



Universidad Michoacana de San Nicolás de Hidalgo

Facultad de Ciencias Físico Matemáticas

Mat. Luis Manuel Rivera Gutiérrez

Opciones americanas y tiempos óptimos

T E S I S

Que para obtener el título de:

Licenciado en Ciencias Físico Matemáticas

P r e s e n t a :

J U A N P A B L O M A L D O N A D O L O P E Z

Asesores:

Dr. Eugenio P. Balanzario

Dra. Tatjana Vukasinac

Morelia Michoacán, Septiembre de 2006

Agradecimientos

Nunca habría terminado de escribir esta tesis sin el apoyo y la paciencia de Tanja. Eugenio siempre estuvo bien dispuesto a apoyarme y a opinar sobre mi trabajo. La dedicación de Tanja y los comentarios de Eugenio fueron muy valiosos.

En la carrera tuve excelentes profesores la mayor parte del tiempo. A todos ellos les agradezco mucho el apoyo. El personal de la dirección (Paty, Celerino, Rosy, Alma, Vera, Lupita) estuvo siempre muy bien dispuesto a apoyarme con toda la burocracia.

Malú merece mención aparte, por el apoyo y la amistad que me ha brindado desde la olimpiada. También Jorge Luis y Joaquín por la confianza que me han tenido para dejarme educar a sus niños.

Mi familia también fue un apoyo muy importante. El apoyo y la motivación de mi mamá han sido importantes para mí. Es un gran ejemplo para mí en muchos aspectos.

Estoy en deuda con mis amigos (Marcos, Iván, Paco, Ana, Ahtziri, Javier, Pancho, Pacheco, Pedro, Miguel, Fany, Carlitos, Pepe, etc.). Sin las fiestas, las sobremesas, los cafés, los viajes a la playa, etc. me habría vuelto loco. Especialmente agradezco a Ahtziri, no sólo su amistad, sino por haberme enseñado a aprender matemáticas al estilo Banach. También agradezco a Javier por su amistad, sus comentarios sarcásticos y su increíble habilidad para ayudar a la gente a perder el tiempo. Con Javier y Ahtziri aprendí que las matemáticas son mucho más interesantes cuando se trabaja en equipo.

Por último, agradezco sinceramente a la gente que participó en el seminario del INI-NEE, a los olímpicos, a mis alumnos y a mis compañeros. Teniéndolos a ellos como público es como poco a poco he ido aprendiendo a comunicarme mejor matemáticamente.

Introducción

Un derivado financiero es un instrumento que depende (deriva) de otro y que actúa como una especie de seguro contra los cambios de valor del activo o del bien del cual se deriva. Para nosotros, es de especial interés un tipo de derivados financieros llamado opciones.

Una opción es un documento que nos da el derecho (pero no la obligación) de comprar o vender una acción a un precio que se fija previamente (precio de ejercicio), independientemente de la variación del precio de la acción en el mercado. Se tiene derecho de ejercer la opción hasta que concluya el llamado tiempo de maduración, que también está especificado en la opción. Una primera clasificación de las opciones es en *put* y *call*.

Una opción de tipo *put* es el documento que da el derecho al poseedor a vender una acción (del inglés, *put to*). Una opción de tipo *call* permite al poseedor comprar una acción (del inglés, *call for*) al precio que se especifica.

Es claro que si se es el poseedor de una opción *call* y el precio en el mercado de la acción es superior que el precio de ejercicio, uno puede ejercer la opción y vender inmediatamente la acción, obteniendo como ganancia la diferencia de precios. Pasa algo parecido con una opción *put*, si el precio en el mercado está por debajo del precio de ejercicio, uno ejerce la opción y vende la acción mejor que como se podría vender en el mercado.

Las opciones de ambos tipos también pueden clasificarse de acuerdo a su fecha de ejercicio. Una opción europea puede ejercerse únicamente a la fecha de maduración, en cambio, una opción americana puede ejercerse durante cualquier momento antes de la fecha de maduración.

En un caso real, si el poseedor de la opción decide no ejercerla por que no le conviene, debe pagar una pequeña cantidad al vendedor de la opción, de cualquier manera, esta cantidad no es comparable con la pérdida que re-presentaría ejercer una opción que no conviniera a los intereses del poseedor.

El problema de valuación de opciones consiste en encontrar un valor para la opción que sea justo tanto para el comprador como para el vendedor.

