $u'(x_2) = r_{q2}$.

En general considérese x_p ; efectúese la partición de D' en los siguientes conjuntos: las clases de indiferencia de $x_1, x_2, \dots, x_{p-1}$ (el número de conjuntos así obtenidos es a lo sumo $p-1$), los intervalos de la forma $]x_{p1}, x_{p1}[$, o $]x_{pm}, x_{pm+1}[$, o $]x_{p(p-1)}, x_{p(p-1)}[$ donde $m < n$ implica que $x_{pm} (\leq) x_{pn}$ (el número de intervalos no vacíos es a lo sumo p). Pueden ocurrir dos casos:

si $x_p \sim x_{p'}$ donde $p' < p$, tómesese $q_p = q_{p'}$ y $u(x_p) = r_{q_p}$;

si x_p está en uno de los intervalos, digamos $]x_{p'}, x_{p''}[$, considérese el intervalo correspondiente $]r_{q_{p'}}, r_{q_{p''}}[$ de Q' y selecciónese en él el número racional de menor rango r_{q_p} ; tómesese $u'(x_p) = r_{q_p}$.

La función u' es creciente. (1) y (2), con el procedimiento de selección del número racional de menor rango, implican que todo elemento de Q' acaba por ser tomado.

3. Extensión de D a X .

La función de utilidad a ser definida en X se designará por u . Si x' es un elemento de X , se escribirá $D_{x'} = \{x \in D \mid x (\leq) x'\}$ y $D^{x'} = \{x \in D \mid x' (\leq) x\}$.

Si x es un elemento mínimo de X , tómesese $u(x) = a$.

Si x es un elemento máximo de X , tómesese $u(x) = b$.

Se demostrará que estos dos números son iguales.

(1) Si x' es un elemento de D_x , y x'' un elemento de $D^{x'}$, se tiene que $x' (\leq) x''$. Por lo tanto, si r' es un elemento de $u'(D_x)$ y r'' un elemento de $u'(D^{x'})$ se tiene que $r' \leq r''$. De esto se deduce que $\sup u'(D_x) \leq \inf u'(D^{x'})$.

(2) $\sup u'(D_x) < \inf u'(D^{x'})$ no puede ocurrir puesto que en ese caso cualquier número racional entre ellos no sería un valor tomado por u' .

Tómesese para $u(x)$ es valor común del $\sup$ y del $\inf$.

Si $x \in D$ se tiene $u(x) = u'(x)$ y u es efectivamente una extensión de u' de D a X ; en particular

$Q' \subset u(x) \subset [a, b]$.

Se puede ver entonces que u es creciente.

4. Continuidad de u .

Se demostrará que si c es un número real cualquiera, la imagen inversa de $[c, \rightarrow[$ dada por u es cerrada en X . Una demostración similar se aplicaría a $] \leftarrow, c]$. En virtud de que una función de S a R es continua en S si y sólo si la imagen inversa de todo intervalo de R de la forma $] \leftarrow, y]$ o $[y, \rightarrow[$ es cerrado en S , se habrá así demostrado que u es continua en X . Si t es un número real, se escribirá

$$X_t = \{x \in X \mid u(x) \leq t\} \text{ y } X^t = \{x \in X \mid u(x) \geq t\}$$

Basta ahora considerar el caso en que c pertenece a $]a, b[$. Entonces el intervalo $[c, \rightarrow[$ es la intersección de los intervalos $[r, \rightarrow[$, donde $r \in Q'$ y $r \leq c$. Tomando las imágenes inversas dadas por u , se obtiene $X^c = \bigcap_{\substack{r \in Q' \\ r \leq c}} X^r$. Sea x un punto de X tal que

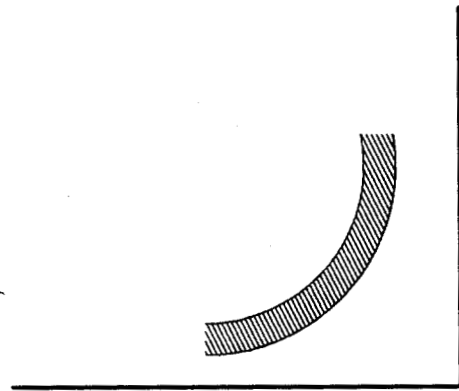
$$u(x) = r; \quad X^r = \{x' \in X \mid x (\leq) x'\}$$

que es cerrado en X en virtud de (a). Por lo tanto X^c , como intersección de conjuntos cerrados en X es cerrado en X . Esto concluye la demostración.

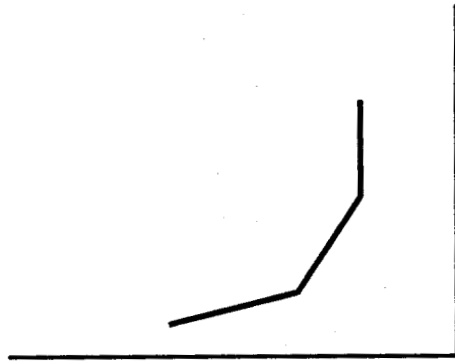
Dado un orden de preferencias ($\geq$) en X , podemos definir las siguientes relaciones *bajo la suposición de que X es un conjunto convexo*. Para dos puntos cualesquiera $x, y \in X$ es posible tener

- i) $x (\geq) y$ implica $tx + (1 - t)y (\geq) y$, $0 < t < 1$, donde $x \neq y$.
- ii) $x (>) y$ implica $tx + (1 - t)y (>) y$, $0 < t < 1$, donde $x \neq y$.
- iii) $x (\geq) y$ implica $tx + (1 - t)y (>) y$, $0 < t < 1$, donde $x \neq y$.
- iv) $x (=) y$ implica $tx + (1 - t)y (\geq) y$, $0 < t < 1$.

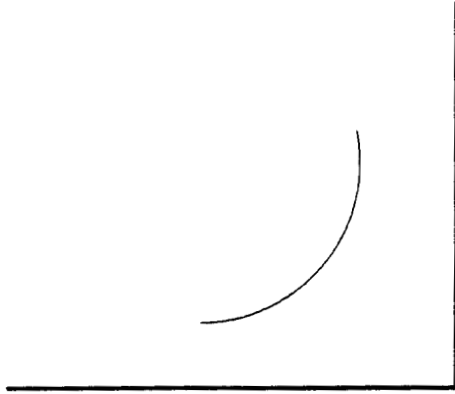
Un orden de preferencias es llamado *débilmente convexo* si (i) se cumple; se llama *convexo* cuando se cumplen (ii) y (iv) se cumplen, y *estrictamente convexo* cuando tienen lugar (ii) y (iii). Estos ordenes de preferencias se ilustran en la siguiente figura



CONVEXIDAD DEBIL



CONVEXIDAD



CONVEXIDAD ESTRICTA

Obsérvese que la convexidad débil en el orden de preferencias permite la existencia de curvas de indiferencia constituidas por "bandas", mientras que en los otros tipos de convexidad, las curvas de indiferencia no pueden tener superficie. Se puede advertir, que si el orden de preferencias es representado por un función de valores reales (como la función de utilidad), la concavidad estricta de la función de utilidad corresponde a la convexidad estricta en el orden de las preferencias, y la concavidad de la función de utilidad corresponde a la convexidad débil de el orden de preferencias.

Cabe hacer las siguientes observaciones: (1) La condición (i) implica que todos los contornos superiores de X deben ser convexos. (2) Si el orden de preferencias es continuo, (ii) implica (i). (3) Si el orden de preferencias es continuo, (iii) implica (ii).

CAPITULO 5. TEOREMAS DE SCARF, BROUWER Y KAKUTANI.

La prueba del teorema de Brouwer (enunciado más adelante) se sustenta en la construcción de una proyección continua de un conjunto en sí mismo, para posteriormente demostrar que, por lo menos, un punto permanece invariante bajo la proyección. Para nuestros fines basta limitar la atención a los dominios que son subconjuntos de espacios vectoriales de dimensiones finitas. El carácter del dominio en que se define la proyección es crucial para la validez de cualquier teorema de punto fijo. Por ejemplo, si consideramos el conjunto definido por una circunferencia, una rotación de $\pi/4$ es una transformación continua que no tiene punto fijo; en cambio, si consideramos todo el círculo, tanto el interior como la circunferencia, tal rotación deja invariante al centro. Aunque se conocen teoremas más generales, bastará con suponer que el dominio es un conjunto convexo (una condición que elimina el caso de la circunferencia).

Es necesario limitar también nuestra atención a conjuntos cerrados y acotados. Por ejemplo, proyectemos en sí mismo el intervalo abierto $(0,1)$, moviendo cada punto la mitad de su distancia al límite superior, es decir, proyectando x en $(x + 1)/2$. Esta proyección no tiene punto fijo en el intervalo abierto; en cambio, si lo cerramos agregando los puntos extremos, entonces 1 es un punto fijo. Es esencial también el carácter acotado: la proyección de x en $x + 1$ proyecta en sí misma toda la línea real, pero no tiene punto fijo.

Así pues, nos limitaremos a proyecciones de conjuntos convexos compactos en sí mismos. Podemos elaborar la teoría básica para un caso especial, donde el dominio sea un simplex, y luego extenderla

al caso general de un conjunto convexo compacto. Veamos ahora la definición de un simplex de manera ligeramente distinta a la que se presentó anteriormente.

Definición: Un simplex es la cápsula convexa de un conjunto finito de vectores linealmente independientes.

Definición: El *simplex fundamental* en el espacio n es el conjunto de n vectores

$$S_n = \{x \mid x_i \geq 0 \quad \sum x_i = 1\}.$$

Puede verificarse que S_n es la envoltura convexa de los n vectores unitarios, e_i ($i=1, \dots, n$), donde e_i es un vector con 1 en el lugar i y 0 en el resto.

Teorema (Brower): Si $f(x)$ es una proyección continua del simplex fundamental, S_n , es sí mismo, entonces existe $x^* \in S_n$ tal que $f(x^*) = x^*$.

Examinaremos a continuación un algoritmo combinatorio que no sólo genera una prueba del teorema de Brouwer, sino también un método para calcular el punto fijo a cualquier grado de aproximación que se desee. El algoritmo logra en realidad mucho más; prueba el teorema más general de Kakutani, que es esencial para demostrar la existencia del equilibrio en el modelo general que se discutirá después.

Revisaremos el método en una forma más o menos intuitiva antes de elaborarlo con precisión; se citarán también algunos resultados de Programación Lineal en el transcurso de la prueba.

Sea Y un conjunto de m elementos de vectores de n dimensiones, donde $m > n$ y se supone que Y incluye los n vectores e_i . Si $b \geq 0$ es cualquier vector dado, nos interesa la solución de

$$\sum_{y \in Y} y w(y) = b \quad w(y) \geq 0 \quad (1)$$

(a) En Y hay por lo menos un subconjunto, B (una base) de n vectores linealmente independientes, tal que (1) es válido con $w(y) = 0$ para $y \in Y \setminus B$. Para ver esto, sea $B = I$, la matriz identidad.

(b) Supongamos $y' \in Y \setminus B$ para alguna base B . Preguntamos si hay otra base, B' , que incluya y' y todos los elementos de B menos uno, tal que (1) pueda satisfacerse con $w(y) = 0$, para todo $y \in Y \setminus B'$. Adiviértase que $B \setminus \{y''\} = B' \setminus \{y'\}$, donde y'' es el elemento en $B \setminus B'$.

Definición: Decimos que B' es la inserción de y' en B si B y B' son ambas bases de (1) y $B' \setminus B = \{y'\}$.

Por la definición de B hay números $r(y)$ tales que

$$y' = \sum_{y \in B} yr(y). \quad (2)$$

Sea $\theta \geq 0$ un escalar y escribamos (1) como

$$\sum_{y \in B} yw(y) - \theta y' + \theta y' = b. \quad (1')$$

Sustituimos (2) en (1):

$$\sum_{y \in B} y[w(y) - \theta r(y)] + \theta y' = b. \quad (1'')$$

Sea $w(y, \theta) = w(y) - \theta r(y)$ para $y \in B = \theta$, e igual a 0 en los demás casos.

Supongamos $r(y) \leq 0$, para toda $y \in B$. Entonces (1) se satisfaría para θ arbitrariamente grande con $w(y) = w(y, \theta)$. Excluimos esta posibilidad suponiendo que el conjunto de soluciones de (1) es acotado de modo que $r(y) > 0$ para algún $y \in B$. Sea

$$\theta^* = \min_{\substack{y \in B \\ r(y) > 0}} \frac{w(y)}{r(y)} = \frac{w(y'')}{r(y'')}.$$

Si este mínimo no es único con $w(y) > 0$, entonces haciendo $\theta = \theta^*$ en (1'') obtenemos una solución nueva de (1):

$$\sum_{y \in B} yw(y) = b \quad w'(y) = w(y) - \theta^* r(y)$$

$$y \in B, w(y) = \theta^*$$

Cuando el mínimo no es único y/o $w(y) = 0$, se dirá que hay *degeneración*.

(c) Trataremos la degeneración en la forma siguiente. Digamos que el vector w_1 es *lexicográficamente mayor* que el vector w_2 si $w_{1j} > w_{2j}$ cuando $j = \min \{i \mid w_{1i} \neq w_{2i}\}$. Escribimos esto como $w_1 >_L w_2$. Sea ahora que $w(y)$ tenga $(n+1)$ dimensiones con las coordenadas numeradas $(0, \dots, n)$, y consideremos la ecuación y las desigualdades:

$$\sum_{y \in B} yw(y) = (b, I), \quad w(y) >_L 0 \text{ para todo } y \in B, \quad (3)$$

donde B es una base. Si puede satisfacerse (3), entonces B es una base, y advertimos que la solución de (3) contiene como la coordenada numerada $\langle 0 \rangle$ la solución del problema (1).

Supongamos ahora que $y' \in Y \setminus B$ y que hay una nueva base B' donde $B \setminus \{y'\} = B' \setminus \{y'\}$. Entonces, existe un vector v tal que.

$$\sum yw(y) - y'v + y'v = \sum y[w(y) - vr(y)] + y'v = (b, I) \quad (3)$$

y

$$v = \frac{w(y')}{r(y')}$$

Ciertamente, $v >_L 0$. Así que sólo necesitamos demostrar que y es el único elemento eliminado de B , es decir, que

$$\min_{\substack{y \in B \\ r(y) > 0}} \frac{w(y)}{r(y)}$$

es único.

Sin no, entonces $w(y'') = [r(y'')]/[r(y_1)]w(y_1)$, para $y_1 \neq y''$. Sea W la matriz $n \times n$ cuyas hileras son los vectores de n dimensiones $\cdot w(y)$ formados al eliminar de $w(y)$ el primer elemento. Sea $[y]_B W$ la matriz formada para $y \in B$. Entonces, por (3), $[y]_B W = I$. Postmultiplicamos ambos miembros de la ecuación por el vector k y advertimos que $Wk = 0$ implica $k = 0$, de modo que W es no singular. Entonces es imposible que dos hileras de W , $\cdot w(y'')$, $\cdot w(y_1)$, sean proporcionales, y así, el mínimo debe ser único.

El algoritmo requiere que consideremos otro conjunto, que por el momento escribiremos como Z , y también por el momento supondremos que tiene como elementos vectores z de n dimensiones. En Z nos interesan ciertos subconjuntos con la propiedad de que todos sus elementos se encuentran "cercaños" entre sí. Llamaremos a tales conjuntos *conjuntos primitivos de Z* . Consideraremos una proyección de uno a uno desde subconjuntos de Z hasta subconjuntos de Y .

Supongamos que hay un $D \subset Z$ tal que la proyección lleva a D hasta una base B en Y . Supongamos también que puede escogerse D en forma tal que todos sus elementos menos uno sean los mismos que en el conjunto primitivo en Z . Podríamos tratar de lograr que ambos conjuntos "concurden" (contengan los mismos elementos) eliminando el elemento que no concuerde del conjunto primitivo y encontrando otro que lo sustituya, o insertando en D el elemento del conjunto primitivo que no se encuentra en D , y eliminando por consiguiente un elemento de D . En virtud de la proyección, la última operación implica el encuentro de una nueva base para el problema (1); la primera operación implica el encuentro de un nuevo conjunto primitivo a partir de un conjunto primitivo dado. Si ocurre que siempre hay una forma única de elección de estas operaciones (como sucede con los cambios de la base) y si los pasos nunca forman un ciclo, podemos demostrar que, finalmente, se logrará hacer

concordar D con un conjunto primitivo de Z .

Para ver cómo este algoritmo puede conducir a un teorema de punto fijo, sea $b = e'$ en (1), donde e' es el vector de hilera de n dimensiones con uno en todos los lugares. Sea Z un subconjunto finito del simplex de n dimensiones, y $f(z)$ una proyección continua del simplex en sí mismo, de un sólo valor. Sea $y(z)$ la proyección de Z a Y mencionada antes, y escribimos

$$y(z) = f(z) - z + e.$$

Supongamos que cuando el algoritmo termina tenemos

$$\sum_{z \in P} y(z)w(z) = \sum_{z \in P} [f(z) - z + e]w(z) = e'.$$

donde P es el conjunto primitivo en Z donde termina el algoritmo. Dado que $ef(z) = ez = 1$, tenemos $e[f(z) - z + e'] = n$, de modo que premultiplicando (4) por e obtenemos

$$n \sum_{z \in P} w(z) = n$$

o sea

$$\sum_{z \in P} w(z) = 1 \quad (6)$$

Supongamos en seguida que los elementos de P están ahora tan "próximos" que podemos tomarlos por idénticos. Aquí estamos pensando en una operación limitante que, por supuesto, requiera que Z sea denso. Entonces, en vista de (5), (4) se convierte en $f(z) - z = 0$, que es el punto fijo de Brouwer.

Estas observaciones pueden servir para motivar el análisis que sigue, pero hay muchas complicaciones que hemos pasado por alto y que explican algunos de los procedimientos especiales que adoptamos.

(a) Podemos definir un conjunto primitivo de Z como sigue. Sea P un subconjunto de Z de n elementos. Sea z' el vector formado al hacer cada uno de sus componentes iguales al componente correspondiente más pequeño de cualquier vector P . Consideremos el conjunto tal que $z \geq z'$. Entonces, podemos decir que P es primitivo si el conjunto que acabamos de definir tiene un interior vacío. En tal caso, no hay elementos de P y, por lo tanto, no hay elementos de Z que se encuentren "entre" los elementos de P .

Sin embargo, surgen dificultades si hay elementos de P con coordenadas iguales a cero y si hay en P más de un vector que tenga una coordenada más pequeña dada. El algoritmo nos exige eliminar un elemento de P y remplazarlo en forma única por otro a fin de formar un nuevo conjunto primitivo.

(b) Se recordará que Y contiene los n vectores e_i y que nos interesa proyectar puntos de otro conjunto en Y . Convendrá remplazar Z por el conjunto $X \cup J$, donde X es un conjunto de vectores de n dimensiones y J el conjunto de los enteros $1, \dots, n$. En nuestra proyección asociaremos el vector e_j en Y con cualquier $j \in J$.

Para formar un conjunto primitivo en $X \cup J$ que nos permita pasar de un conjunto al otro, debemos inducir un orden en los elementos de $X \cup J$. Para evitar la degeneración deseamos evitar los "empates". Podemos lograr esto mediante una permutación cíclica de la serie $1, \dots, n$.

Para hacer una ilustración, supóngase que los vectores de X son de tres dimensiones y $J = 1, 2, 3$. Consideremos una permutación cíclica de los índices, $2, 3, 1$, que podríamos llamar un "ordenamiento conforme al 2" porque ahora el 2 es el primer término. Entonces, en este ordenamiento $1 > 3$, puesto que 1 viene después que 3. Por supuesto, también $1 > 2$.

Consideremos ahora un ordenamiento conforme al 2 de las coordenadas de elementos de X . Es decir, la coordenada con el número $\langle 2 \rangle$ es la primera de cada vector, la coordenada con el número 3 es la segunda, y la coordenada con el número 1 es la

tercera. Tomemos en cuenta en el ordenamiento a dos elementos x , x' en X , con $x \neq x'$. Ciertamente deben diferir, por lo menos en una coordenada. Encontraremos la más pequeña de tales coordenadas en el ordenamiento conforme al 2, es decir, la primera coordenada en el ordenamiento conforme al 2 en que difieran. Supongamos que tiene el número $\langle 3 \rangle$. Entonces diremos que $x >_2 x'$ si $x_3 > x'_3$ y que $x < x'$ si $x_3 < x'_3$.

Para comparar un elemento de J con un elemento de X se requiere una convención, y escogemos la siguiente. En un ordenamiento al 2, consideramos $2 <_2 x$, para todo $x \in X$, y todos los demás enteros, es decir, 1 ó 3 son mayores al 2 que cualquier elemento de X .

En esta ilustración la elección de un ordenamiento al 2 es arbitraria; podríamos haber tomado un ordenamiento al 3 ó al 1 de manera análoga.

Advertimos que para vectores desiguales en X se tiene $x >_2 x'$ ó $x <_2 x'$ de modo que están estrictamente ordenados. Además, hay un ordenamiento completo al 2 de todos los elementos de $X \cup J$, y dado que el mismo es estricto, hay un elemento mínimo al 2 único.

Con estas convenciones, sea Q un subconjunto de n elementos de $X \cup J$. Sea $\bar{x}^i(Q)$ el elemento mínimo conforme al i de Q . Entonces, decimos que Q es primitivo si no hay x en $X \cup J$ tal que $x >_i \bar{x}^i(Q)$, para todo i . Si hay un x tal decimos que el mismo domina en Q . En nuestro ejemplo i asume los valores 1,2,3.

Para ilustrar, sea X el conjunto de tres vectores: $x_1 = (3,7,6)$, $x_2 = (3,5,1)$, $x_3 = (5,4,1)$. Sea también $x^4 = 1$, $x^5 = 2$, $x^6 = 3$.

Supongamos que Q tiene los elementos x^1 , x^2 , y x^3 . Entonces $\bar{x}^1(Q) = x^4$, $\bar{x}^2(Q) = x^1$, $\bar{x}^3(Q) = x^2$. Entonces, ciertamente $x^1 >_1 \bar{x}^1(Q)$ para toda i , y el conjunto no puede ser primitivo.

Sea ahora que Q tenga los elementos x^1 , x^4 y x^6 . Entonces $\bar{x}^1(Q) = x^4$, $\bar{x}^2(Q) = x^1$, $\bar{x}^3(Q) = x^6$. Tenemos que ningún elemento de Q domina en Q . Tenemos también que $x^5 >_2 \bar{x}^2(Q)$, $x^2 >_2 \bar{x}^2(Q)$ y $x^3 >_2 \bar{x}^2(Q)$, de modo que ningún elemento fuera de Q domina en Q . Por lo

tanto (x^1, x^4, x^6) es un conjunto primitivo.

Adviértase que cuando consideramos el conjunto (x^1, x^2, x^4) tuvimos $x = \bar{x}^1(Q)$ para más de un i y encontramos que el conjunto no era primitivo. En realidad, puede demostrarse que si Q es primitivo, $x = \bar{x}^i(Q)$ para exactamente un i . Adviértase también que un conjunto primitivo no puede componerse de elementos j solamente, ya que por definición $x >_1 i$, si $x \in X$. Supongamos que partimos del conjunto primitivo $Q: (x^1, x^4, x^6)$ y se nos pide que eliminemos de allí a x^4 . El problema consiste en encontrar un remplazo para este elemento de modo que se forme un nuevo conjunto primitivo Q' . Llamemos x'' a este elemento -que todavía no encontramos- y supongamos primero que en verdad puede formarse un nuevo conjunto primitivo Q' . Procederemos por pasos.

(a) x^4 es mínimo al 1 en Q . Demostramos que x'' no puede ser mínimo al 1 en Q' . Esto puede verse porque los conjuntos $Q \setminus \{x^4\}$ y $Q' \setminus \{x''\}$ son el mismo. Por lo tanto, si x'' es mínimo al 1 en Q' , tenemos que

$$\bar{x}^2(Q') = x^1 \quad \bar{x}^3(Q') = x^6$$

Pero también $x'' \neq x^4$, de modo que $x'' >_1 x^4$ o bien $x'' <_1 x^4$. En el primer caso, dado que $x'' >_2 x^1$ y $x'' >_3 x^6$, x'' dominaría Q , lo que contradice que es primitivo. En el segundo caso $x^4 >_2 x^1$, $x^4 >_3 x^6$, y $x^4 >_1 x'' = \bar{x}^1(Q')$, lo que contradice que Q' es primitivo. Por lo tanto $x'' \neq \bar{x}^1(Q')$.

(b) Entonces $\bar{x}^1(Q') = x^1$ ó $\bar{x}^1(Q') = x^6$. Pero $x^6 >_1 x$, para todo $x \in X$, $i \neq 3$, y por lo tanto $\bar{x}^1(Q') = x^1$. (Esto demuestra que si hubiésemos eliminado a x^1 de Q no se habría podido encontrar ningún nuevo conjunto primitivo. Porque x^1 es mínimo al 2 y ni x^4 ni x^6 podrían ser mínimos al 2 en Q' . Es decir, no podemos proceder con el algoritmo si, tras de eliminar un elemento de Q ,

sólo nos quedan elementos de J . No sólo es x^1 el elemento mínimo al 1 en Q' , sino que también es el elemento mínimo al 1 en $Q \setminus \{x^4\} = (x^1, x^6)$.)

(c) Puesto que hemos demostrado que x^1 es mínimo al 1 en Q' y sabemos que x^6 es mínimo al 3 en Q' , debe ser que x'' sea mínimo al 2 en Q' . Para encontrar x'' consideremos el conjunto R de elementos de $X \cup J$ que son mayores al 1 que x^1 y mayores al 3 que x^6 . Ciertamente este conjunto tiene, por lo menos, un miembro, puesto que x^5 satisface los criterios. También x'' debe ser un miembro de este conjunto, que hemos definido y llamaremos x^{**} . Entonces, $x^{**} >_1 x^1$, $x^{**} >_3 x^6$. Si $x^{**} >_2 x''$, entonces x^{**} dominaría $Q' = (x^1, x^6, x^4)$, de modo que $x^{**} = x''$.

Dado que $x^2 <_1 x^1$ en nuestro ejemplo, x^2 no es candidato para ser x'' . Pero $x^3 >_1 x^1$ y $x^3 >_3 x^6$, de modo que x^3 es un candidato al igual que x^5 . Pero $x^5 <_2 x$ para todo x en X , de modo que $x^5 <_2 x^3$, $x^3 = x''$, ó $Q' = (x^1, x^6, x^3)$. Se puede verificar directamente que Q' es en verdad primitivo. Pero también podemos argumentar como sigue.

Supongamos que hay un x que domina Q' . Entonces, ciertamente x debe ser mayor al 1 que x^1 y mayor al 3 que x^6 y por lo tanto $x \in R = \{x \mid x \in X \cup J, x >_1 x^1, x >_3 x^6\}$. Pero también $x >_2 x''$, y esto contradice la definición de x'' como elemento R máximo al 2.

Introduciremos ahora varios conceptos nuevos como parte de la prueba en cuestión (las pruebas tradicionales de los teoremas de punto fijo requieren también la introducción de conceptos peculiares a las mismas, pero diferentes a los aquí expuestos).

Sea X un conjunto de vectores no negativos y J el conjunto de enteros $1, \dots, n$. El conjunto $X \cup J$ es entonces un conjunto bien definido, aunque algo extraño, compuesto de n vectores y n enteros. Estamos interesados en conjuntos Q que sean subconjuntos de n elementos de $X \cup J$. La idea es que los elementos enteros de Q , es decir, el conjunto $J \cap Q$, designen ciertas dimensiones que se tratan en forma especial. Entonces, dado que Q tiene n

elementos, hay exactamente tantos miembros de $J \setminus Q$ como de $X \cap Q$, y los pondremos en correspondencia uno a uno en formas apropiadas.

Queremos ordenar los elementos de X de acuerdo con su coordenada número i , pero al igual que en nuestro tratamiento de la degeneración, complementaremos las reglas para evitar los empates; es decir, si las coordenadas número i son iguales, ordenarán un elemento por encima de otro de acuerdo con los valores relativos de algunas otras coordenadas. También deseamos ordenar los elementos de X en relación con los de J ; convenimos en que h mismo es "menor en la coordenada número i " que cualquier elemento de X , pero que un entero $j \neq i$ es "mayor en la coordenada número i " que cualquier elemento de X .

Aunque el ordenamiento deseado puede introducirse de varias formas, resultará más conveniente, por razones de computación, ordenar los enteros de acuerdo con una permutación cíclica en la que el entero i ocupe el primer lugar; luego ordenamos los elementos de X de acuerdo con el componente de orden menor que rompa un empate, donde los componentes se ordenan de acuerdo con esta permutación cíclica. Específicamente, introduciremos un ordenamiento al i sobre J ,

$$j <_i j' \text{ si } i \leq j < j' \text{ o } i \leq j, j' < i \text{ o } j < j' < i \quad (6)$$

Esto equivale a poner los enteros $i, \dots, n$ antes de los enteros $1, \dots, i-1$ y a ordenarlos luego. La relación " $j <_i j'$ " se leerá " j es menor al i que j' ".

Ahora ordenamos los elementos de X de acuerdo con la coordenada número i si ello es posible; en caso contrario, por causa de algún empate, los ordenamos de acuerdo con la primera coordenada desigual en el ordenamiento al i de las coordenadas.

Para dos elementos, x y x' en X , sea $I(x, x')$ el conjunto de coordenadas i para el que $x_i \neq x'_i$, y sea j el elemento mínimo al i de $I(x, x')$. Entonces definimos que

$$x <_i x' \text{ significa que } x_j < x'_j \quad (7)$$

Adviértase que si $x_i \neq x'_i$, entonces el propio i es el elemento mínimo al i de $I(x, x')$ y $x <_i x'$ ó $x' <_i x$ dependiendo de que $x_i < x'_i$ ó $x'_i < x_i$.

Completamos la definición del ordenamiento al i para permitir comparaciones de elementos de J con elementos de X .

$$i <_i x \quad \text{para todo } x \text{ en } X \quad (8)$$

$$x <_i j \quad \text{para todo } x \text{ en } X \text{ y todo } j \neq i, j \in J \quad (9)$$

El ordenamiento al i es en realidad un ordenamiento estricto, de modo que para $x \neq x'$ es válida una y sólo una de las relaciones $x <_i x'$ y $x' <_i x$. Se sigue que para cualquier subconjunto de $X \cup J$ hay un mínimo al i que es único. Específicamente, para cualquier subconjunto de n elementos, Q , de $X \cup J$, sea

$$\bar{x}^i(Q) \text{ el elemento de } Q \text{ mínimo al } i \quad (10)$$

Notamos algunas consecuencias de estas definiciones. Por (8) y (9), si $i \in J \cap Q$, entonces $\bar{x}^i(Q) = i$. Si $i \notin J \cap Q$, de modo que $i \in J \setminus Q$, entonces $J \cap Q$ tiene a lo sumo $n - 1$ elementos, de modo que Q contiene, por lo menos, un elemento de X . En consecuencia, por (9), es imposible que $\bar{x}^i(Q) = j$ para $j \in J$, con $j \neq i$.

$$\text{Si } i \in J \cap Q, \text{ entonces } \bar{x}^i(Q) = i;$$

$$\text{Si } i \in J \setminus Q, \text{ entonces } \bar{x}^i(Q) \in X \cap Q \quad (11)$$

Si $\bar{x}^i(Q) > x_i$ para algún x en $X \cap Q$, ciertamente no podría ser verdad que $\bar{x}^i(Q)$ es el mínimo al i .

$$\text{Si } i \in J \setminus Q, \text{ entonces } \bar{x}^i(Q) = \min_{x \in Y \cap Q} x_i \quad (12)$$

Asociaremos con Q un cono contenido en el ortante no negativo. Definimos un vector $\xi(Q)$ por las relaciones de coordenadas:

$$\xi_i(Q) = \begin{cases} 0 & \text{si } i \in J \cap Q \\ \bar{x}^i(Q) & \text{si } i \in J \setminus Q \end{cases} \quad (13)$$

Luego definimos un cono convexo

$$T(Q) = \{x \mid x \geq \xi(Q)\} \quad (14)$$

Las definiciones anteriores equivalen a afirmar que escogemos $m \leq n$ elementos de X y un número correspondiente de direcciones de coordenadas; entonces $T(Q)$ es el cono más pequeño similar al ortante no negativo que contiene estos puntos y no restringido en todas las demás direcciones de coordenadas. Ahora introduciremos la siguiente

Definición: Un elemento x de X domina a un subconjunto de n elementos, Q , de $X \cup J$ ($x \text{ dom } Q$), si y sólo si $x >_i \bar{x}^i(Q)$ para todo i . Decimos que Q es primitivo si no está dominado por ningún elemento de X .

Si $T(Q)$ tuviese un elemento de X en su interior, entonces ciertamente sería dominado.

Si Q es primitivo, entonces el interior de $T(Q)$ es disjunto de X (15)

Lo recíproco no es cierto a causa de la posibilidad de empates en cualquier coordenada, que se rompen examinando otra coordenada. Tomemos en cuenta primero algunas propiedades elementales de los conjuntos primitivos. Supongamos que Q no contuviese elementos de X , es decir que $Q \subset J$. Dado que tanto Q como J tienen n elementos,

$Q = J$. Entonces $\bar{x}^i(Q) = i$ para todo i ; pero, por (8), $x \in Q$ para cualquier $x \in X$.

Si Q es primitivo, $X \cap Q$ es no vacío (16)

Supongamos que $x \in Q$, $x \neq \bar{x}^i(Q)$. Por (10), entonces $x >_i \bar{x}^i(Q)$. Por lo tanto, si Q es primitivo, para todo $x \in Q$, $x = \bar{x}^i(Q)$ para algún i , porque de otro modo $x \in Q$. Dado que tanto J como Q tienen n elementos, es imposible que $x = \bar{x}^i(Q)$ para más de un i , porque entonces habría algún $x' \neq \bar{x}^i(Q)$, para todo i .

Si Q es primitivo, entonces para todo $x \in Q$,

$x = \bar{x}^i(Q)$ para exactamente un i (17)

Se sigue de (16) que

$T(Q) \cap X$ no es vacío (18)

Un conjunto primitivo es, en algunos sentidos sugerentes, análogo a una base en programación lineal; el conjunto $J \cap Q$ es análogo al conjunto de vectores "sobrantes". Un paso fundamental en la programación lineal es la transformación mediante la cual se elimina de una base un elemento y se reemplaza por otro. En general, este reemplazo es único. La definición siguiente introduce un paso análogo.

Definición: Decimos que Q' es un reemplazo de x' en Q si Q y Q' son conjuntos primitivos y $Q \setminus Q' = \{x'\}$.

En otras palabras, x' es eliminado del conjunto primitivo y reemplazado por otro elemento de $X \cup J$, digamos x'' , en forma tal que el nuevo conjunto, Q' , también es primitivo.

Teorema: Si Q es primitivo y $x' \in Q$, entonces no hay remplazo de x' en Q si $Q \setminus \{x'\} \subset J$, y exactamente un remplazo en caso contrario.

Prueba.

Suponemos que existe tal remplazo y encontramos las condiciones necesarias. Sea x'' el único elemento de $Q' \setminus Q$, es decir, el nuevo elemento en Q' . Por (17), podemos suponer sin pérdida de generalidad que

$$x' = \bar{x}^1(Q).$$

Primero, supongamos que $x'' = \bar{x}^1(Q)$; demostraremos que esto es imposible. Por (17), tendríamos para $i > 1$ que

$$\bar{x}^1(Q) \in Q \setminus \{x\} \quad \text{y} \quad \bar{x}^1(Q') \in Q' \setminus \{x''\}.$$

Dado que $\bar{x}^1(Q)$ es mínimo al i en Q , ciertamente es mínimo al i en el subconjunto $Q \setminus \{x\}$, y de igual modo $\bar{x}^1(Q')$ es mínimo al i en $Q' \setminus \{x''\}$, que, sin embargo, es el mismo conjunto, es decir

$$\bar{x}^1(Q) = \bar{x}^1(Q') \quad \text{para } i > 1$$

Por (10) y (17), $x >_1 \bar{x}^1(Q) = \bar{x}^1(Q')$ para $i > 1$, y de igual modo $x'' >_1 \bar{x}^1(Q)$ para $i > 1$. Dado que $x' \neq x''$, se tiene $x' >_1 x''$ o bien $x'' >_1 x$. Puesto que $x' = \bar{x}^1(Q)$, $x'' = \bar{x}^1(Q')$, tenemos en el primer caso que $x' \text{ dom } Q'$ y en el segundo caso que $x'' \text{ dom } Q$, y ambos resultados contradicen el carácter primitivo que hemos supuesto en Q y Q' . Así pues, la suposición de que $x'' = \bar{x}^1(Q')$ conduce a una contradicción.

Entonces, $\bar{x}^1(Q')$ debe pertenecer a $Q' \setminus \{x''\} = Q \setminus \{x\}$. Sea

$$x^* = \bar{x}^1(Q') \quad (19)$$

Dado que $x^* \in Q$, $x^* = \bar{x}^1(Q)$, para algún i , pero $x^* \neq x' = \bar{x}^1(Q')$, tenemos que $i > 1$. Sin pérdida de generalidad podemos suponer que

$$x^* = \bar{x}^2(Q') \quad (20)$$

Si $x^* \in J$, entonces debemos tener, por (11), $x^* = 1$, de (19), y $x^* = 2$, de (20), lo que sería una contradicción. Por lo tanto,

$$x^* \in X \cap [Q \setminus \{x'\}] \quad (21)$$

Este enunciado muestra ya que la existencia de un remplazo de x' en Q requiere que $X \cap [Q \setminus \{x'\}]$ sea no vacío. Por lo tanto, si $Q \setminus \{x'\} \subset J$, entonces no hay tal remplazo, lo que confirma la primera parte del teorema.

Por (10), x^* es mínimo al 1 en Q' ; por (21),

$$x^* \text{ es mínimo al 1 en } X \cap [Q' \setminus \{x''\}] = X \cap [Q \setminus \{x'\}] \quad (22)$$

de modo que x^* se define en forma única y por un procedimiento constructivo.

Por (19) y (17), x^* no es mínimo al 2 en Q' . Por (20), x^* es mínimo al 2 en Q y por lo tanto en $Q \setminus \{x'\} = Q' \setminus \{x''\}$; en consecuencia

$$x'' = \bar{x}^2(Q') \quad (23)$$

Ahora $Q \setminus [\{x'\} \cup \{x^*\}] = Q' \setminus [\{x''\} \cup \{x^*\}]$. Para $i > 2$, $\bar{x}^1(Q)$ es el mínimo al i en el primero de estos conjuntos y $\bar{x}^1(Q')$ lo es en el segundo.

$$\bar{x}^1(Q) = \bar{x}^1(Q') \quad \text{para } i > 2 \quad (24)$$

Buscamos ahora una caracterización de x'' . Definimos el conjunto

$$R = \{x \mid x \in X \cup J, x >_1 x^*, x >_1 \bar{x}^i(Q) \text{ para } i > 2\} \quad (25)$$

Primero, advertimos que R es necesariamente no vacío, ya que $2 \in R$; por (9), $2 >_1 x$ para todo $x \in X$ si $i \neq 2$, mientras que por (6), $2 >_1 i$ para todo $i \in J$, $i \neq 2$.

Advertimos también que, por (23) y (17), $x'' \in R$. Sea x^{**} el elemento de R máximo al 2. Por (19) y (24), se sigue de la definición de R que,

$$x^{**} >_1 \bar{x}^i(Q') \quad \text{para } i \neq 2$$

Si $x^{**} >_2 x''$, entonces, por (23), $x^{**} \text{ dom } Q'$, lo que contradice la hipótesis. Por lo tanto, $x^{**} = x''$, por definición de un elemento máximo al 2.

$$x'' \text{ es el elemento de } R \text{ máximo al } 2 \quad (26)$$

Por (22), x^* se define en forma única, y por lo tanto ocurre lo mismo con x'' , por (25) y (26). Así pues, hay a lo sumo un remplazo de x' en Q .