En 1973, Fischer Black y Myron Scholes obtuvieron una fórmula para calcular el valor

justo de una opción europea. Ellos utilizaron matemáticas de alto nivel para deducir dicha fórmula, sin embargo, en 1979, Cox, Ross y Rubinstein obtuvieron la misma fórmula a través de un modelo discreto relativamente simple. Discutiremos un poco este modelo porque es la base para la teoría que desarrollaremos sobre opciones americanas.

Para una opción americana tenemos dos problemas que están generando un campo de investigación muy activo. El primer problema se refiere al tiempo óptimo de ejercicio de una opción. Es claro que si se observan ciertas tendencias en el mercado al alza o a la baja del precio de una acción tenemos de alguna manera más información para poder decidir cuál es el mejor momento para ejercer la opción (pensando desde la perspectiva del comprador). Por otro lado, está el problema de valuación de opciones. Para este caso, no existe una fórmula cerrada como en el caso europeo. Además, pensando que el comprador va a tomar una estrategia inteligente para ejercer su opción, el vendedor debe estar preparado para poder responder al comprador. Nos centraremos en el problema de encontrar tiempos óptimos y en la manera que, conociendo estos tiempos óptimos, el vendedor debe prepararse para responder al comprador.

Ninguno de nuestros problemas centrales puede resolverse eficazmente con computadora, razón por la cual es necesario buscar soluciones alternativas. Como veremos más adelante, la necesidad de buscar otros métodos viene del crecimiento exponencial de las variables involucradas.

En el primer capítulo, introduciremos algunos conceptos de probabilidad esenciales para la comprensión del material expuesto. En el segundo capítulo introduciremos algunas ideas sobre martingalas y procesos estocásticos. El valor de una acción cambia con el tiempo de manera aleatoria y, en general, impredecible, sin embargo, haciendo ciertas suposiciones sobre esta variación en el precio de la acción y observando la evolución en el precio desde el momento inicial hasta el día de hoy, podemos tomar decisiones sobre ejercer o no la opción. El tercer capítulo introduce el modelo discreto de Cox, Ross y Rubinstein que utilizaremos para describir la variación en el precio de la acción. En el cuarto capítulo construiremos una herramienta que ayudará al comprador a escoger estrategias óptimas para maximizar sus ganancias, la envolvente de Snell. En el último capítulo, discutiremos la cobertura de reclamos contingentes para el caso de opciones americanas: el vendedor de la opción tiene que estar preparado para poder responder al comprador de la opción (esto es, tiene que disponer de suficientes activos para evitar llegar él mismo a la bancarrota), suponiendo que el comprador buscará siempre maximizar sus ganancias.

Como cualquier modelo matemático, el modelo que presentamos está sujeto a muchas restricciones que a primera vista parecen demasiado desconectadas de la realidad. Hay

que tener en cuenta que es válido suponer condiciones ideales por simplicidad, pero que estas condiciones no se mantendrán por mucho tiempo. Este fue el gran error de Black y Scholes en su vida práctica: una compañía que fundaron ellos, que se dedicaría a aprovecharse de las oportunidades de arbitraje en el mercado (i.e. enriquecimiento sin inversión) quebró espectacularmente al poco tiempo, estando cerca de crear un colapso en el mercado mundial en 1997. La razón esencial del fracaso fue un cambio imprevisto en la política rusa que no estaba contemplado en el modelo.

Índice general

1. Probabilidad	7
1.1. Conceptos básicos	7
2. Martingalas y procesos estocásticos	15
2.1. Martingalas en conjuntos discretos	15
2.2. Descomposición de Doob	17
3. El modelo Cox-Ross-Rubinstein	21
3.1. Una pequeña introducción al modelo	21
3.2. Valuación de opciones en el modelo CRR	26
4. La envolvente de Snell	29
4.1. Tiempos de paro	29
4.2. La envolvente de Snell y tiempos de paro óptimo	31
4.3. Procesos interrumpidos	36
4.4. Caracterización de tiempos óptimos	38
4.5. La envolvente de Snell y los procesos de Markov	46
5. Reclamos contingentes	51
5.1. Portafolios y estrategias de mercado	52
5.2. Mercados completos y valuación de opciones	54
5.3. Opciones americanas	59

Capítulo 1

Probabilidad

1.1. Conceptos básicos

Como hemos mencionado ya en la introducción, en este capítulo presentaremos algunos conceptos básicos de probabilidad que nos permitirán comprender el material que expondremos.