Dado que Q' , definido como $[Q \setminus \{x'\}] \cup \{x''\}$, tiene n elementos, basta demostrar que es primitivo. Primero demostraremos que $\bar{x}^1(Q)$ son en verdad como se dan en (19), (23) y (24). Si $1 \in Q$, entonces $x' = 1$, de modo que, en todo caso, $1 \notin Q \setminus \{x'\}$. Se sigue entonces de (22) y (9) que x^* es mínimo al 1 en $Q \setminus \{x'\} = Q' \setminus \{x''\}$. Por (25) y (26), $x'' >_1 x^*$, lo que confirma (24).

Por (25), x^* no pertenece a R . Supongamos que $x'' >_2 x^*$. Dado que $x^* = \bar{x}^2(Q)$, por (20), tenemos, por (25) que $x'' >_1 \bar{x}^i(Q)$ para $i > 1$. Dado que $x^* >_1 \bar{x}^2(Q)$, por (20), y $x'' >_1 x^*$, por definición de

R, tenemos, por la transitividad, que $x'' >_1 \bar{x}^1(Q)$, de modo que $x'' \text{ dom } Q$, lo que contradice la hipótesis. Por lo tanto, $x^* >_2 x''$.

De nuevo por (20), x^* es mínimo al 2 en Q y por lo tanto en $Q \setminus \{x'\} = Q' \setminus \{x''\}$ de modo que (23) debe darse.

Ahora puede demostrarse que Q' es primitivo. Si $x \text{ dom } Q'$, se sigue de (19), (24) y (25) que $x \in R$, mientras que también $x >_2 \bar{x}^2(Q')$, lo que contradice (26). Esto completa la prueba del teorema.

Corolario: Si Q' es el remplazo de x' en Q y x'' es el único elemento de $Q' \setminus Q$, entonces Q es el remplazo de x'' en Q .

Esto se sigue inmediatamente del carácter único del remplazo y de su definición.

Observación: Se ha demostrado que el remplazo único de x' en Q puede lograrse en un número finito de pasos, por (22), (25) y (26). La determinación de x^* en (22) implica una búsqueda entre $n-1$ alternativas a lo sumo, y en la práctica se tomará generalmente a n como un número fijo (por ejemplo, el número de bienes). Sin embargo, supondremos que el conjunto X es muy grande; en realidad, en la determinación de un punto fijo, la aproximación es mejor cuanto mayor sea el conjunto X . A primera vista, parecería que la determinación del conjunto R requiere una exploración en el total de $X \cup J$, lo que no es muy atrayente. Esto es cierto en general, pero el cálculo puede limitarse a un ejercicio menor si X es el conjunto de elementos de S_n con coordenadas racionales y un denominador fijo, es decir

$$X_n = \left\{ x \mid x_i = \frac{a_i}{N}, a_i \text{ enteros no negativos, } \sum_1 x_i = 1 \right\}$$

Teorema (Scarf): Sea X un conjunto de n vectores, $y(x)$ una función

que proyecta X en un conjunto Y de n vectores, I el conjunto de vectores unitarios y J el conjunto de enteros $1, \dots, n$. Si $b \geq 0$ y si el conjunto de desigualdades

$$\sum_{y \in Y \cup I} y w(y) = b \quad w(y) \geq 0$$

tiene un conjunto acotado de soluciones, entonces existen una base B para el sistema de desigualdades, y un conjunto primitivo Q para el conjunto $X \cup J$ tal que B es la imagen de Q bajo la proyección $y(x)$, donde se entiende que $y(j) = e^j$ para $j \in J$.

Prueba:

Para fines de esta prueba, llamaremos base a un subconjunto C de n elementos de $X \cup J$ si su imagen bajo la proyección $y(x)$, denotada por $y(C)$, es una base. Decimos que una base C concuerda con un conjunto primitivo Q si $Q = C$. Tratamos de demostrar que existe un par (C, Q) tal que C concuerda con Q . Decimos que una base C casi concuerda con un conjunto primitivo Q si $C \setminus Q = \{1\}$; es decir, $n-1$ elementos de $B = y(C)$ corresponden a elementos de Q , pero B contiene también el primer vector unitario.

Este método de prueba consiste en principiar con un par (C_0, Q_0) que casi concuerda, y luego continuar con una serie de transformaciones que conviertan un par que casi concuerda, o en otro par que casi concuerda, o en un par que concuerda. Se demostrará que la sucesión nunca puede repetirse; dado que hay un número finito de pares posibles, el proceso debe terminar en un par que concuerda, lo que no sólo demuestra el teorema, sino que además pone al descubierto un algoritmo para encontrar una base.

Si C es una base, entonces $y(C)$ es una base en su conjunto apropiado, $Y \cup I$. Dado que ambos conjuntos tienen n elementos, la proyección $y(x)$ tiene una inversa única para $x \in C$. Si $x' \notin C$, definimos C' como la inserción de x' en C , si C' es la inversa de la inserción de $y(x')$ en $y(C)$. En mayor detalle, si y'' es el elemento eliminado en esta inserción, entonces $y'' = y(x'')$ para

algún x'' único, y definimos C' como $[C \setminus \{x''\}] \cup \{x'\}$.

Puede ocurrir que $y(x')$ sea ya de hecho un elemento de $y(C)$; es decir, $y(x) = y(x'')$ para algún $x'' \in C$. No puede haber más de un x'' tal, puesto que los elementos de $y(C)$ son distintos. La inserción de $y(x)$ en $y(C)$ no tiene significado en este caso, o más precisamente, la "inserción" deja simplemente invariable la base, pero todavía podemos definir a C' exactamente en la misma forma, y C' es diferente de C .

Si C casi concuerda con Q , decimos que el par (C', Q') es *adyacente* a (C, Q) si cualquiera de estas dos cosas es posible:

- (a) $C = C'$, Q' es el remplazo de $x(C, Q)$ en Q ;
- (b) $Q = Q'$, C' es la inserción de $x(C, Q)$ en C .

De acuerdo con la teoría de Programación Lineal, siempre es posible la transformación (b). Demostramos que la transformación (a) es posible si y sólo si $C \neq J$.

Supongamos primero que $C = J$. Entonces $C \cap Q$ debe consistir en los enteros $2, \dots, n$. Dado que C y Q no concuerdan, $x(C, Q)$ debe ser un elemento de X . Por lo tanto, $Q \setminus \{x(C, Q)\} \subset J$, y por el teorema arriba citado, no es posible ningún remplazo.

Supongamos ahora $C \neq J$. Entonces C contiene un elemento, x , de X . Dado que $C \setminus Q = \{1\}$, $x \in Q$. Pero $x \neq x(C, Q)$, por (27), de modo que $Q \setminus \{x(C, Q)\}$ contiene un elemento de X , y es posible un remplazo. Hemos demostrado que si $C \neq J$, son posibles las transformaciones (a) y (b), mientras que si $C = J$, sólo (b) es posible.

(C, Q) tiene un par adyacente si $C = J$,
dos pares adyacentes de otro modo. (28)

Si (C', Q') es adyacente a (C, Q) y C' casi concuerda con Q' , entonces se sigue del último Corolario que (C, Q) es adyacente a (C', Q') .

Observamos también que si $C = J$, sólo hay un Q posible que casi

concuerde con C . Como vimos antes, Q debe contener $2, \dots, n$. Entonces $\bar{x}^1(Q) = i$, $i > 1$, mientras que $\bar{x}^1(Q) \in X$. Sea x^* el elemento de X máximo al 1. Entonces $x^* >_1 i$ ($i > 1$); si también $x^* >_1 \bar{x}^1(Q)$, entonces x^* domina Q , y Q no sería primitivo. Por lo tanto, si $C = J$ y C casi concuerda con Q , entonces Q se compone de x^* y los enteros $2, \dots, n$; llamamos Q_0 a este conjunto.

Ahora generamos una sucesión de pares, (C_v, Q_v) , como sigue. Sea $C_0 = J$ y Q_0 como acaba de especificarse. Entonces hay precisamente un par adyacente a (C_0, Q_0) ; lo llamamos (C_1, Q_1) . En general, si $C_v \neq J$, hay dos pares adyacentes a (C_v, Q_v) ; sea (C_{v+1}, Q_{v+1}) adyacente a (C_v, Q_v) y distinto de (C_{v-1}, Q_{v-1}) . Por inducción se verifica que si C_v casi concuerda con Q_v , entonces (C_{v-1}, Q_{v-1}) es adyacente a (C_v, Q_v) , de modo que (C_{v+1}, Q_{v+1}) queda especificado de forma única.

Adviértase también que, por la definición de *adyacente*, si C casi concuerda con Q y (C', Q') es adyacente a (C, Q) , entonces C' concuerda o casi concuerda con Q' .

Ahora demostraremos que la sucesión nunca se repite; es decir, es imposible que $(C_\mu, Q_\mu) = (C_{v+1}, Q_{v+1})$ para $\mu \leq v$. De otro modo escójase el más pequeño de tales v . Puesto que dos pares adyacentes nunca son iguales, es imposible de $\mu = v$, y por la construcción es imposible que $\mu = v - 1$. Por lo tanto $\mu < v - 1$.

Supongamos que $\mu = 0$. Entonces (C_v, Q_v) sería adyacente a (C_0, Q_0) al igual que (C_1, Q_1) . Pero sólo hay un par adyacente a (C_0, Q_0) por lo tanto $(C_1, Q_1) = (C_v, Q_v)$. Dado que $v > 1$, como acabamos de ver, esto contradice la elección de v como el valor más pequeño para el que $(C_\mu, Q_\mu) = (C_{v+1}, Q_{v+1})$, para algún $\mu \leq v$.

Supongamos que $\mu > 0$. Entonces (C_{v-1}, Q_{v-1}) , (C_{v+1}, Q_{v+1}) y (C_v, Q_v) serían todos adyacentes a (C_μ, Q_μ) . En vista de (28) dos tendrían que ser iguales; pero tal igualdad contradice la elección de v .

En consecuencia, todos los elementos de la sucesión (C_v, Q_v) son distintos. La sucesión puede continuarse mientras que C_v casi concuerda con Q ; por otra parte, sólo hay un número finito de bases posibles y sólo un número finito de pares (C, Q) que casi

concuerdan. Por lo tanto, dado que todos los elementos de la sucesión concuerdan o casi concuerdan, y que la sucesión termina en el primer caso, se sigue que en un número finito de pasos el algoritmo alcanza un par que concuerde.

El teorema de Scarf provee una prueba notablemente rápida del teorema de punto fijo de Brouwer. Para cualquier entero dado, N , consideramos que el conjunto X es el conjunto X_N definido en la última observación y luego hacemos una elección apropiada de $y(n)$.

Lema: Sea Q_N cualquier conjunto primitivo para $X_N \cup J$, y sea x_n cualquier elemento de $X_N \cap Q_N$. Escogiendo una subsucesión apropiada, podemos encontrar un subconjunto propio, J^* , de J , y un punto, x^* , tal que $J \cap Q_N = J^*$ y x_n tienda a x^* a lo largo de la subsucesión, y $x_i = 0$ para todo $i \in J^*$. La subsucesión y el punto x^* son independientes de la forma en que se escoja x_n en $X_N \cap Q_N$.

Prueba:

Recordemos la definición de $T(Q_N)$ en (14). Dado que $X_N \subset S_n$ se sigue de (18) que $T(Q_N) \cap S_n$ es no vacío; sea x un elemento. Entonces $ex = 1$ y, por (14),

$$e\xi(Q_N) \leq 1 \quad (29)$$

Sea U el subconjunto del ortante no negativo, consistente en todos los puntos que se encuentran en el simplex fundamental por debajo del mismo, es decir,

$$U = \{x \mid x \geq 0, \quad ex \leq 1\}.$$

El subconjunto $T(Q_N)$ es geoméricamente semejante al ortante no negativo, y por lo tanto el conjunto $U \cap T(Q_N)$ es semejante al conjunto U .

Considérese, en particular, el radio de la vecindad más grande contenida en $T(Q_N) \cap S_n$, donde la vecindad se toma en relación con

S_n ; el radio de la vecindad más grande en relación con S_n , que es disjunta de X_n , tiende a 0 cuando N tiende a infinito. Dado que, por (15), el interior de $T(Q_N)$ es disjunto de X_N , lo mismo es cierto para cualquier vecindad contenida en $T(Q_N) \cap S_n$. De todo esto se sigue que

$$\lim_{N \rightarrow \infty} e\xi(Q_N) = 1 \quad (30)$$

Para cada N , $J \cap Q_N$ es un subconjunto propio de J ; dado que sólo hay un número finito de tales conjuntos, debe haber, por lo menos, un conjunto, J^* , tal que $J \cap Q_N = J^*$ para un número infinito de N . De ahora en adelante, consideremos sólo tal N . La sucesión $\xi(Q_N)$ es acotada; por lo tanto podemos escoger otra subsucesión tal que

$$\xi(Q_N) \rightarrow x^* \quad (31)$$

Por la definición de $\xi(Q)$ en (13), $\xi_i(Q_N) = 0$ para $i \in J$, de modo que ciertamente $x_i^* = 0$ para $i \in J^*$, como se propone. Por último, dado que $x_n \in X_N \cap Q_N$, tenemos a la vez $x_n \in S_n$ y $x_n \in T(Q_N)$, de modo que

$$ex_n = 1 \quad x_n \geq \xi(Q_N).$$

Entonces se sigue de (29)-(31) que $x_n \rightarrow x^*$ a lo largo de la subsucesión.

Ahora enunciaremos un lema que muestra, en general, la forma en que puede alcanzarse el límite en la aplicación del teorema de Scarf; este lema no se utiliza sólo en la prueba de los teoremas de punto fijo de Brouwer y Kakutani, sino también para probar algunos otros resultados importantes.

Lema: Si $y(x)$ es una proyección acotada de S_n y los conjuntos de soluciones a las desigualdades

$$\sum y(x)w(x) = b \quad w(x) \geq 0$$

para todos los subconjuntos finitos X de S_n están acotados, entonces existe un subconjunto propio, J^* , de J , $u^* \geq 0$, $v_j^* \geq 0$ ($j = 1, \dots, m$), donde m es el número de elementos en $J \setminus J^*$, x^* en S_n e y^j ($j = 1, \dots, m$) donde para cada j , y^j es un punto límite de $y(x)$ a medida que x se aproxima a x^* , tal que

$$u^* + \sum_{j=1}^m v_j^* y^j = b, \quad (a)$$

$$\begin{aligned} u_i^* &= 0 \text{ para } i \in J \setminus J^*, \\ x_i^* &= 0 \text{ para } i \in J^*. \end{aligned} \quad (b)$$

Prueba:

Por el teorema de Scarf, para cada N existe Q_N , que es a la vez un conjunto primitivo y una base para las desigualdades

$$\sum y(x)w(x) = b \quad w(x) \geq 0$$

Entonces, por la definición de una base, podemos encontrar $w_N(x) \geq 0$, $x \in Q_N$, tal que

$$\sum_{x \in Q_N} y(x)w_N(x) = b \quad (32)$$

Escogemos la subsucesión del lema 1 y el J^* correspondiente. Sea $u_1^n = w_N(i)$ para $i \in J^*$, $u_1^n = 0$ para $i \in J \setminus J^*$. Dado que $y(i) = e_i$,

$$\sum_{x \in J^*} y(x)w(x) = \sum_{i \in J^*} u_1^n e_i = \sum_{i=1}^n u_1^n e_i = u^n \quad (32)$$

donde u^n es el vector de componentes u_1^n . Enumeramos los m elementos de $Q_N \setminus J^*$ arbitrariamente como $x^{1N}, \dots, x^{mN}$, y hacemos

$wn(x^{jN}) = v_j^N$. Entonces (32) puede escribirse

$$u^N + \sum_{j=1}^m v_j^N y(x^{jN}) = b,$$

donde $u^N \geq 0$, $v_j^N \geq 0$, $u_i^N = 0$ para $i \in J \setminus J^*$. Por el lema 1 existe x^* tal que $x^{jN} \rightarrow x^*$ para todo j en una subsucesión apropiada, y $x_i^* = 0$ para $i \in J^*$. Por la hipótesis de acotamiento del lema, las sucesiones u^N , v_j^N , e $y(x^{jN})$ son todas acotadas; por lo tanto, escogiendo una nueva subsucesión, podemos asegurar que

$$u^N \rightarrow u^* \quad v_j^N \rightarrow v_j^* \quad y(x^{jN}) \rightarrow y^j,$$

de modo que se dan (a) y (b); que $u_i^* = 0$ para $i \in J \setminus J^*$ se sigue del hecho de que $u_i^N = 0$ para tales i para todo N .

Prueba del Teorema de Brouwer: Ahora aplicamos el lema 2 con una elección apropiada de $y(x)$ y b . El teorema de Brouwer supone que $f(x)$ es una proyección continua de S_n en sí mismo. Por lo tanto $ef(x) = 1 = ex$, para todo $x \in S_n$.

Definimos ahora

$$y(x) = f(x) - x + e, \quad b = e', \quad \text{para } x \in S_n. \quad (33)$$

donde e' es el vector de columna cuyos componentes son todos 1. Primero verificamos que las hipótesis de acotamiento son válidas; $y(x)$ es una función acotada. Sea X cualquier subconjunto finito de S_n , y consideremos las ecuaciones y desigualdades

$$\sum y(w)w(x) = e' \quad w(x) \geq 0.$$

Dado que $y(i) = e_i$ para $i \in J$, haciendo $u_i = w(i)$ y u el vector con componentes u_i , podemos escribir el sistema como

$$u + \sum_{x \in X} y(x)w(x) = e' \quad u \geq 0, w(x) \geq 0, \text{ para todo } x \in X \quad (34)$$

Por (33), $ey(x) = ef(x) - ex + ee' = ee' = n$, para todo $x \in X$.
También, $eu \geq 0$, $eu = 0$ si y sólo si $u = 0$. Si en (34)
multiplicamos por e la expresión de la izquierda,

$$eu + n \sum_{x \in X} w(x) = n,$$

de modo que

$$\begin{aligned} \sum_{x \in X} w(x) &\leq 1 & w(x) &\geq 0, \text{ para todo } x \in X; \\ \sum_{x \in X} w(x) &= 1 & \text{ si y sólo si } u &= 0, \end{aligned} \quad (35)$$

y las soluciones para $w(x)$ (con $x \in X$) son, en particular,
acotadas; por (34), lo mismo debe aplicarse a u .

Por lo tanto podemos aplicar el Lema 2 de modo que

$$u^* + \sum_{j=1}^m v_j^* y^j = e' \quad (36)$$

donde se aplica también el lema 2 (b), y por (35)

$$\sum_{j=1}^m v_j^* \leq 1 \text{ y la igualdad implica que } u^* = 0 \quad (37)$$

Aquí, y^j es un punto limitante de $y(x)$ cuando x tiende a x^* . Por
(33) entonces

$$y^j + x^* - e' \text{ es una cota de } f(x) \text{ cuando } x \text{ tiende a } x^* \quad (38)$$

Dado que suponemos continua a $f(x)$,

$$y^j + x^* - e = f(x^*) \text{ para todo } j \quad (39)$$

Sea ahora d el vector de hilera, con $d_i = 0$ para $i \in J^*$, $d_i = 1$

para $i \in J \setminus J^*$. Por el último lema inciso (b), $du^* = 0$. También, $di \neq e_i$ sólo para $i \in J^*$ y por lo tanto sólo si $x_i^* = 0$, debido al mismo lema y al mismo inciso; en consecuencia $dx^* = ex^* = 1$. Por último, $d \leq e$, de modo que $df(x) \leq ef(x) = 1$. Por (38) ó (39),

$$d(y^j + x^* - e) \leq 1 \text{ para cada } j.$$

Dado que $dx^* = 1$, este enunciado puede escribirse $dy^j \leq de$. Si premultiplicamos (36) por d , encontramos

$$(de') \sum_{j=1}^m v_j^* \geq de.$$

En conjunción con (37), vemos que

$$\sum v_j^* = 1 \quad u^* = 0.$$

Si sustituimos en (36), tenemos

$$y^* = e' \quad \text{donde} \quad y^* = \sum_{j=1}^m v_j^* y^j \quad y$$

$$\sum v_j^* = 1, \quad v_j^* \geq 0, \text{ para todo } j \quad (40)$$

Si multiplicamos (39) por v_j^* y luego sumamos a través de j , tenemos

$$y^* + x^* - e' = f(x)^*,$$

lo cual, en vista de (40), establece el teorema de Brouwer.

Ahora enunciaremos y probaremos la generalización del teorema de Brouwer debida a Kakutani.

Teorema (Kakutani): Sea C un conjunto convexo compacto, y $\Phi(x)$ una correspondencia semicontinua superiormente definida en C tal que

$\Phi(x) \subset C$ y $\Phi(x)$ sea convexa, para cada $x \in C$. Entonces existe un $x^* \in C$ tal que $x^* \in \Phi(x^*)$.

Adviértase que si $\Phi(x)$ consiste en un solo punto para cada x , esto se reduce al teorema de Brouwer para un dominio convexo compacto.

Prueba:

La prueba anterior del teorema de Brouwer basta, con ligeras modificaciones, para establecer el teorema de Kakutani para el caso en que C es simplex fundamental, S_n . Entonces, el caso general puede tratarse de la siguiente forma.

Para cada $x \in S_n$, escogemos $f(x)$ como un elemento arbitrario de $\Phi(x)$. Luego definimos $y(x)$ como antes. Hasta (38), el argumento no utiliza la continuidad de $f(x)$ y así sigue siendo válido; también sigue siendo válido (40). Dado que $f(x) \in \Phi(x)$, que es semicontinua superiormente, se sigue de (38) y de la definición de semicontinuidad superior que

$$y^j + x^* - e \in \Phi(x^*), \text{ para cada } j.$$

Dado que, por (40),

$$v_j^* \geq 0, \text{ para cada } j, \text{ y}$$

$$\sum_{j=1}^m v_j^* = 1$$

y puesto que $\Phi(x^*)$ es convexa por hipótesis,

$$\sum_{j=1}^m v_j^* (y^j + x^* - e) \in \Phi(x^*),$$

lo que puede escribirse como

$$y^* + x^* - e' \in \Phi(x).$$

Por (40), esto establece el teorema de Kakutani.

Para extender esto al caso en el que el dominio es un conjunto convexo compacto cualquiera, advertimos primero que, en realidad, no hay pérdida de generalidad al suponer que C es un subconjunto de un simplex fundamental. Supongamos que la dimensión de C es $n-1$; esta es la dimensión del espacio lineal más pequeño que contiene C . Escogemos un sistema coordenado en este espacio de $n-1$ dimensiones; esto puede tomarse como una transformación lineal de las variables originales. Cambiando el origen, podemos suponer que C está contenido en el ortante no negativo. Dado que C es acotado, ex es acotado en C ; por otra transformación sencilla podemos suponer que $ex \leq 1$. Por último, proyectamos C linealmente en un subconjunto convexo compacto de S_n , y tiene la misma dimensión. Puesto que todas las transformaciones son lineales, siguen siendo válidas las hipótesis del teorema de Kakutani.

Ahora escogemos x^0 como cualquier punto de C interno en relación con el espacio lineal definido por la ecuación $ex = 1$, y hacemos que $p(x)$ sea la función calibradora de C en relación con este espacio. Entonces $p(x)$ es continua, y $p(x-x^0) \leq 1$ si y sólo si $x \in C$ para x en este espacio y en particular para $x \in S_n$. Sea $q(x) = \max [p(x-x^0)]$ de modo que $x \in C$ si y sólo si $q(x) = 1$. Definimos una proyección,

$$T(x) = x^0 + \frac{x - x^0}{q(x)};$$

que proyecta un elemento de C en sí mismo y un elemento x de $S_n \setminus C$ en el elemento de C que está más cerca de x en el segmento lineal que une a x^0 con x .

Dado que $q(x)$ está acotado al extremo opuesto de 0, $T(x)$ es continua. Por último, definimos la correspondencia

$$\Phi_1(x) = \Phi[T(x)] \text{ para todo } x \text{ en } S_n.$$

Para x en C , esto es simplemente la correspondencia original, Φ_1 es semicontinua superiormente, puesto que Φ es semicontinua superiormente y $T(x)$ es continua; $\Phi_1(x) = \Phi_1[T(x)]$, donde $T(x) \in C$ para todo $x \in S_n$, de modo que para todo x tal, $\Phi_1(x)$ es convexo y $\Phi_1(x) \subset C \subset S_n$.

Por lo tanto, $\Phi_1(x)$ satisface todas las hipótesis del teorema de Kakutani para el simplex fundamental, y hay un punto, $x^* \in S_n$, tal que $x^* \in \Phi_1(x^*)$. Dado que, por construcción, $\Phi_1(x) \subset C$ para todo $x \in S_n$, $x^* \in C$; entonces $T(x^*) = x^*$ y $\Phi_1(x^*) = \Phi(x^*)$, de modo que el teorema de Kakutani se aplica al conjunto C .

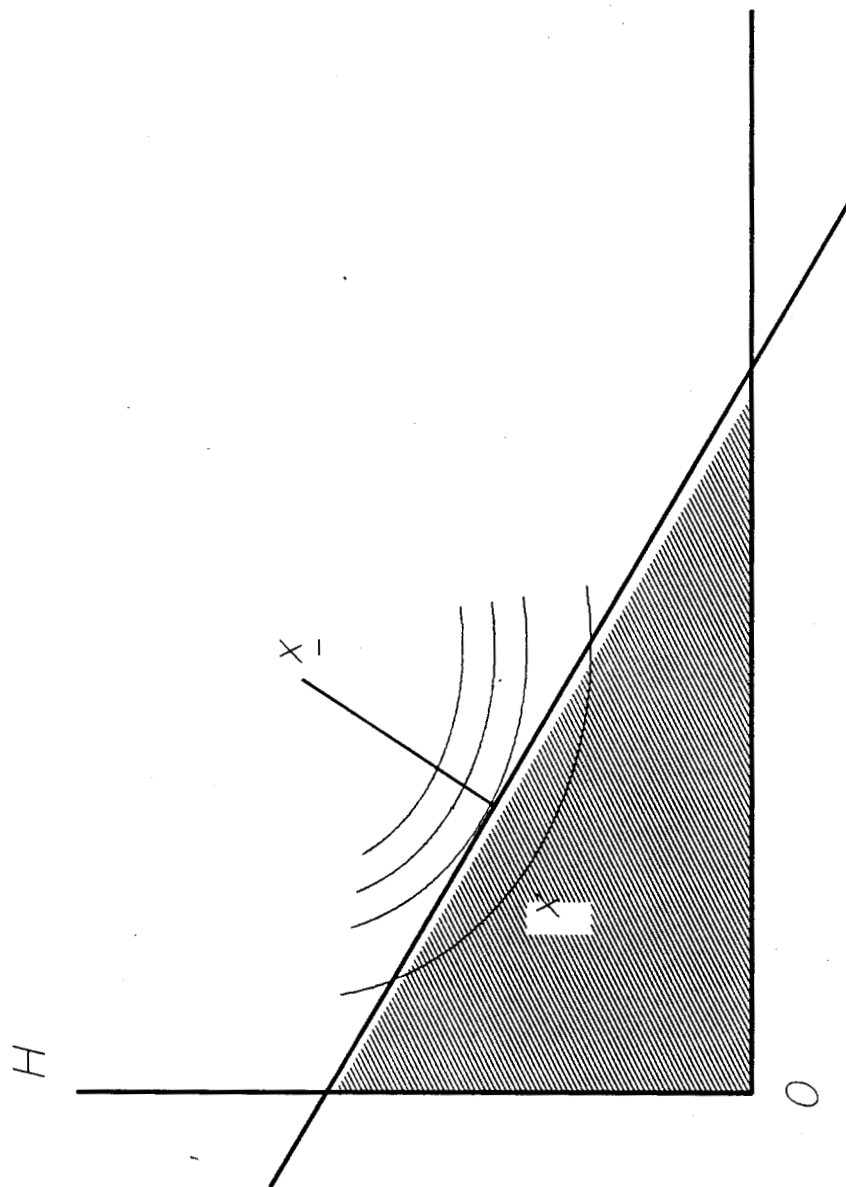
CAPITULO 6. EQUILIBRIO Y OPTIMOS DE PARETO.

Cuando se discutió el problema del análisis de actividades, se demostró que un punto en el conjunto de producción que maximiza el beneficio conforme a un determinado vector de precios, es un punto eficiente; la convexidad en el conjunto de producción juega un papel crucial en este caso. En un mercado competitivo, los productores son incapaces de alterar los precios de mercado, y en consecuencia, deben maximizar su beneficio para el vector de precios prevaleciente. Esto nos sugiere la existencia de una relación relativamente fuerte entre precios competitivos y puntos eficientes. Pero antes de profundizar en esta relación, debemos analizar cuál es el comportamiento, en este sentido, del consumidor.

En el mercado competitivo, también el consumidor debe aceptar los precios prevalecientes, ya que en función de estos, maximizará su satisfacción dentro de su conjunto de consumo. Asumiendo que productores y consumidores son las dos unidades básicas de la economía, y que la oferta total de bienes es igual a la demanda total, una cuestión natural surge al preguntarse si la economía permite, en estas circunstancias, un cierto óptimo social. Al introducir a los consumidores en la discusión, nos damos cuenta de que el beneficio individual de un productor no puede ser tomado como un indicador de *bienestar social*. La búsqueda del bienestar social nos lleva, de manera casi directa, al concepto de *óptimo de Pareto*, que es el estado en donde nadie puede estar mejor sin afectar a los demás (mientras la oferta total de bienes sea igual a la demanda total). Como se presume que la satisfacción de un

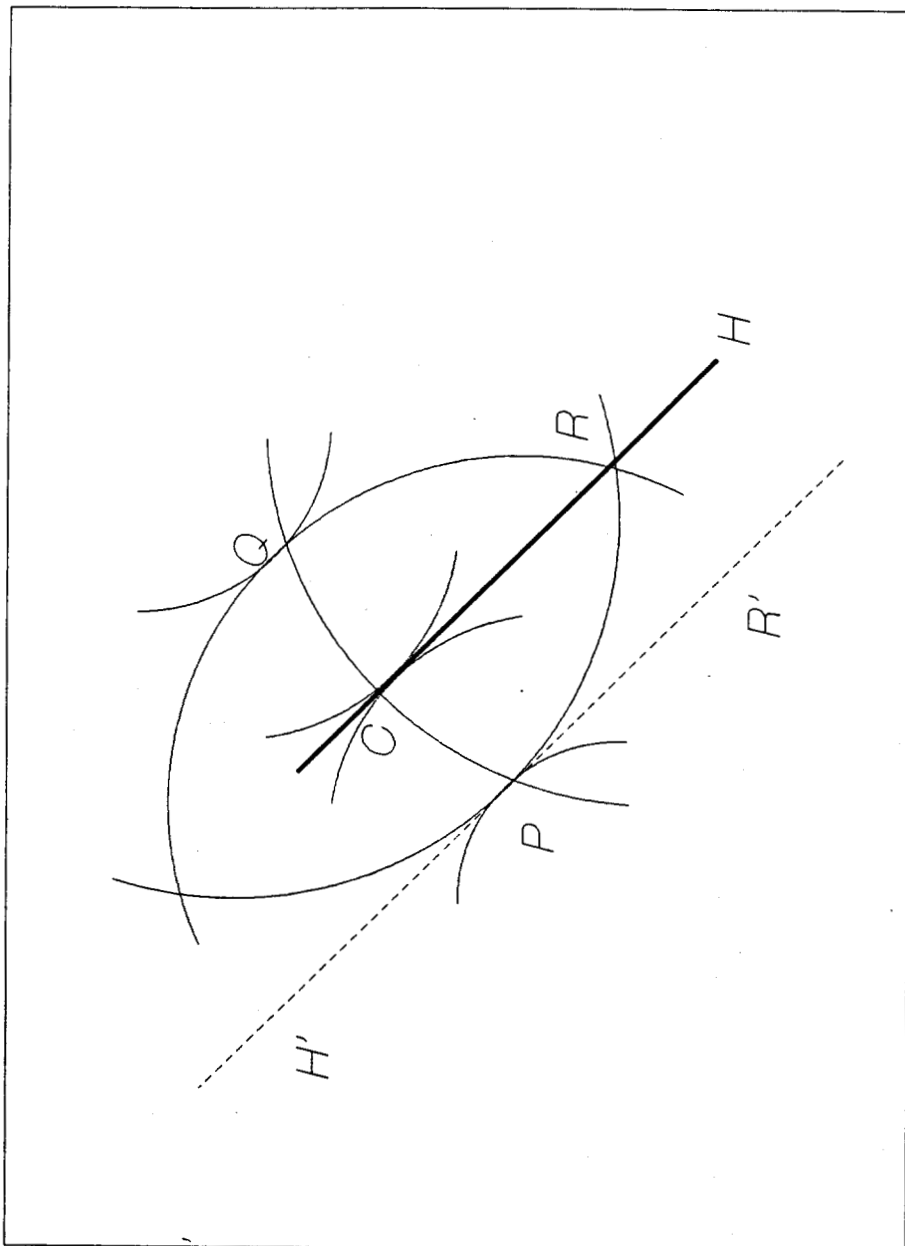
consumidor se deriva del ejercicio de su actividad como tal, los términos "mejor que" y "peor que" se refieren solamente al orden de preferencias individual.

La siguiente figura ilustra el problema de la elección de un consumidor, cuyo conjunto de consumo se ubica en el cuadrante no negativo de $\mathbb{R}^2$. Si se encuentra un vector de precios (en H) donde el precio de cada bien es positivo, se puede observar que el consumidor debe maximizar su satisfacción eligiendo algún punto dentro de la región sombreada.



Considérese ahora una economía en donde sólo participan dos consumidores (y ningún productor). Los consumidores intercambian bienes entre sí, y suponiendo que sólo existen dos bienes y que nuevamente el conjunto de consumo es el cuadrante no negativo de $\mathbb{R}^2$, podemos representar la situación conforme al siguiente diagrama

OB



OA

Cualquier distribución de bienes es posible mientras el punto que la representa se encuentre dentro del rectángulo dibujado. El punto R , por ejemplo, no constituye un óptimo de Pareto, debido a que es posible incrementar el bienestar de un individuo sin afectar el del otro; para ver esto, imagínese un movimiento dentro del "lente" formado por las curvas de indiferencia individuales que pasan por R . Por otro lado, cualquier punto sobre la curva PQ es un óptimo de Pareto. El punto C representará, por su parte, un equilibrio competitivo, siempre que los precios indicados por H prevalezcan (cada individuo maximiza su satisfacción en el sentido descrito anteriormente). Nótese que en el punto C , la línea de precios es tangente a cada una de las curvas de indiferencia. De hecho, esta condición es suficiente para garantizar la maximización.

Como puede apreciarse en la figura, los puntos que representan un óptimo de Pareto (en este caso P), son aquellos en donde las curvas de indiferencia individuales son tangentes entre sí; a la colección de puntos de este tipo se le llama curva contractual. Adviértase que en cada óptimo de Pareto se puede dibujar una línea tangente a la curva de indiferencia de cada consumidor. En el punto C , H es dicha línea, y en el punto P lo es H' .

Regresando al tratamiento formal, sea x_i un vector n -dimensional de consumo y y_j un vector n -dimensional de producción. Los elementos negativos de y_j denotan insumos y los positivos productos. Sean

$$x = \sum x_i \quad y \quad y = \sum y_j,$$

y sean

$$X = \sum X_i \quad y \quad Y = \sum Y_j.$$

Dado un precio p , hemos dicho que el beneficio del j -ésimo productor es $p \cdot y_j$. Se presupone la existencia de una dotación inicial de bienes para cada consumidor, que denotaremos por α_i , y análogamente

$$\alpha = \sum \alpha_i.$$

Definiremos ahora Factibilidad, Óptimo de Pareto, y Equilibrio Competitivo.

Definición: Un conjunto de vectores de consumo $\{x_i\}$ es *factible*, si existe un correspondiente conjunto de vectores de producción $\{y_j\}$ tal que $x = y + \alpha$.

Definición: Un $\{x_i\}$ factible es un *óptimo de Pareto (o.p.)* si no existe un $\{x'_i\}$ factible tal que $x'_i (\geq) x_i$ para todo $i=1, \dots, m$ con $x'_i (>) x_i$ para al menos un i .

Definición: Un conjunto de vectores $[p, \{\underline{x}_i\}, \{\underline{y}_j\}]$ es un *equilibrio competitivo (e.c.)*, si $\underline{x}_i \in X$ para $i=1, 2, \dots, m$, $\underline{y}_j \in Y$ para $j = 1, 2, \dots, k$, y se cumplen

(i) $\underline{x}_i (\geq) x_i$ para todo $x_i \in X_i$ tal que $p \cdot x_i \leq p \cdot \underline{x}_i$, para $i = 1, 2, \dots, m$.

(ii) $p \cdot \underline{y}_j \geq p \cdot y_j$ para todo y_j en Y_j ,
con $j = 1, 2, \dots, k$.

(iii) $\underline{x} = \underline{y} + \alpha$.

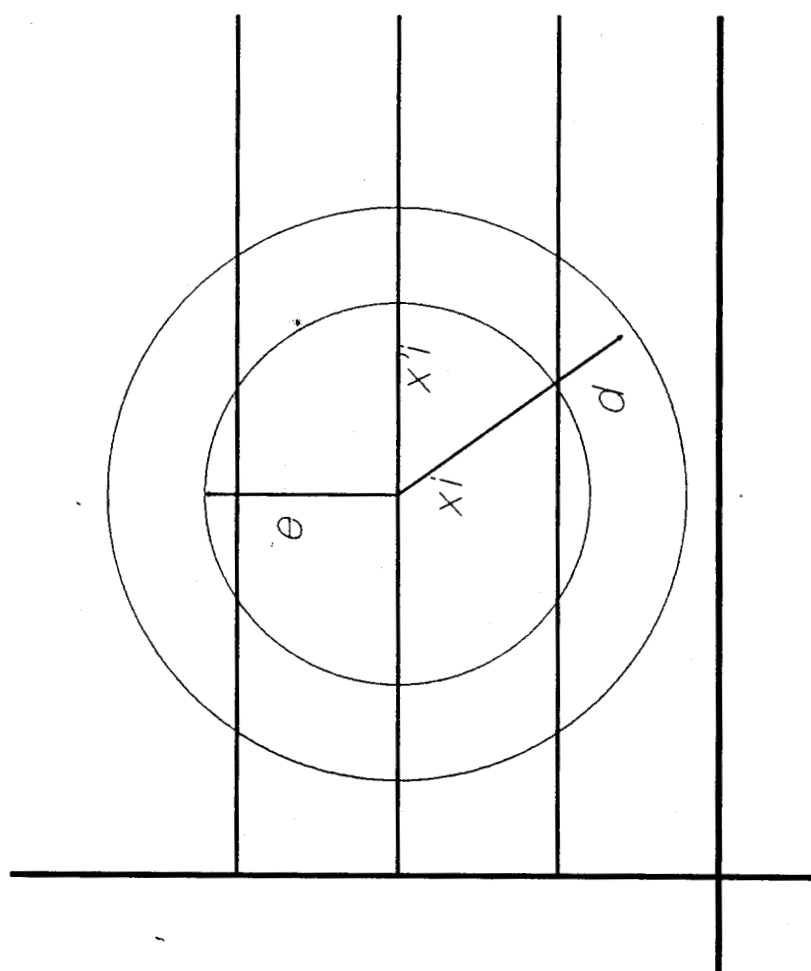
Definición: Para un precio p dado, un punto $\underline{x}_i \in X_i$ será llamado una *elección* del i -ésimo consumidor, si

$$\underline{x}_i (\geq) x_i \text{ para todo } x_i \in X_i \text{ con } p \cdot x_i \leq p \cdot \underline{x}_i$$

Definición: Un punto $x_i \in X_i$ se conoce como *punto de insaciabilidad local*, si existe un $\delta > 0$ con $B_\delta(x_i) \cap (X_i \setminus x_i) \neq \emptyset$, tal que para todo ε , $0 < \varepsilon < \delta$, con $B_\varepsilon(x_i) \cap (X_i \setminus x_i) \neq \emptyset$,

tengamos que $x'_1 (>) x_1$ para algún $x'_1 \in B_\varepsilon(x_1) \cap X_1$, donde $B_\delta(x_1)$ y $B_\varepsilon(x_1)$ son bolas abiertas con centro en x_1 y radio δ y ε respectivamente.