Definición 1.1.1. Sea Ω un conjunto. La **potencia** de Ω , denotada por $\wp(\Omega)$, es el conjunto cuyos elementos son los subconjuntos de Ω .

Observación 1.1.2. Si Ω es un conjunto finito con N elementos, la cardinalidad de $\wp(\Omega)$ es 2^N .

Definición 1.1.3. Un **espacio de probabilidad finito** es un par ordenado $(\Omega, \mathbb{P})$ donde Ω es un conjunto finito no vacío, llamado **espacio muestral** y $\mathbb{P} : \wp(\Omega) \rightarrow \mathbb{R}$ una función que satisface:

(1) Para todo $A \subseteq \Omega$.

$$0 \leq \mathbb{P}(A) \leq 1.$$

(2) $\mathbb{P}(\Omega) = 1$.

(3) Si A, B son subconjuntos ajenos de Ω , se tiene que:

$$\mathbb{P}(A \cup B) = \mathbb{P}(A) + \mathbb{P}(B).$$

Decimos que los subconjuntos de Ω son **eventos**.

Observemos que la última propiedad puede extenderse para uniones finitas de subconjuntos ajenos por parejas, utilizando inducción matemática. Este hecho lo utilizamos en el siguiente lema, el cual a su vez nos facilitará algunos cálculos.

Definición 1.1.4. Una colección de subconjuntos $E_1, E_2, \dots, E_k$ de Ω es una **partición** si se satisfacen las siguientes propiedades:

$$(1) \bigcup_{i=1}^k E_i = \Omega.$$

$$(2) E_i \cap E_j = \emptyset, i \neq j.$$

Los subconjuntos $E_1, E_2, \dots, E_k$ se llaman **bloques** de la partición.

Lema 1.1.5. Sea $E_1, E_2, \dots, E_k$ una partición de Ω . Entonces, para todo $A \subseteq \Omega$,

$$\mathbb{P}(A) = \sum_{i=1}^k \mathbb{P}(A \cap E_i).$$

Demostración. Tenemos que $A = \bigcup_{i=1}^k A \cap E_i$. Además, se tiene que

$$(A \cap E_i) \cap (A \cap E_j) = \emptyset, i \neq j$$

de donde

$$\mathbb{P}(A) = \sum_{i=1}^k \mathbb{P}(A \cap E_i).$$

□

Definición 1.1.6. Una **variable aleatoria** X es una función

$$X : \Omega \rightarrow \mathbb{R}.$$

En nuestro caso, una variable aleatoria corresponde a los valores finales de los activos en los que estamos invirtiendo, dado que estos activos tuvieron una "historia" (i.e. subieron y bajaron de precio a lo largo de un cierto intervalo de tiempo).

Se acostumbra escribir $\{X = x_i\}$ en lugar de $\{\omega_j \in \Omega | X(\omega_j) = x_i\}$. De la misma manera, se escribe $\mathbb{P}(X = x_i)$ en lugar de $\mathbb{P}(\{\omega_j \in \Omega | X(\omega_j) = x_i\})$.

Definición 1.1.7. Dos eventos A y B son **independientes** si

$$\mathbb{P}(A \cap B) = \mathbb{P}(A)\mathbb{P}(B).$$

Esta definición de independencia puede parecer un poco artificial, pero quedará claro más adelante.

Definición 1.1.8. Sea X variable aleatoria. El **valor esperado o esperanza** de X está dado por:

$$\mathcal{E}_{\mathbb{P}}(X) = \sum_{i=1}^n X(\omega_i) \mathbb{P}(\omega_i).$$

Pensemos en la idea detrás de esta definición. La esperanza es un promedio de los valores que puede tomar la variable aleatoria, pesados con respecto a su probabilidad.

Consideremos ahora una variable aleatoria X y supongamos que toma un conjunto finito de valores $\{x_1, \dots, x_m\}$ (todo el tiempo consideraremos espacios muestrales finitos, por lo que es válido hacer esta suposición). Observemos que los conjuntos $\{X = x_1\}, \{X = x_2\}, \dots, \{X = x_m\}$ inducen una partición sobre Ω . Llamamos a esta partición la **partición inducida por X** . Utilizaremos esta idea para demostrar el siguiente:

Teorema 1.1.9. *Sea $RV(\Omega)$ el espacio de las variables aleatorias definidas sobre Ω . Entonces el funcional dado por $\mathcal{E} : RV(\Omega) \rightarrow \mathbb{R}$ es lineal.*

Demostración. De la definición es inmediato que $\mathcal{E}(aX) = a\mathcal{E}(X)$ para X una variable aleatoria y $a \in \mathbb{R}$.