Obsérvese que la definición anterior supone la existencia de al menos un bien divisible (ya que ε puede asumir cualquier valor entre 0 y δ), y la no saciedad del consumidor con respecto a este bien (es decir, existe un punto tal que $x'_1 (>) x_1$ para cada ε). En la siguiente figura se ilustra el caso en el que un bien es perfectamente divisible y otro es indivisible. Aquí, el conjunto de consumo está constituido por la colección de líneas horizontales del cuadrante no negativo.



INSACIABILIDAD LOCAL

Definición: El orden de preferencias ($\succeq$) se llama *localmente insaturado* si, dado un punto de insaciabilidad local x_1 , $x'_1 (=) x_1$ implica que x'_1 es también un punto de insaciabilidad local.

Introduciremos ahora el siguiente supuesto

(a-1) El orden de preferencias ($\succeq$) es localmente insaturado para cada consumidor (suponiendo que existe al menos un bien perfectamente divisible).

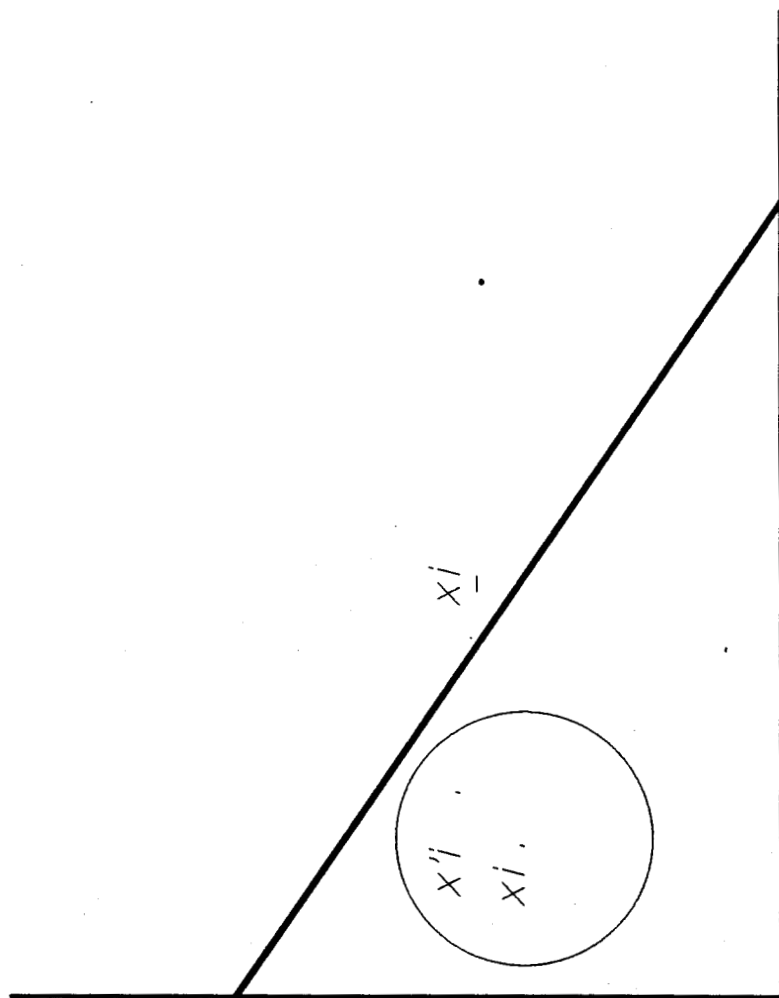
Lema: Sea $\underline{x}_1$ un punto de insaciabilidad local elegido para el i -ésimo consumidor a un precio p dado. Entonces, bajo (a-1),

- (i) $x_1 (>) \underline{x}_1$ implica que $p \cdot x_1 > p \cdot \underline{x}_1$.
- (ii) $x_1 (=) \underline{x}_1$ implica que $p \cdot x_1 \geq p \cdot \underline{x}_1$.

Prueba:

- (i) Supóngase lo contrario; esto es, $p \cdot x_1 \leq p \cdot \underline{x}_1$. Ya que $\underline{x}_1$ es un punto elegido, $\underline{x}_1 (\succeq) x_1$, que es una contradicción.
- (ii) Sea $x_1 (=) \underline{x}_1$ y supóngase que $p \cdot x_1 < p \cdot \underline{x}_1$. Ya que $\underline{x}_1$ no es un punto de saciedad local, tampoco lo es x_1 (por a-1). De aquí que para todo ε , $0 < \varepsilon < \delta$, existe $x'_1 \in X_1$ y $x'_1 \in B_\varepsilon(x_1)$ tal que $x'_1 (>) x_1$, lo cual implica que $x_1 (>) \underline{x}_1$ por la transitividad del orden de preferencias. Podemos ahora elegir x'_1 suficientemente cerca de x_1 para que $p \cdot x'_1 < p \cdot \underline{x}_1$ (que es posible debido a que $p \cdot x_1$ es una función continua). Esto contradice la suposición de que $\underline{x}_1$ es un punto elegido bajo p .

La siguiente figura ilustra el punto (ii). Adviértase que la elección del x'_1 suficientemente cerca de x_1 para que $p \cdot x'_1 < p \cdot \underline{x}_1$ necesita la suposición de que al menos un bien es divisible.



Teorema 2-6: Sea $[p, \{x_i\}, \{y_j\}]$ un equilibrio competitivo tal que x_i sea un punto de insaciabilidad local para todo $i=1,2,\dots,m$. Suponga que se cumple (a-1) para todo i . Entonces, $[\{x_i\}, \{y_j\}]$ es un óptimo de Pareto.

Prueba: Supongamos que $[\{x_i\}, \{y_j\}]$ no es un óptimo de Pareto. Entonces existe $[\{x_i\}, \{y_j\}]$ tal que $x_i \in X_i$, $i=1,2,\dots,m$, $y_j \in Y_j$, $j=1,2,\dots,k$, y

- (i) $x = y + \alpha$
- (ii) $x_i (\geq) \underline{x}_i$ para todo $i=1,2,\dots,m$
- (iii) $x_i (>) \underline{x}_i$ para algún i .

y del lema anterior tenemos que

$$\sum p \cdot x_i > \sum p \cdot \underline{x}_i \quad \text{o bien} \quad p \cdot x > p \cdot \underline{x}$$

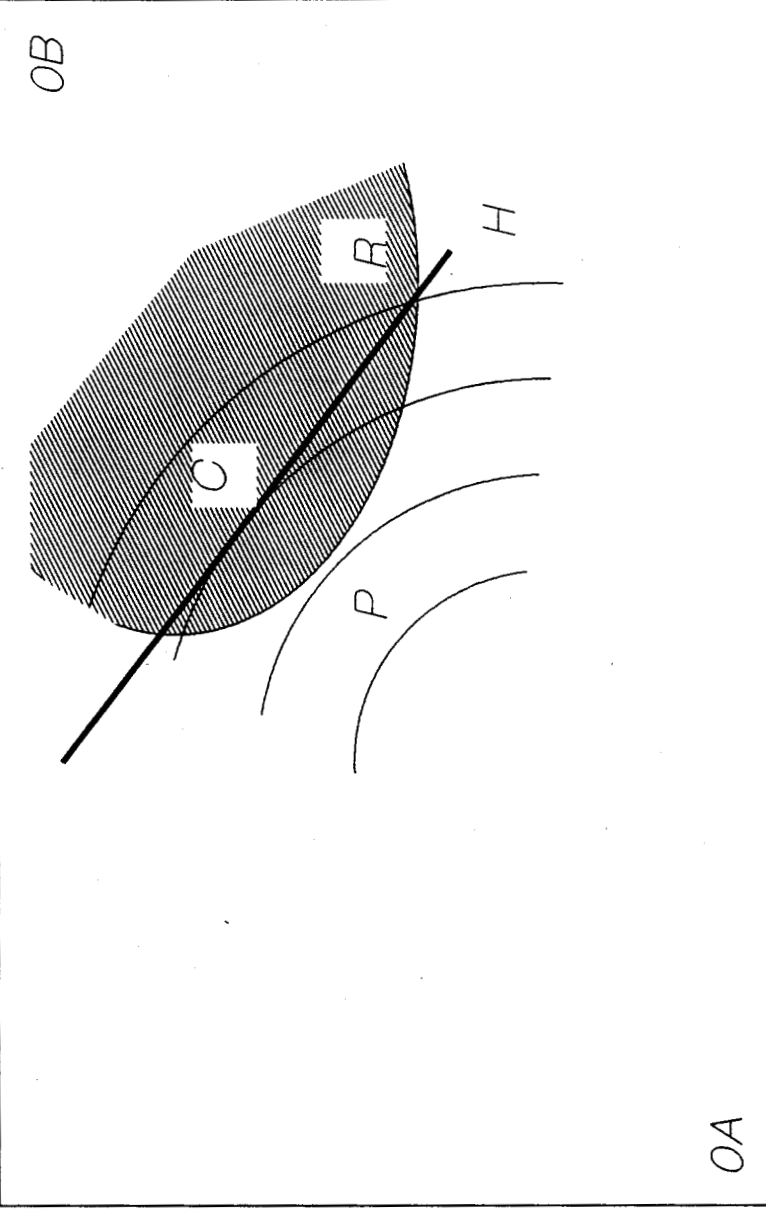
pero la condición (iii) de equilibrio competitivo requiere que $p \cdot \underline{x} = p \cdot \alpha$. De ahí tenemos que $p \cdot x > p \cdot y + p \cdot \alpha$. La condición (ii) de e.c. requiere que $p \cdot y_j \geq p \cdot z_j$ para todo $z_j \in Y_j$, $j = 1,2,\dots,k$. En particular, $p \cdot y_j \geq p \cdot y_j$, $j = 1,2,\dots,k$. O bien, sumando sobre j , $p \cdot y \geq p \cdot y$. Por lo tanto, tenemos ahora que

$$p \cdot x > p \cdot y + p \cdot \alpha$$

o bien $p \cdot (x - y - \alpha) > 0$. Esto contradice la factibilidad de $\{x_i\}$. Q.E.D.

Cabe hacer algunas observaciones. Primero, ya que x_i es un punto de insaciabilidad elegido, tenemos que $p \neq 0$. Por otro lado, obsérvese que en las demostraciones de el lema y el teorema anteriores, nunca se asumió convexidad para el conjunto de producción. Se puede construir un ejemplo en donde se vea que un

equilibrio competitivo no nos lleva a un óptimo de Pareto cuando (a-1) no se cumple. En el siguiente diagrama (Edgeworth-Bowley), se representa a una economía de intercambio puro entre dos agentes. El O_A representa el origen para el individuo A, y el O_B representa el origen para el individuo B. El punto de dotaciones iniciales se representa por R . Las curvas de indiferencia están dibujadas convexas hacia el origen. Nótese que el individuo A se sacia sobre y por arriba de su curva de indiferencia, y cada punto en la región de saciedad (en este caso sombreada) provee al individuo del un mismo nivel de satisfacción.



E.C. NO IMPLICA O.P.

El punto C es un equilibrio competitivo para el precio prevaleciente representado por H . Pero C no es un óptimo de Pareto, ya que el individuo B puede aumentar su satisfacción al desplazarse del punto C al punto P sin afectar al otro individuo. El punto P sí es un óptimo de Pareto.

En otro orden de ideas, hay que recordar que en la definición que se dio de equilibrio competitivo, nada se especificó acerca de cómo obtiene cada consumidor el ingreso que le permite comprar bienes con valor igual a $p \cdot x_i$. Una manera, es el caso típico en el que el consumidor recibe lo que vale su dotación inicial de recursos α_i (con $\alpha_i \in \mathbb{R}^n$ y $\sum \alpha_i = \alpha$) y reparte $\theta_{1i}, \theta_{2i}, \dots, \theta_{ki}$ de sus ingresos al 1o., 2o., ..., k -ésimo productor respectivamente (con $\theta_{ji} \in \mathbb{R}$, $\theta_{ji} > 0$ y $\sum \theta_{ji} = 1$). Las dotaciones iniciales α_i son cantidades dadas de bienes que el consumidor posee *a priori*, y θ_{ji} se interpreta como la porción de la j -ésima empresa que es propiedad del i -ésimo consumidor.

Esto describe lo que podríamos llamar *economía de propiedad privada*. Podemos ahora adecuar la definición de equilibrio competitivo al caso particular de una economía de propiedad privada.

Definición: Un conjunto de vectores $[p, \{x_i\}, \{y_j\}, \{\theta_{ji}\}]$ es un *equilibrio competitivo de una economía de propiedad privada (e.c.e.p.p.)* si $x_i \in X_i$ para $i=1, 2, \dots, m$, $y_j \in Y_j$ para $j = 1, 2, \dots, k$, y

- (i) $\bar{x}_i (\geq) x_i$ para todo $x_i \in X_i$ tal que $p \cdot x_i \leq M_i$, donde $M_i = p \cdot \alpha_i + \sum_{j=1}^k \theta_{ji} p \cdot y_j$, para $i = 1, 2, \dots, m$.
- (ii) $p \cdot \bar{y}_j \geq p \cdot y_j$ para $y_j \in Y_j$, con $j = 1, 2, \dots, k$.
- (iii) $\bar{x} = \bar{y} + \alpha$

Se puede ver que si $[p, \{x_i\}, \{y_j\}, \{\theta_{ji}\}]$ es un e.c.e.p.p. entonces es un e.c. Para verificar que todo e.c. puede derivarse de algun

e.c.e.p.p., considérese lo siguiente. Al i -ésimo consumidor, otórguesele una dotación inicial de recursos $\alpha_i = \underline{x}_i - (1/m)\underline{y}$ y tómese $\theta_{ji} = 1/m$, y observe que $\sum \theta_{ji} = 1$, $\alpha + \underline{y} = \underline{x}$, etc., donde $[p, \{\underline{x}_i\}, \{\underline{y}_i\}]$ es un e.c.

Examinemos ahora un teorema más, relacionado con T2-6. Dado que la economía se encuentre en un estado óptimo de Pareto, queremos saber si existe o no un vector de precios tal, que pueda ser usado para sostener un equilibrio competitivo (permitiendo, si es necesario, la redistribución de los recursos).

La respuesta es afirmativa bajo un conjunto de suposiciones más fuertes que las hechas para establecer T2-6. Se mostrará a continuación, que la herramienta fundamental para establecer el nuevo teorema es el *teorema de separación*, y que la pendiente del hiperplano de separación nos dará el vector de precios buscado.

Definición: Llamamos al orden de preferencias ($\succeq$) en X_i *convexo* si

- (i) El conjunto α_i es convexo
- (ii) $x (>) x'$ implica que $tx + (1 - t)x' (>) x$, para $0 < t < 1$.
- (iii) $x (=) x'$ implica que $tx + (1 - t)x' (\succeq) x'$ para $0 < t < 1$.

Aquí, la convexidad de α_i es necesaria para que los enunciados (ii) e (iii) tengan significado. La convexidad de X_i presupone la divisibilidad perfecta de todos los bienes.

Definición: Sean

$$C_i(\underline{x}_i) = \{x_i \in X_i \mid x_i (\succeq) \underline{x}_i\}$$

$$\bar{C}_i(\underline{x}_i) = \{x_i \in X_i \mid x_i (>) \underline{x}_i\}$$

La convexidad en el orden de preferencias implica la convexidad para $C_i(\underline{x}_i)$ y $\bar{C}_i(\underline{x}_i)$ para todo $\underline{x}_i \in X_i$.

Ahora introducimos los siguientes supuestos:

- (a-2) El orden de preferencias es convexo para todo $i = 1, 2, \dots, m$.
- (a-3) El conjunto Y es convexo

(a-4) Dados un punto $\underline{x}_1$ y un vector de precios prevalecientes $\underline{p}$, existe $x'_1 \in X_1$ tal que $\underline{p} \cdot x'_1 < \underline{p} \cdot \underline{x}_1$.

(a-5) Para cada $i = 1, 2, \dots, m$, el conjunto $\{x_1 \in X_1 \mid x_1 (\leq) x'_1\}$ es cerrado para todo $x'_1 \in X_1$ (es decir, si $\{x_1^q\}$ es una sucesión en X_1 tal que $x_1^q (\leq) x'_1$ y $x_1^q \rightarrow x_1^0$, entonces tenemos que $x_1^0 (\leq) x'_1$).

Adviértase que (a-3) no requiere que los conjuntos de producción individuales Y_j sean convexos. Al supuesto (a-4) se le conoce como supuesto de mínima riqueza.

Definición: Se dice que el i -ésimo consumidor se encuentra no saciado en $\underline{x}_1$, si existe un $x_1 \in X_1$ tal que $x_1 (>) \underline{x}_1$.

Teorema 2-6: Supóngase que $\{[\underline{x}_1], [\underline{y}_j]\}$ es un óptimo de Pareto tal que existe al menos un consumidor que se encuentra no saciado. Entonces, bajo las suposiciones (a-2) y (a-3), existe un vector de precios $\underline{p} \neq 0$ tal que

- (i) $\underline{p} \cdot \underline{x}_1 \leq \underline{p} \cdot x_1$ para todo $x_1 \in X_1$ con $x_1 (\geq) \underline{x}_1$, $i = 1, 2, \dots, m$.
- (ii) $\underline{p} \cdot \underline{y}_j \geq \underline{p} \cdot y_j$ para todo $y_j \in Y_j$, $j = 1, 2, \dots, k$.
- (iii) $\underline{x} = \underline{y} + \alpha$.

Obsérvese que si bien este teorema no pide que se cumplan (a-4) y (a-5), tampoco nos asegura que para todo óptimo de Pareto podamos disponer de un vector de precios tal que se pueda usar en un e.c. La condición (i) establece que cada consumidor minimiza su gasto sobre el conjunto no peor que $\underline{x}_1$, pero esto no implica necesariamente que maximice su satisfacción sobre su conjunto presupuestal. Para probar lo anterior, usaremos (a-4) y (a-5).

Ahora demostramos T2-6:

Prueba:

Sin pérdida de generalidad, podemos suponer que el primer consumidor se encuentra en un estado de no saciedad para $\underline{x}_1$. Sea

$$\bar{C}(\underline{x}_1, \underline{x}_2, \dots, \underline{x}_m) = \bar{C}_1(\underline{x}_1) + \sum_{i=2}^m C_1(\underline{x}_i)$$

Por simplicidad de notación, abreviaremos $\bar{C}(\underline{x}_1, \underline{x}_2, \dots, \underline{x}_m)$ por $\bar{C}$. Por (a-2), $\bar{C}$ es convexo. Sea $W = \{w \mid w = y + \alpha, y \in Y\}$. Ya que Y es convexo por (a-3), W es también convexo. Por la definición de óptimo de Pareto, $z \in \bar{C}$ implica que $z \notin W$. De ahí que $\bar{C}$ y W sean dos conjuntos ajenos no vacíos. Por el teorema de separación de Minkowski, existen un $p \neq 0$ y un número real α tales que

$$(a) \quad p \cdot w \leq \alpha \quad \text{para todo } w \in W$$

y

$$(b) \quad p \cdot z \geq \alpha \quad \text{para todo } z \in \bar{C}$$

Probaremos ahora que $p \cdot \underline{x} = \alpha$.

Como $\underline{x} = \underline{y} + \alpha$ y $\underline{y} \in Y$, tenemos que $\underline{x} \in W$, así que $p \cdot \underline{x} \leq \alpha$.

Supóngase que $x' \in \bar{C}$; entonces, podemos encontrar un $x'_i \in \bar{C}_1(\underline{x}_i)$ y un $x'_i \in \bar{C}_1(\underline{x}_i)$, $i = 1, 2, \dots, m$, tales que

$$x' = \sum_{i=1}^m x'_i$$

Ahora sea $x_i(t) = tx'_i + (1-t)\underline{x}_i$, $0 < t < 1$, $i = 1, 2, \dots, m$, y

$$x(t) = \sum_{i=1}^m x_i(t)$$

Por (a-2), $x_1(t) \in \bar{C}_1(\underline{x}_1)$ y $x_i(t) \in \bar{C}_1(\underline{x}_i)$, $i = 2, 3, \dots, m$.

Entonces $x(t) \in \bar{C}$. Supóngase que $p \cdot \underline{x} < \alpha$. Por (b), obtenemos que $p \cdot z > p \cdot \underline{x}$ para todo $z \in \bar{C}$. En particular, $p \cdot x(t) > p \cdot \underline{x}$ para $0 < t < 1$. Ya que escogiendo un t suficientemente pequeño, $x(t)$ puede estar arbitrariamente cerca de $\underline{x}$, encontramos una contradicción. Por lo tanto tenemos que $p \cdot \underline{x}$ no es menor que α . Esto, junto con

$p \cdot \underline{x} \leq \alpha$ nos da que $p \cdot \underline{x} = \alpha$ (esto es, el hiperplano separador de los conjuntos W y $\bar{C}$, pasa por $\underline{x}$). Entonces (a) y (b) pueden reescribirse de la siguiente manera

$$(a') \quad p \cdot w \leq p \cdot \underline{x} \quad \text{para todo } w \in W$$

y

$$(b') \quad p \cdot z \geq p \cdot \underline{x} \quad \text{para todo } z \in \bar{C}$$

El hecho de que $w \in W$ significa que w puede escribirse como

$$w = \sum_{j=1}^k y_j + \alpha$$

donde $y_j \in Y_j$, $j = 1, 2, \dots, k$.

Por lo tanto, de (a') obtenemos

$$p \cdot [\sum y_j + \alpha] \leq p \cdot \underline{x}, \quad \text{para todo } y_j \in Y_j, \quad j = 1, 2, \dots, k$$

pero $\underline{x} = \underline{y} + \alpha$ (factibilidad de o.p.). Entonces

$$p \cdot [\sum y_j + \alpha] \leq p \cdot [\sum \underline{y}_j + \alpha] \quad \text{para todo } y_j \in Y_j, \quad j = 1, 2, \dots, k$$

o bien

$$\sum_{j=1}^k p \cdot \underline{y}_j \geq \sum_{j=1}^k p \cdot y_j$$

Tómese $j = j_0$ fijo y sea $y_j = \underline{y}_j$ para todo $j \neq j_0$. Entonces $p \cdot \underline{y}_{j_0} \geq p \cdot y_{j_0}$ para todo $y_{j_0} \in Y_{j_0}$. Ya que la elección de j_0 es arbitraria, esto prueba la parte (ii) del teorema. La condición (iii) se satisface por la factibilidad del o.p.

Similarmemente, de (b') obtenemos

$$(c) \quad p \cdot x_1 \geq p \cdot \underline{x}_1 \quad \text{para todo } x_1 \in \bar{C}_1(\underline{x}_1)$$

y

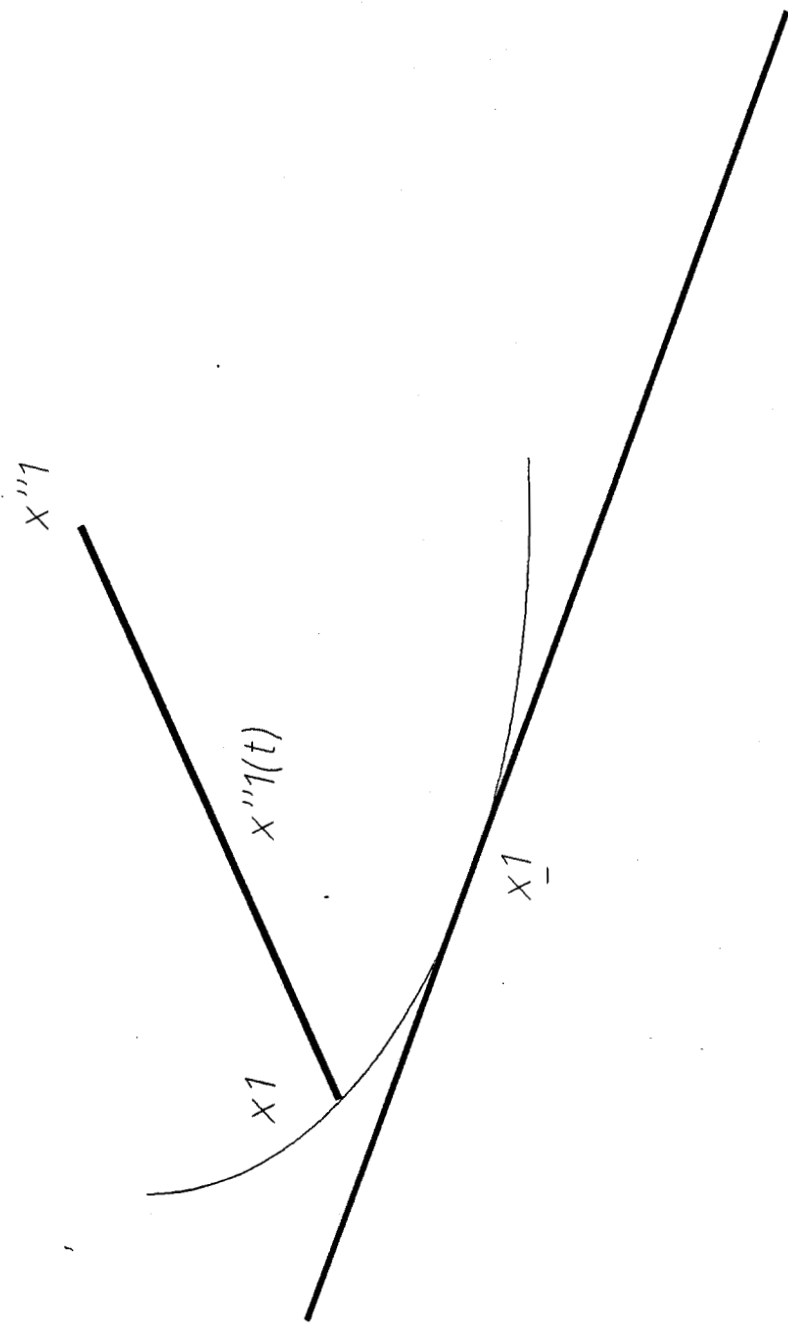
$$(d) \quad p \cdot x_1 \geq p \cdot \underline{x}_1 \quad \text{para todo } x_1 \in C_1(\underline{x}_1), \quad i = 1, 2, \dots, m$$

Queremos llegar a que $p \cdot x_1 \geq p \cdot \underline{x}_1$ para todo $x_1 \in C_1(\underline{x}_1)$. Para ello, tenemos que mostrar que $x_1 (=) \underline{x}_1$ implica que $p \cdot x_1 \geq p \cdot \underline{x}_1$. Ya que $\underline{x}_1$ no es un punto de saciedad, podemos encontrar un punto $x''_1 (>) x_1 (=) \underline{x}_1$. Sea $x''_1(t) = tx''_1 + (1 - t)x_1$, $0 < t < 1$. Entonces, por convexidad de $(\geq)$ (a-2), $x''_1(t) (>) x_1$, $0 < t < 1$, así que $x''_1(t) (>) \underline{x}_1$. Entonces, de (c), $p \cdot x''_1(t) \geq p \cdot \underline{x}_1$. Luego, por continuidad de la función $p \cdot x_1$, se sigue que $p \cdot x_1 \geq p \cdot \underline{x}_1$. De este modo obtenemos

$$(e) \quad p \cdot x_1 \geq p \cdot \underline{x}_1 \quad \text{para todo } x_1 \in C_1(\underline{x}_1), \quad i = 1, 2, \dots, m$$

Esto prueba la condición (i) del teorema.

Q.E.D.



Una ilustración de $x''(t)$ y $x'(t)$

Corolario: Si se cumple (a-4) con respecto a $\underline{x}_i$ y a $\underline{p}$ en el teorema anterior, y además se cumple (a-5), entonces para todo óptimo de Pareto $[\{x_i\}, \{y_i\}]$, existe un $\underline{p} \neq 0$ tal que $[\underline{p}, \{x_i\}, \{y_i\}]$ es un equilibrio competitivo.

Prueba: Es suficiente mostrar que la condición (i) de e.c. se cumple. En el teorema anterior obtuvimos que

$$(e) \quad \underline{p} \cdot x_i \geq \underline{p} \cdot \underline{x}_i \quad \text{para todo } x_i \in C_i(\underline{x}_i), \quad i = 1, 2, \dots, m$$

Lo que no excluye la posibilidad de que exista un $x_i \in X_i$ tal que $\underline{p} \cdot x_i = \underline{p} \cdot \underline{x}_i$ y $x_i (>) \underline{x}_i$. Mostraremos que esto no puede suceder bajo (a-4) y (a-5). Supongamos que existe un $x_i \in X_i$ tal que

$$\underline{p} \cdot x_i = \underline{p} \cdot \underline{x}_i \quad \text{y} \quad x_i (>) \underline{x}_i$$

De (a-4) tenemos que existe un $x'_i \in X_i$ tal que $\underline{p} \cdot x'_i < \underline{p} \cdot \underline{x}_i$. La relación (e) implica que $\underline{x}'_i (<) \underline{x}_i$. Ahora considérese $z_i = tx'_i + (1 - t)x_i$, $0 < t < 1$. Se puede ver que $\underline{p} \cdot z_i < \underline{p} \cdot x_i = \underline{p} \cdot \underline{x}_i$. Eligiendo un t suficientemente pequeño, podemos hacer que z_i se encuentre arbitrariamente cerca de x_i . Entonces, de (a-5) obtenemos

$$z_i (>) \underline{x}_i, \quad \text{pero} \quad \underline{p} \cdot z_i < \underline{p} \cdot \underline{x}_i$$

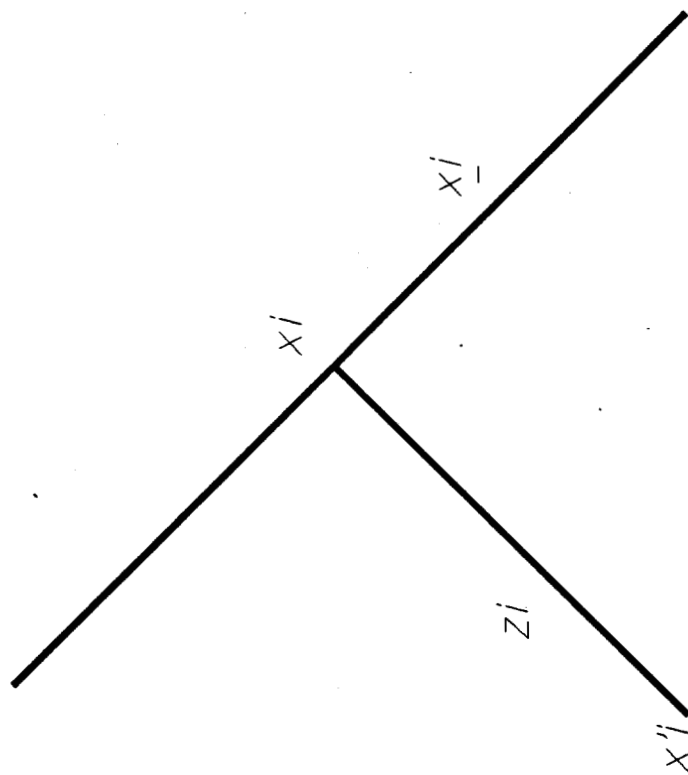


Ilustración de la prueba

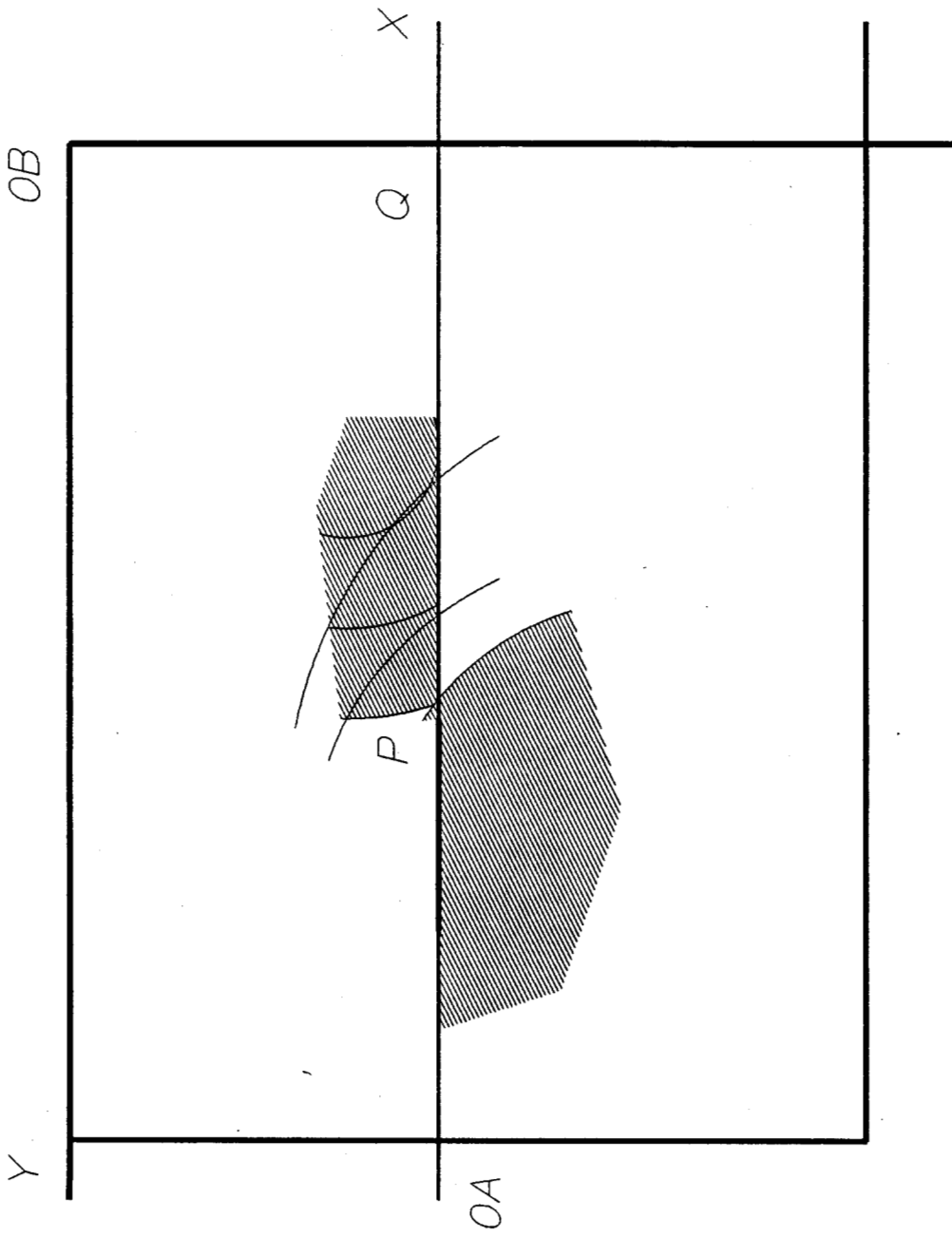
Esto contradice la relación (e). Luego no puede existir un $x_1 \in X_1$ tal que $p \cdot x_1 = p \cdot \underline{x}_1$ y $x_1 (>) \underline{x}_1$. Esto significa que $p \cdot x_1 = p \cdot \underline{x}_1$ implica que $x_1 (\leq) \underline{x}_1$. Nótese que la relación (e) significa (tomando la contraposición) que $p \cdot x_1 < p \cdot \underline{x}_1$ implica que $x_1 (<) \underline{x}_1$. Luego hemos obtenido que

$$p \cdot x_1 \leq p \cdot \underline{x}_1 \text{ implica que } x_1 (\leq) \underline{x}_1$$

En otros términos, $\underline{x}_1 (\geq) x_1$ para todo $x_1 \in X_1$ con $p \cdot x_1 \leq p \cdot \underline{x}_1$. Esto prueba la condición (i) de e.c. Q.E.D.

Adviértase que el equilibrio competitivo en el corolario anterior puede alcanzarse distribuyendo de el ingreso agregado $p \cdot (y + \alpha)$ la cantidad $p \cdot \underline{x}_i$, $i = 1, 2, \dots, m$ a cada consumidor. En otros términos, sin una redistribución de la propiedad, un óptimo de Pareto puede no estar sustentado, en general, por un sistema de precios competitivos.

Como se apuntó anteriormente, el teorema 2-V4 no establece precisamente que un óptimo de Pareto nos pueda llevar a un equilibrio competitivo. Para el corolario anterior, tuvimos que hechar mano de las suposiciones adicionales (a-4) y (a-5). Un ejemplo en donde se muestra que la conclusión del corolario no es válida cuando (a-4) no se cumple, ha sido elaborado por Arrow y se muestra esquematizado en el siguiente diagrama:



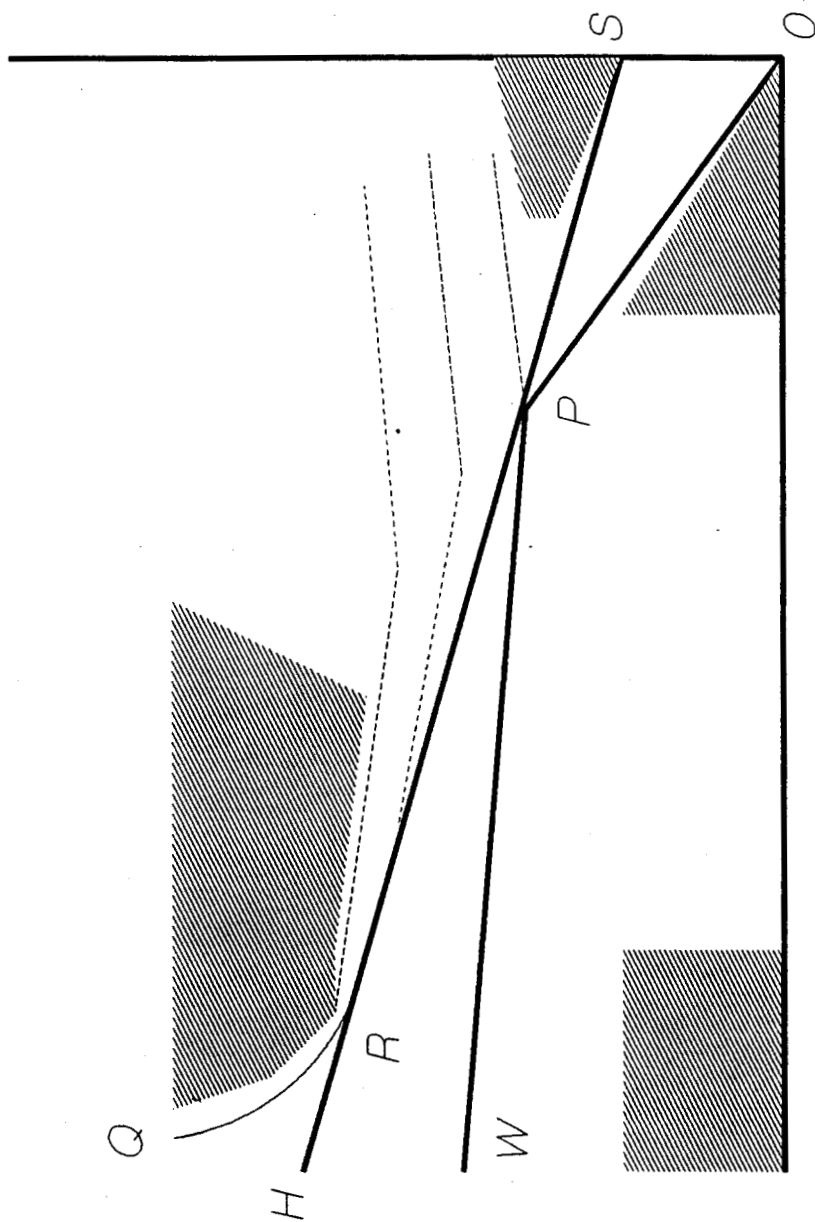
Anomalia de Arrow

El conjunto de consumo del agente A, se extiende hacia la parte superior derecha del punto O_A . El conjunto de consumo del agente B se extiende hacia la parte inferior izquierda del punto O_B . No hay producción en esta economía.