Sea $\{x_1, x_2, \dots, x_n\}$ el rango de X y $\{y_1, y_2, \dots, y_m\}$ el rango de Y . Observemos que a cada evento en Ω le corresponde un único valor de la forma $x_i + y_j$, pues cada evento pertenece a un único bloque de la partición inducida por X y a un único bloque de la partición inducida por Y . Por el lema (??) se tiene que:

$$\begin{aligned}
 \mathcal{E}(X + Y) &= \sum_{i=1}^n \sum_{j=1}^m (x_i + y_j) \mathbb{P}(X = x_i \cap Y = y_j) \\
 &= \sum_{i=1}^n \sum_{j=1}^m x_i \mathbb{P}(X = x_i \cap Y = y_j) + \sum_{i=1}^n \sum_{j=1}^m y_j \mathbb{P}(X = x_i \cap Y = y_j) \\
 &= \sum_{i=1}^n x_i \sum_{j=1}^m \mathbb{P}(X = x_i \cap Y = y_j) + \sum_{j=1}^m y_j \sum_{i=1}^n \mathbb{P}(X = x_i \cap Y = y_j) \\
 &= \sum_{i=1}^n x_i \mathbb{P}(X = x_i) + \sum_{j=1}^m y_j \mathbb{P}(Y = y_j) \\
 &= \mathcal{E}(X) + \mathcal{E}(Y),
 \end{aligned}$$

que es lo que queríamos demostrar. □

La esperanza de una variable aleatoria es una constante que aproxima bien (en algún sentido, que puntualizaremos enseguida) a la variable aleatoria. Esta propiedad es esencial para tener una comprensión intuitiva de muchos de los resultados que demostraremos. Para ello, demostramos antes el siguiente lema.

Definición 1.1.10. Sea X una variable aleatoria con valor esperado μ . La **varianza** de X es

$$\sigma_X^2 = \text{Var}(X) = \mathcal{E}((X - \mu)^2)$$

y la **desviación estándar** la raíz cuadrada positiva de la varianza

$$\sigma_X = SD(X) = \sqrt{\text{Var}(X)}.$$

La desviación estándar es el análogo estadístico del momento de inercia. El momento de inercia es una medida de dispersión de un conjunto de masas que están a cierta distancia del centro de inercia. En nuestro caso, las probabilidades de cada uno de los eventos son las masas y la distancia es la diferencia entre la esperanza (el análogo al centro de inercia de la distribución) y cada uno de los distintos valores de la variable aleatoria. En este sentido, la esperanza aproxima de una buena manera a los valores de la variable aleatoria. Es válido entonces pensar en la esperanza como el promedio de valores de la variable aleatoria.

Lema 1.1.11. Si escribimos $\mathcal{E}(X) = \mu$, para todo $c \in \mathbb{R}$

$$\mathcal{E}[(X - \mu)^2] \leq \mathcal{E}[(X - c)^2]$$

y se cumple la igualdad si y sólo si $c = \mu$.

Demostración.

$$\begin{aligned} \mathcal{E}[(X - c)^2] &= \mathcal{E}[(X - \mu) + (\mu - c)^2] \\ &= \mathcal{E}[(X - \mu)^2] + \mathcal{E}[(\mu - c)^2] + 2\mathcal{E}[(X - \mu)(\mu - c)] \end{aligned}$$

, por linealidad obtenemos que:

$$\mathcal{E}[(X - \mu)(\mu - c)] = (\mu - c)\mathcal{E}[(X - \mu)]$$

Pero $\mathcal{E}(X - \mu) = 0$ (esto es inmediato por la definición). Se sigue que

$$\begin{aligned} \mathcal{E}(X - c)^2 &= \mathcal{E}[(X - \mu) + (\mu - c)^2] \\ &= \mathcal{E}(X - \mu)^2 + \mathcal{E}(\mu - c)^2 \geq \mathcal{E}(X - \mu)^2 \end{aligned}$$

con la igualdad si y sólo si $c = \mu$.