El punto P, es un óptimo de Pareto. La línea tangente a las curvas de indiferencia de ambos consumidores pasa por O_A, P, Q (el contorno superior para cada consumidor en el punto P es la región sombreada). Pero esta recta no puede representar una línea de precios que sustente un equilibrio competitivo en P. Con esta línea de precios, el precio del bien X es cero, y el consumidor A no está maximizando su satisfacción sujeto a su presupuesto en P. Puede incrementar su satisfacción moviéndose hacia donde está Q (lo cual es posible dado que puede obtener gratis cantidades del bien X). Nótese que en el punto P, el valor de su plan de consumo es cero y no tiene punto alguno de posible consumo que se encuentre por debajo de $OAPQ$. En otras palabras, el supuesto (a-4) no se satisface. Nótese también que todos los demás supuesto si se cumplen.

Este ejemplo no considera la presencia de producción en la economía. Un ejemplo en donde se introduce esta actividad ha sido elaborado por Koopmans. Trata de una economía en donde sólo intervienen un productor y un consumidor. Al haber sólo un consumidor, un óptimo de Pareto será el punto en donde este agente alcanza su máximo dada la condición de factibilidad, esto es, dada la oferta total de bienes en la economía (recordemos que denotamos la oferta agregada con W en la prueba del teorema 2-V4). Sólo existen dos bienes, que son alimento y trabajo. El conjunto de consumo X es la región que se encuentra sobre la curva QRPS. Las curvas de indiferencia del consumidor están representadas por las líneas punteadas. Supóngase que la curva de indiferencia que pasa por el punto P es tangente a la línea PR (pero termina en P sin seguir sobre PR). Cualquier punto sobre la línea RP (excepto P) es mejor que P para el consumidor. El punto P es el único óptimo de

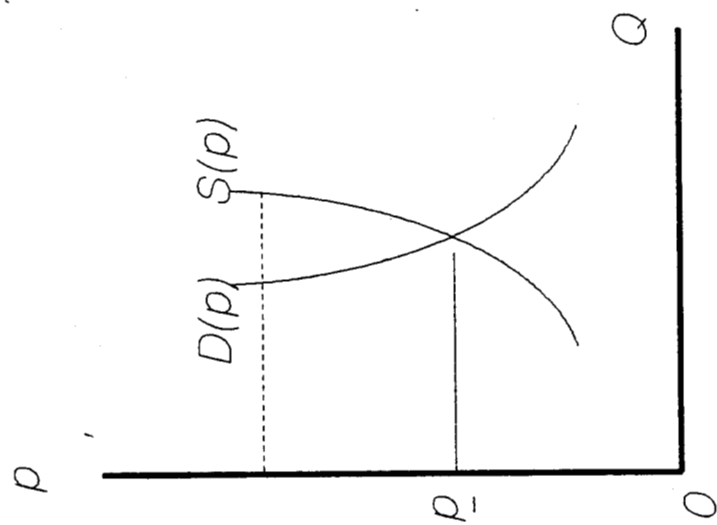
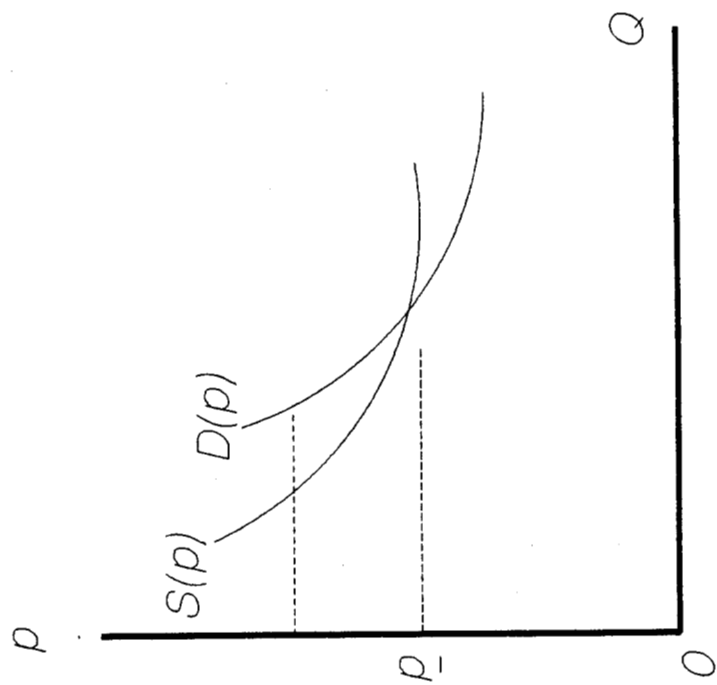
Pareto, debido a que es el único punto factible (el conjunto X y W tienen exclusivamente intersección en P).



La línea H (que pasa por PR) es la única que separa los conjuntos X y W . (Recordemos que en las pruebas del teorema 2-V4 y de su corolario la pendiente del hiperplano separador de X y W nos da el vector de precios que sustenta un equilibrio competitivo). Pero esta línea contiene un punto (digamos R) en X . En otras palabras, el punto R es mejor que P pero tienen el mismo valor. Luego, si el precio representado por H prevalece el consumidor incrementará su satisfacción cambiando al productor su paquete de bienes representado por P por el que representa R . No existe punto alguno de X por debajo de H . Entonces, el supuesto (a-4) se viola nuevamente.

CAPITULO 7. ESTABILIDAD Y CONDICIONES DE SUSTITUABILIDAD.

Considérese el mercado competitivo aislado de un bien (digamos el bien A). Supóngase que la demanda de A, denotada por $D(p)$, es una función de su precio p , y supóngase también que la oferta de A, denotada por $S(p)$, es una función de su precio. Un precio de equilibrio, $\bar{p}$, es un precio tal que $D(\bar{p}) = S(\bar{p})$. La cuestión de si existe o no este $\bar{p}$, es justamente el problema de *existencia del equilibrio*. Esta existencia quedará garantizada cuando las curvas de oferta y de demanda se crucen en algún punto. Si se cruzan en dos o más puntos se dirá que existen dos o más equilibrios, según sea el caso. Supóngase que el precio de A, p , se aleja de cierto precio de equilibrio (digamos $\bar{p}$). Cabe preguntarse qué sucede con la trayectoria de p a través del tiempo. Esto nos lleva al problema de *estabilidad*. Para poder tratarlo, introduciremos una hipótesis fundamental: una excesiva demanda de A producirá un aumento en su precio, p , y una excesiva oferta producirá una reducción de su precio. Más formalmente, diremos que $dp/dt \geq 0$ si $D(p) - S(p) \geq 0$ y $dp/dt \leq 0$ si $D(p) - S(p) \leq 0$. Para hacer una descripción paulatina del problema, por lo pronto pensaremos que existe un precio de equilibrio único tal que $S(\bar{p}) = D(\bar{p})$. Puesto de esta forma, la cuestión de la estabilidad puede visualizarse con cierta facilidad a partir del análisis gráfico. La primera de las siguientes dos figuras ilustra el caso de un equilibrio estable, mientras que en la segunda se puede ver el caso de un equilibrio inestable. Imagínese la dirección hacia donde se movería p cuando se encuentra distante del equilibrio $\bar{p}$ en cada diagrama.



Para considerar un problema un poco más general, supóngase que cierta economía se puede describir a través de un sistema de n ecuaciones en donde aparecen las *relaciones de equilibrio* de la economía. Sea $x = (x_1, x_2, \dots, x_n)$ el conjunto de variables que se deben determinar por el sistema de n ecuaciones. Sean $f_i(x) = 0$, $i = 1, 2, \dots, n$, las relaciones de equilibrio de la economía. El valor de x que satisface el sistema anterior (suponiendo que existe solución) será llamado *valor de equilibrio* de x , y se denotará por $\underline{x}$. El problema de la estabilidad se refiere a lo que puede pasar con la trayectoria de x a través del tiempo, cuando x se desvía de $\underline{x}$. Particularmente, queremos indagar si x converge a $\underline{x}$ con el tiempo. Un ejemplo de sistema de equilibrio es el sistema de mercados competitivos, en donde f_i denota el exceso de demanda para el i -ésimo bien, y x_i denota el precio del mismo. Supóngase ahora que el valor inicial de x está dado por x_0 . Asumiremos que x_0 no es un valor de equilibrio, es decir, que $f(x_0) \neq 0$. Supóngase además que esta situación genera un cierto proceso de ajuste a partir del cual, la trayectoria de x con respecto al tiempo, denotada por $x(t)$, queda descrita por el siguiente sistema de ecuaciones:

$$F_i[x(t), t] = 0, \quad i = 1, 2, \dots, n$$

con $x(0) = x_0$. Este sistema de ecuaciones puede verse como el generado por un sistema de ecuaciones diferenciales que describe los procesos de ajuste. El análisis de estabilidad tendrá que ver entonces con la solución del sistema anterior y con la posible convergencia de $x(t)$ a un valor de equilibrio $\underline{x}$. Concretamente, el sistema de ecuaciones diferenciales que genera a las F_i puede escribirse como

$$\frac{dx_i(t)}{dt} = h_i[f_i(x(t))].$$

En el caso de un equilibrio competitivo, esto significa que el movimiento del i -ésimo precio, $x_i(t)$, es una función de la demanda excedente del i -ésimo bien (f_i).

Considérese ahora un equilibrio competitivo que quede descrito por el siguiente sistema de ecuaciones:

$$f_i(p_1, p_2, \dots, p_n) = 0, \quad i = 1, 2, \dots, n \quad [\text{o bien } f(p) = 0]$$

donde p_i denota el precio de el i -ésimo bien y f_i denota la demanda excedente para el mismo. Como se apuntó anteriormente, la hipótesis fundamental en el análisis de estabilidad para mercados competitivos es que el exceso de demanda del i -ésimo bien hace que el precio del mismo aumente, mientras que el exceso de oferta del i -ésimo bien origina una disminución de su precio.

Para tratar el caso de un mercado con muchos bienes, distinguiremos dos tipos de "estabilidad". Un equilibrio en el mercado del j -ésimo bien se llamará *imperfectamente estable* si los mercados de todos los demás bienes se encuentran en equilibrio y existe estabilidad en el j -ésimo mercado. Sea $\underline{p}$ un vector de precios de equilibrio [esto es, $f(p) = 0$, suponiendo que $\underline{p}$ existe], y supóngase que el vector de precios p dista de este $\underline{p}$. Entonces el equilibrio del j -ésimo mercado se llamará *imperfectamente estable* si

$$(i) \quad f_i(p) = 0 \text{ para toda } i \neq j$$

y

$$(ii) \quad p_j > \underline{p}_j \text{ implica que } f_j(p) < 0 \quad \text{y} \quad p_j < \underline{p}_j \text{ implica que } f_j(p) > 0$$

El equilibrio de j -ésimo mercado será *perfectamente estable* si se cumplen las condiciones de estabilidad imperfecta sin importar el número de otros mercados ajustados al equilibrio, o en otros términos, sin que importe si se ajustan o no otros precios para mantener el equilibrio en el mercado en cuestión [esto es, para i

$\neq j$, $f_i(p) = 0$ o bien p_i es constante]. Si el equilibrio en cada mercado es imperfectamente estable, entonces el sistema es imperfectamente estable, o bien, si el equilibrio en cada mercado es perfectamente estable, entonces el sistema es perfectamente estable.

Derivando el sistema $f(p) = 0$ en algún punto de equilibrio (digamos p) con respecto a p_j y aplicando las definiciones de estabilidad perfecta y estabilidad imperfecta, se obtiene la siguiente condición para estabilidad perfecta del equilibrio del sistema (después de repetir para $j = 1, 2, \dots, n$)

$$a_{ij} < 0, \begin{vmatrix} a_{ii} & a_{ij} \\ a_{ji} & a_{jj} \end{vmatrix} > 0, \begin{vmatrix} a_{ii} & a_{ij} & a_{ik} \\ a_{ji} & a_{jj} & a_{jk} \\ a_{ki} & a_{kj} & a_{kk} \end{vmatrix} < 0, \dots$$

para toda $i, j, k, \dots$, de el conjunto de índices $\{1, 2, \dots, n\}$, siendo $a_{ij} = \partial f_i / \partial p_j$, evaluada en p , con $i, j = 1, 2, \dots, n$. Cuando una matriz $A = [a_{ij}]$ de dimensiones $n \times n$ satisface la condición anterior, se le conoce como matriz Hicksiana.

Los conceptos de estabilidad perfecta y estabilidad imperfecta pueden dar la impresión de poseer cierta artificialidad. Mediante un escrutinio más cuidadoso de los mismos, se puede ver que en realidad no tienen tanta relevancia para la descripción del problema de estabilidad que se está tratando aquí. En lugar de ahondar más en esta cuestión, pasaremos a visualizar el problema con un enfoque más reciente.

Consideremos ahora un modelo de economía competitiva en donde existen sólo tres bienes. Demostraremos estabilidad global utilizando los llamados *diagramas de fase*. Esta técnica se emplea con frecuencia para diversas aplicaciones en otras ramas de la economía como son Teoría del Crecimiento y Macroeconomía. La discusión que sigue servirá tanto para ilustrar el empleo de la técnica de diagramas de fase, como para probar estabilidad global en una economía de tres bienes.

Sea $f_i(p)$, donde $p = (p_1, p_2, p_3)$, la función de demanda excedente del i -ésimo bien. Considérese el modelo de equilibrio competitivo descrito por

$$f_i(p_1, p_2, p_3) = 0, \quad i = 1, 2, 3$$

o bien, más brevemente, $f_i(p) = 0$, $i = 1, 2, 3$. Asumiremos que

(a-1) (Sustituabilidad bruta) $\partial f_i / \partial p_j > 0$ para todos los valores de p , $i \neq j$, $i, j = 1, 2, 3$.

(a-2) (Homogeneidad) $f_i(p)$, $i = 1, 2, 3$, son homogéneas de grado cero

(a-3) (Ley de Walras) $\sum p_i f_i(p) = 0$ para todo p .

(a-4) $p_i > 0$, $i = 1, 2, 3$.

(a-5) Existe al menos un p tal que $f_i(p) = 0$, para todo i .

Bajo el supuesto (a-2), eligiéremos normalizar el vector de precios p de modo tal que siempre tengamos $p_3 = 1$.

El proceso de ajuste dinámico queda descrito por

$$\dot{p}_i = k_i f_i(p_1, p_2, p_3), \quad k_i > 0, \quad i = 1, 2.$$

Nótese que $f_1 = 0$ y $f_2 = 0$ implican que $f_3 = 0$, debido a la Ley de Walras (es decir, si los primeros dos mercados llegan al equilibrio, el tercer mercado queda también en equilibrio). Por lo

tanto es suficiente considerar el proceso de ajuste para los primeros dos mercados para hacer el análisis de estabilidad.

El problema radica en encontrar si el sistema anterior de ecuaciones diferenciales [que escribiremos como $p(t; p_0)$], converge al vector de precios de equilibrio p , con p definido por $f_i(p) = 0$, $i = 1, 2, 3$. A través de la técnica de diagramas de fase, podemos examinar la trayectoria con respecto al tiempo de $p(t; p_0)$ sin resolver explícitamente las ecuaciones diferenciales. Esta técnica se basa (para el presente caso) en el hecho de que cada curva $[p_i = 0]$ o bien $[f_i(p) = 0]$ con $i = 1, 2$, divide el plano (p_1-p_2) en dos regiones: la región en donde $p_i > 0$, y la región en donde $p_i < 0$ (recuérdese que $p_3 = 1$ siempre). Podemos omitir la consideración de p_3 ó p_3 . Esto nos permitirá tratar el problema en el plano de dos dimensiones.

Primero veamos la forma de las curvas $[f_i(p) = 0]$. Ambas tienen pendientes positivas y se intersectan sólo una vez [de donde el punto de equilibrio, es decir, el punto en donde $f_i(p) = 0$ para toda i , si existe, es único]. Más aún, podemos asegurar que la curva $f_2(p) = 0$ intersecta a la curva $f_1(p) = 0$ "por su lado izquierdo". Examinando los signos de p_1 y de p_2 en las cuatro regiones que definen estas dos curvas, podremos asegurar la estabilidad global de p . Procederemos a hacer esta descripción con más detalle.

Obsérvese en primer término que la siguiente ecuación se cumple bajo el supuesto de homogeneidad de grado cero (a-2) [denotaremos en adelante $\partial f_i / \partial p_j$ por f_{ij}]:

$$\sum_j f_{ij} p_j = 0 \quad \text{para todo } p, \text{ con } i = 1, 2, 3.$$

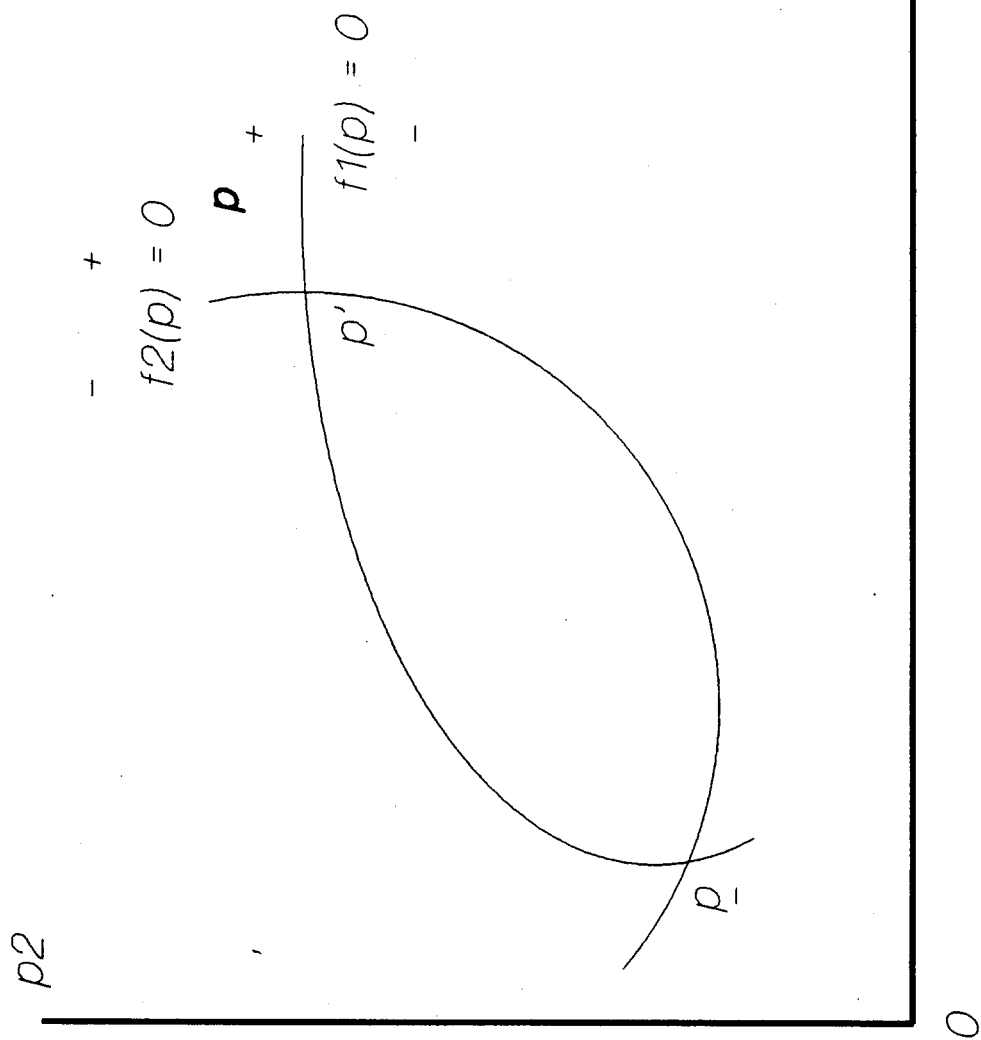
Y en vista de (a-1) y (a-4), obtenemos que

$$f_{ii} < 0 \quad \text{para todo } p, \text{ con } i = 1, 2, 3.$$

Considérense los valores de p_1 y p_2 para los cuales $f_1 = 0$. Esto define una curva en el plano (p_1-p_2) . Queremos obtener la pendiente de esta curva, es decir, dp_2/dp_1 . Esto se puede hacer derivando $f_1 = 0$, es decir $f_{11}dp_1 + f_{12}dp_2 = 0$. De este modo, $dp_2/dp_1 = -f_{11}/f_{12}$, que es una cantidad positiva debido al hecho de que $f_{11} < 0$ para toda i , bajo el supuesto (a-1). Análogamente, obtenemos la pendiente de la curva definida por $f_2 = 0$ en el plano (p_1-p_2) derivando $f_2 = 0$, esto es, $f_{21}dp_1 + f_{22}dp_2 = 0$. Entonces tenemos que $dp_2/dp_1 = -f_{21}/f_{22}$, que es nuevamente una cantidad positiva. Por lo tanto ambas curvas $[f_1(p) = 0]$ y $[f_2(p) = 0]$ tienen pendiente positiva.

En seguida examinaremos los signos de p_1 en la región definida por estas dos curvas. Como $f_{12} > 0$ para todo p debido a (a-1), del lado izquierdo de la curva $[f_1(p) = 0]$, $p_1 > 0$. Del lado derecho de la curva $[f_1(p) = 0]$, $p_1 < 0$. Como $f_{21} > 0$ para todo p , tendremos también que del lado izquierdo de $[f_2(p) = 0]$, $p_2 < 0$, y del lado derecho $p_2 > 0$. A los mismos resultados llegaríamos partiendo de que $f_{11} < 0$ y $f_{22} < 0$.

Por (a-5), existe al menos un punto de equilibrio. Así que las dos curvas $[f_1(p) = 0]$ y $[f_2(p) = 0]$ se intersectan por lo menos una vez. Veremos ahora que precisamente esta intersección es única, es decir, que sólo existe un punto de equilibrio. Para demostrarlo, supóngase lo contrario, es decir que hay un punto p' tal que $p' > 0$, $f(p') = 0$ y que $p' \neq p$, en donde $p'_3 = p_3 = 1$. Entonces las curvas $[f_1(p) = 0]$ y $[f_2(p) = 0]$ se intersectan por lo menos dos veces, justamente en los puntos p' y p . Por lo tanto, en alguno de esos puntos, digamos en p' , la curva $[f_1(p) = 0]$ debe tocar a la curva $[f_2(p) = 0]$ por el lado izquierdo. Esto se ilustra en la siguiente figura.



Unicidad del Equilibrio

Ahora considérese un punto p de modo tal que $p_1 > p'_1$ y $p_2 > p'_2$ (con $p_3 = p'_3 = p_3 = 1$). Entonces $f_1(p) > 0$ y $f_2(p) > 0$. Derivando $f_3(p)$ obtenemos que

$$df_3 = f_{31}dp_1 + f_{32}dp_2$$

Compárense ahora los puntos p' y p . Como $f_3(p') = 0$, la ecuación anterior implica que $f_3(p) > 0$ ($f_{31} > 0$, $f_{32} > 0$, $dp_1 > 0$, $dp_2 > 0$, por lo tanto $df_3 > 0$), así que $f_i(p) > 0$ para toda i . Esto, a la luz de (a-4), contradice la Ley de Walras (a-3). De aquí que podamos establecer la unicidad del punto de equilibrio, es decir, que la curva $[f_2(p) = 0]$ intersecta solamente por la izquierda a la curva $[f_1(p) = 0]$ (siendo p el punto de intersección).

Podemos ahora dibujar el diagrama fase que se ilustra en la siguiente figura, y en donde se puede observar la estabilidad global de p . Aparecen ahí las trayectorias con respecto al tiempo del vector de precios (p_1, p_2) correspondientes a dos valores iniciales distintos (p^0 y p^0).

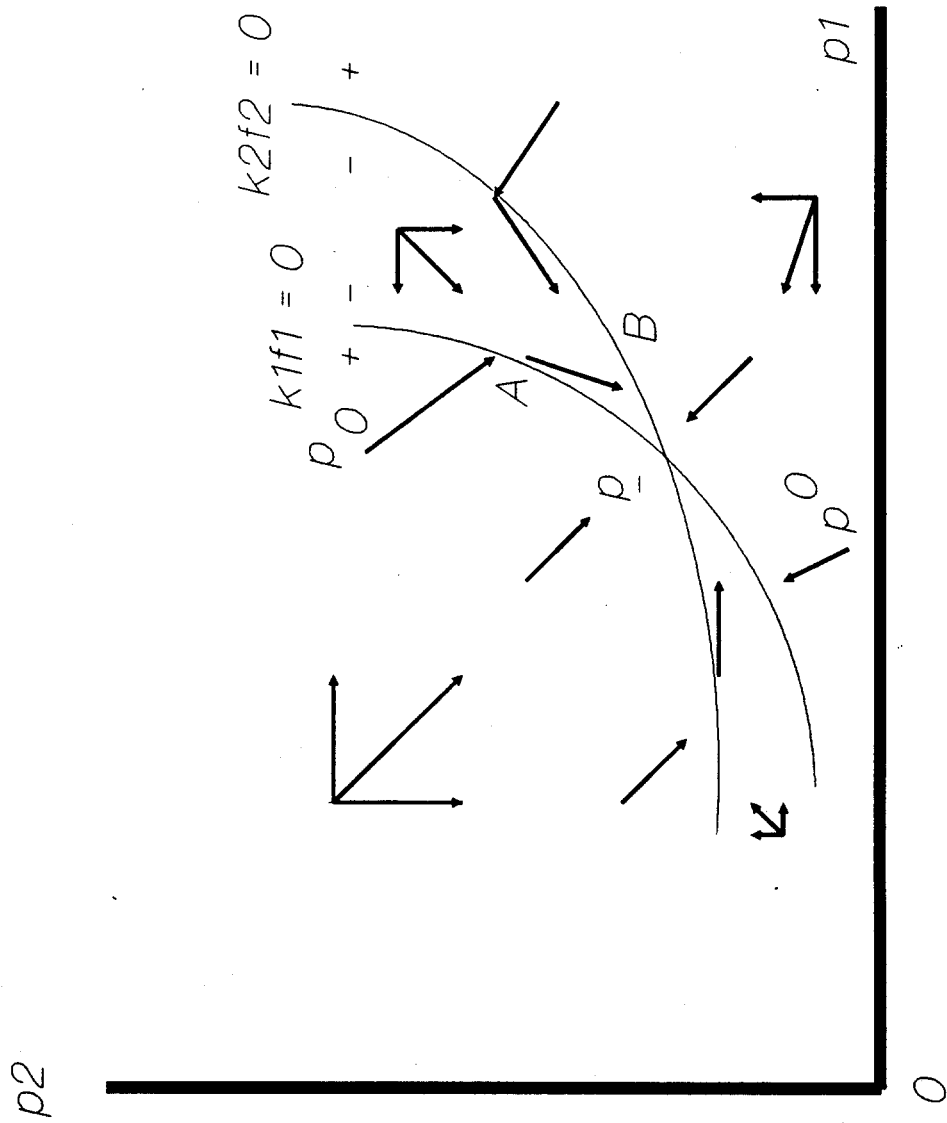


Diagrama Fase

Por ejemplo, considérese el punto p_0 de la figura. En p_0 , $f_1 > 0$ y $f_2 < 0$ así que $\dot{p}_1 > 0$ y $\dot{p}_2 < 0$. Es decir, p_1 crece mientras p_2 decrece en el tiempo. La trayectoria de $[p_1(t), p_2(t)]$ eventualmente toca a la curva $[f_1 = 0]$ en el punto A, en donde f_2 es todavía negativo. Por lo tanto, en el punto A, $\dot{p}_1 = 0$ y $\dot{p}_2 < 0$, de modo tal que la trayectoria se sigue por la región en donde $f_1 < 0$ y $f_2 < 0$, donde tanto p_1 como p_2 decrecen con el tiempo. En la terminología de ecuaciones diferenciales, se califica con frecuencia a puntos de equilibrio de este tipo como *nodos impropios*.

Por lo anterior, debe quedar claro el importante papel que juega la hipótesis de sustituibilidad bruta en el análisis de estabilidad. Se ha establecido estabilidad global bajo las hipótesis de homogeneidad de grado cero, Ley de Walrass y sustituibilidad bruta.

Resulta natural cuestionar qué tanto se podrán relajar estas suposiciones. Particularmente, examinaremos el caso de sustituibilidad bruta, ya que recientemente se han hecho intentos de cambiar esta hipótesis por alguna más plausible sobre la función de utilidad (como semiconcavidad, es decir, convexidad con respecto al origen de las curvas de indiferencia).

Considérese para el efecto, una economía de intercambio puro que comprenda a tres consumidores y tres bienes. Sea x_{ij} el consumo del bien j ($j = 1, 2, 3$) que lleva a cabo el i -ésimo agente, $i = 1, 2, 3$. Supóngase que las funciones de utilidad, u_i , queda descrita de la siguiente manera:

$$u_1(x_{11}, x_{12}, x_{13}) = \min \{x_{11}, x_{12}\}$$

$$u_2(x_{21}, x_{22}, x_{23}) = \min \{x_{22}, x_{23}\}$$

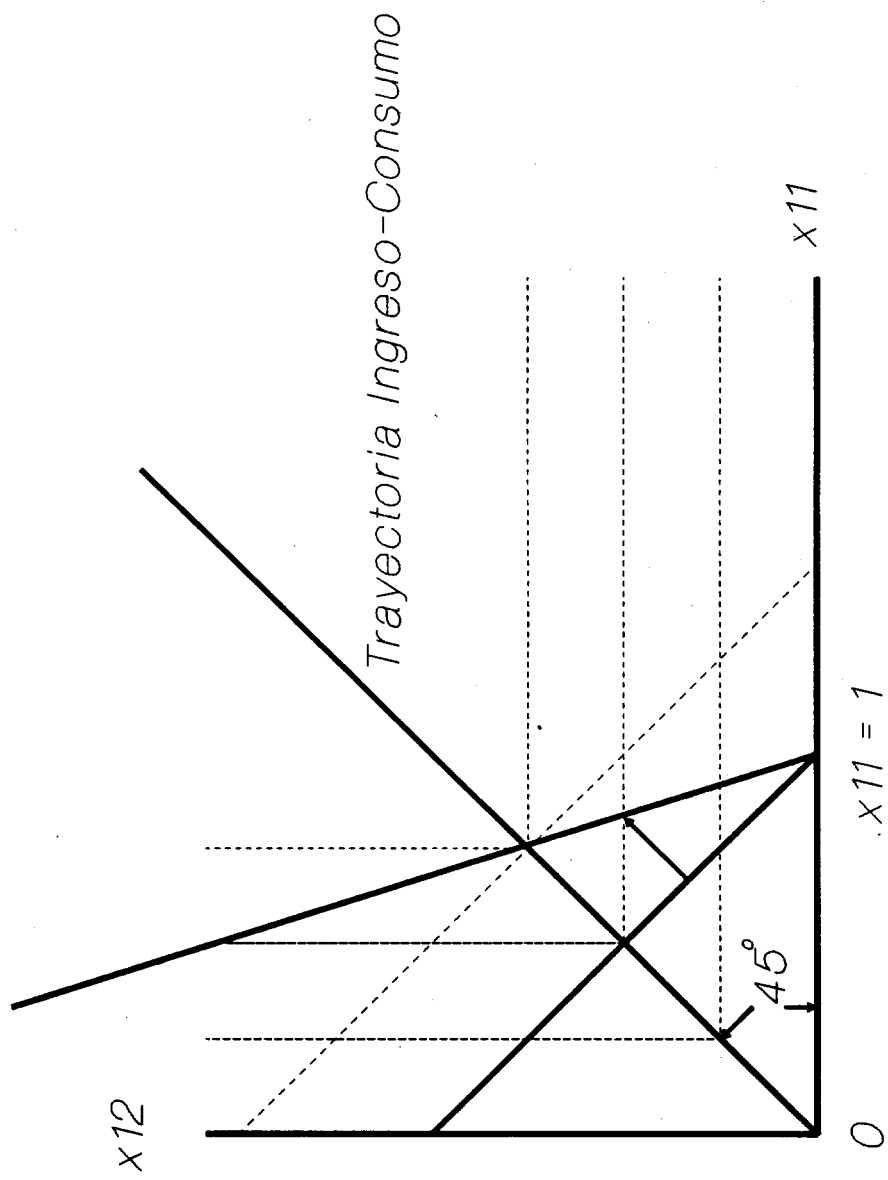
$$u_3(x_{31}, x_{32}, x_{33}) = \min \{x_{31}, x_{33}\}$$

Es decir, cada individuo desea sólo dos bienes en proporciones fijas (uno a uno). Puede verse que dicha función de utilidad nos

lleva a unas curvas de indiferencia convexas con respecto al origen. Sea α_{ij} la dotación inicial del bien j que posee el agente i . Supóngase que

$$\alpha_{ii} = 1, \quad \text{y} \quad \alpha_{ij} = 0, \quad i \neq j$$

Esto es, que el agente i sólo posee una unidad del bien i , y ninguna de los otros bienes. La ilustración de la curva de indiferencia y la línea del presupuesto aparecen en la siguiente figura



Adviértase que, dadas las especificaciones de las funciones de utilidad, la curva de ingreso-consumo de cada individuo en la figura anterior es una línea recta con pendiente de 45°, de modo que $x_{11} = x_{12}$, $x_{22} = x_{23}$, y $x_{31} = x_{33}$. Considérese el cambio en el precio que se indica por la flecha dibujada en el diagrama (un decremento en el precio del bien 2). Podrá observarse que prevalece solamente un "efecto de ingreso", mientras que no hay efecto de sustitución alguno. Por lo tanto, para estas curvas de indiferencia, el cambio en el precio es totalmente absorbido por el efecto de ingreso.

La demanda excedente para el bien 1 puede escribirse como

$$\begin{aligned} x_1 - \alpha_1 &= (x_{11} + x_{21} + x_{31}) - (\alpha_{11} + \alpha_{21} + \alpha_{31}) \\ &= (x_{11} + x_{31}) - \alpha_{11} = (x_{11} - \alpha_{11}) + x_{31} \end{aligned}$$

La ecuación de presupuesto para el agente 1 puede escribirse como $p_1 x_{11} + p_2 x_{22} = p_1 \alpha_{11}$. Pero, debido a nuestra convención, $x_{11} = x_{12}$ y $\alpha_{11} = 1$. Entonces $x_{11} = p_1 / (p_1 + p_2)$, así que $x_{11} - \alpha_{11} = -p_2 / (p_1 + p_2)$. Similarmente, x_{31} puede obtenerse de la ecuación de presupuesto del agente 3, $p_1 x_{31} + p_3 x_{33} = p_3 \alpha_{33}$, esto es, $x_{31} = p_3 / (p_3 + p_1)$. Por lo tanto, de la ecuación de demanda excedente tenemos que

$$(a) \quad x_1 - \alpha_1 = \frac{p_2}{p_1 + p_2} + \frac{p_3}{p_3 + p_1}$$

Análogamente obtenemos que

$$(b) \quad x_2 - \alpha_2 = \frac{p_3}{p_2 + p_3} + \frac{p_1}{p_1 + p_2}$$

$$(c) \quad x_3 - \alpha_3 = \frac{p_1}{p_3 + p_1} + \frac{p_2}{p_2 + p_3}$$

De donde se puede ver que existe un único "rayo de precios de equilibrio" $p_1 = p_2 = p_3$.

Escribimos ahora la ecuación de ajuste dinámico como

$$\dot{p}_1(t) = x_1(t) - \alpha_1.$$

Ahora queremos mostrar que $|p(t)| = \text{constante}$ para todo t . Para hacerlo, derivamos $|p(t)|$ con respecto a t . Es decir,

$$(d) \quad \frac{d}{dt} \left[\sum_{i=1}^3 p_i^2(t) \right] = 2 \sum_{i=1}^3 p_i(t) \dot{p}_i = 2 \sum_{i=1}^3 p_i (x_i - \alpha_i) = 0$$

Por lo tanto, concluimos que $|p(t)|^2 = \sum_{i=1}^3 p_i^2(t) = \text{constante}$.

En seguida, mostraremos que $\prod_{i=1}^3 p_i(t) = \text{constante}$. Para hacerlo, derivamos de la siguiente manera

$$\begin{aligned} \frac{d}{dt} \left[\prod_{i=1}^3 p_i(t) \right] &= \dot{p}_1 p_2 p_3 + p_2 \dot{p}_3 p_1 + p_3 p_1 \dot{p}_2 \\ &= (x_1 - \alpha_1) p_2 p_3 + (x_2 - \alpha_2) p_3 p_1 + (x_3 - \alpha_3) p_1 p_2 = 0 \end{aligned}$$

En donde la última igualdad se obtiene utilizando las ecuaciones (a), (b) y (c).

Por último probaremos que el proceso dinámico (d) no es globalmente estable. Para esto, elíjanse los precios iniciales $p_i(0)$, $i = 1, 2, 3$, de manera que $\sum_{i=1}^3 p_i^2(0) = 3$ y $\prod_{i=1}^3 p_i(0) \neq 1$.

Entonces $\sum_{i=1}^3 p_i^2(t) = 3$ y $\prod_{i=1}^3 p_i(t) \neq 1$ para todo t . Como

$\sum_{i=1}^3 p_i^2(0) = 3$, los únicos precios de equilibrio son $p_1 = p_2 = p_3 = 1$. De aquí que la solución del sistema de ecuaciones diferenciales (d) no puede converger al precio de equilibrio p donde $p_1 = p_2 = p_3 = 1$.

BIBLIOGRAFIA.

- [1] Debreu, G., *Theory of Value*, New York, Wiley, 1979.
- [2] Berge, C., *Topological Spaces*, New York, Macmillan, 1963.
- [3] Fenchel, W., *Convex Cones, Sets and Functions*, Princeton, N. J., Princeton University Press, 1973.
- [4] Fleming, W. H., *Functions of several variables*, Adison - Wesley, 1975.
- [5] Halmos, P. R., *Finite Dimensional Vector Spaces*, Princeton, N. J., 1978.
- [6] Kelley, J. L., *General Topology*, New York, Van Nostrand, 1975.
- [7] Nikaido, H., *Introduction to Sets and Mappings in Modern Economics*, North - Holland, 1980.
- [8] Rudin, W., *Principles of Mathematical Analysis*, New York, Mc Graw - Hill, 1974.
- [9] Simmons, G. F., *Introduction to Topology and Modern Analysis*, New York, Mc Graw - Hill, 1963.
- [10] Valentine, F. A., *Convex Sets*, New York, Mc Graw - Hill, 1974.

- [11] Hestenes, M. R., *Optimization Theory*, New York, Wiley, 1985.
- [12] Quirk, J. P., *Comparative statics under Walras' Law: The case of strong dependence*, Review of Economics Studies, 35, 11 - 21.
- [13] Flaschel, P., *Classical and Neoclassical Competitive adjustment Processes*, New School of Social Research, 1986.
- [14] *Notes on Price and Quantity Tâtonnement Dynamics*, Models of Economic Dynamics, Heidelberg, Springer, 49 - 68, 1986.
- [15] Semmler, W., *Competition, Monopoly and Differential Profit Rates*, New York, Columbia University Press, 1989.
- [16] Sonnenschein, H., *Price Dynamics and the Disappearance of Short - Run Profits*, Journal of Mathematical Economics, Vol 8, No 2, 201 - 204, 1989.
- [17] Svensson, L. E. O., *Walrasian and Marshallian Stability*, Journal of Economic Theory, vol. 34, No. 2, 371 - 379, 1984.
- [18] Aoki, M., *Dual Stability in a Cambridge Type Model*, Review of Economic Studies, vol 44, 143 - 151, 1977.
- [19] Dumenil, G., Levy, D., *The Dynamics of Competition: A restoration of the classical analysis*, Cambridge Journal, 1984.
- [20] Goodwin, R. M., *The use of gradient dynamics in linear general disequilibrium theory*, University of Siena, 1984.
- [21] Kaldor N., *Economics without equilibrium*, New York, Sharpe Inc., 1985

[22] Arrow, K. J., *Análisis General Competitivo*, México, Fondo de Cultura Económica, 1977.

[23] Takayama, A., *Mathematical Economics*, Cambridge University Press, 1987.