□

Una aplicación inmediata de la linealidad es el siguiente resultado.

Lema 1.1.12. Si $\mathbb{P}$ es una función de probabilidad fuertemente positiva, esto es,

$$\mathbb{P}(\omega) > 0, \forall \omega \in \Omega$$

y además, X, Y son variables aleatorias que satisfacen

$$X \leq Y, \mathcal{E}(X) = \mathcal{E}(Y)$$

entonces $X = Y$.

Demostración. Tenemos que $0 = \mathcal{E}(Y) - \mathcal{E}(X) = \mathcal{E}(Y - X)$. Pero como $Y(\omega) - X(\omega) \geq 0$, se tiene que :

$$\mathcal{E}(Y - X) = \sum_{\omega \in \Omega} [Y(\omega) - X(\omega)] \mathbb{P}(\omega) \geq 0$$

de donde concluimos que ambas variables aleatorias son iguales. $\square$

La partición inducida por una variable aleatoria tiene una propiedad muy importante: por la manera en que la construimos, la variable aleatoria X es constante en cada uno de los bloques de la partición: en cada bloque una constante diferente. Esta propiedad la expresamos en la siguiente definición.

Definición 1.1.13. Sea $\mathcal{P}$ una partición de Ω . Una variable aleatoria X en Ω es $\mathcal{P}$ -**medible** si X es constante en cada bloque de $\mathcal{P}$.

Observemos que la partición inducida por una variable aleatoria es una partición en la que esta variable aleatoria es medible. Sin embargo, podemos pensar en particiones donde esta misma variable aleatoria sea medible y que no sean necesariamente las inducidas por ella.

Definición 1.1.14. Dados dos eventos, A, B , la probabilidad de que suceda el evento A sabiendo que aconteció el evento B está dada por:

$$\mathbb{P}(A|B) := \frac{\mathbb{P}(A \cap B)}{\mathbb{P}(B)}.$$

Ahora es más clara y natural nuestra definición de independencia de eventos. Si dos eventos A y B son independientes, la probabilidad condicional $\mathbb{P}(A|B)$ es simplemente $\mathbb{P}(A)$, en este caso, recuperamos la definición de independencia que teníamos.

Combinando nuestra nueva definición con la definición de esperanza, tenemos un nuevo concepto.

Definición 1.1.15. La **esperanza condicional** respecto al evento A

$$\mathcal{E}_{\mathbb{P}}(X|A) := \frac{1}{\mathbb{P}(A)} \sum_{\omega \in A} X(\omega) \mathbb{P}(\omega)$$

donde A es un evento tal que $\mathbb{P}(A) > 0$.

Hemos considerado la esperanza de una variable aleatoria sabiendo que ocurrió un evento. Si tenemos una partición $\mathcal{P} = \{B_1, B_2, \dots, B_n\}$ de Ω , tiene sentido pensar en la esperanza condicional respecto a una partición, extendiendo simplemente la definición de esperanza condicional respecto a un evento de la manera natural.

Definición 1.1.16. Sea $(\Omega, \mathbb{P})$ un espacio de probabilidad finito, con una partición $\mathcal{P}$ dada por $\mathcal{P} = \{B_1, B_2, \dots, B_n\}$ tal que $\mathbb{P}(B_i) > 0$ para toda i . La **esperanza condicional** de una variable aleatoria X respecto a la partición $\mathcal{P}$ es una variable aleatoria

$$\mathcal{E}_{\mathbb{P}}(X|\mathcal{P}) : \Omega \rightarrow \mathbb{R},$$

dada por

$$\mathcal{E}_{\mathbb{P}}(X|\mathcal{P}) = \mathcal{E}_{\mathbb{P}}(X|B_1)1_{B_1} + \mathcal{E}_{\mathbb{P}}(X|B_2)1_{B_2} + \dots + \mathcal{E}_{\mathbb{P}}(X|B_n)1_{B_n}$$

donde

$$1_{B_i}(\omega) = \begin{cases} 1 & \omega \in B_i, \\ 0 & \text{en otro caso.} \end{cases}$$

Ejemplo 1.1.17. Consideremos un dado que se lanza una vez. Sea $\Omega = \{1, 2, 3, 4, 5, 6\}$ el conjunto de resultados posibles. Sea X la inclusión de Ω en $\mathbb{R}$. Por último, digamos que $\mathcal{P}$ es la partición

$$\mathcal{P} = \{\{1, 6\}, \{2, 3\}\{4, 5\}\}$$

hacemos ahora

$$B_1 = \{1, 6\}, \quad B_2 = \{2, 3\}, \quad B_3 = \{4, 5\}.$$

De acuerdo a la definición de esperanza condicional, tenemos que:

$$\mathcal{E}_{\mathbb{P}}(X|\mathcal{P}) = \mathcal{E}_{\mathbb{P}}(X|B_1)1_{B_1} + \mathcal{E}_{\mathbb{P}}(X|B_2)1_{B_2} + \mathcal{E}_{\mathbb{P}}(X|B_3)1_{B_3}.$$

por otra parte,

$$\begin{aligned} \mathcal{E}_{\mathbb{P}}(X|B_1) &= 1 \cdot \frac{1}{2} + 6 \cdot \frac{1}{2} = \frac{7}{2}, \\ \mathcal{E}_{\mathbb{P}}(X|B_2) &= 2 \cdot \frac{1}{2} + 3 \cdot \frac{1}{2} = \frac{5}{2}, \\ \mathcal{E}_{\mathbb{P}}(X|B_3) &= 4 \cdot \frac{1}{2} + 5 \cdot \frac{1}{2} = \frac{9}{2} \end{aligned}$$

ya que, para nuestro caso, la probabilidad condicional dado cada uno de los bloques es $\frac{1}{2}$. Entonces,

$$\mathcal{E}_{\mathbb{P}}(X|\mathcal{P})(\omega) = \begin{cases} \frac{7}{2}, & \omega = 1, 6 \\ \frac{5}{2}, & \omega = 2, 3 \\ \frac{9}{2}, & \omega = 4, 5. \end{cases}$$

Introduciremos un lema que nos será de utilidad después, aunque de momento parezca un poco fuera de contexto.

Lema 1.1.18. *Sean X, Y dos variables aleatorias definidas sobre un espacio muestral Ω y sea $\mathcal{P}$ una partición de Ω . Se tiene que:*

$$\max\{\mathcal{E}(X|\mathcal{P}), \mathcal{E}(Y|\mathcal{P})\} \leq \mathcal{E}(\max\{X, Y\}|\mathcal{P})$$

Demostración. Observemos que el funcional $\mathcal{E}(\cdot|\mathcal{P})$ respeta el orden, es decir, si una variable aleatoria es mayor o igual que otra, al tomar la esperanza condicional respecto a una partición, la esperanza condicional de la primera variable sigue siendo mayor o igual que la segunda. Esto es inmediato de la definición. Como $\max\{X, Y\} \geq X, Y$ se sigue inmediatamente la desigualdad. $\square$

Hemos dicho que en cierto sentido la esperanza es el promedio de los valores de la variable aleatoria. Dada una partición, la esperanza condicional con respecto a esta partición es aquella función que es constante en los bloques de la partición y que mejor aproxima a la variable aleatoria (siempre pensando en el sentido de aproximación cuadrática media). Una propiedad muy importante es la siguiente.

Teorema 1.1.19. *(Propiedad de las torres) Si $\mathcal{Q}$ es una partición más fina que $\mathcal{P}$ (es decir, los bloques de $\mathcal{Q}$ están contenidos en los bloques de $\mathcal{P}$), se tiene que:*

$$\mathcal{E}(\mathcal{E}(X|\mathcal{P})|\mathcal{Q}) = \mathcal{E}(X|\mathcal{P}) = \mathcal{E}(\mathcal{E}(X|\mathcal{Q})|\mathcal{P}).$$

En otras palabras, la esperanza puede tomarse en cualquier orden, finalmente lo único que cuenta es la esperanza en la partición más gruesa.

Demostración. Sea B un bloque de $\mathcal{Q}$ y sea A el bloque de $\mathcal{P}$ tal que $B \subset A$. La función $\mathcal{E}(X|\mathcal{P})$ es constante en este bloque, y en particular en B , por lo que $\mathcal{E}(\mathcal{E}(X|\mathcal{P})|B) = \mathcal{E}(X|A)$. Haciendo esto para cada bloque de $\mathcal{Q}$, obtenemos que $\mathcal{E}(\mathcal{E}(X|\mathcal{P})|\mathcal{Q}) = \mathcal{E}(X|\mathcal{P})$. La otra igualdad es análoga. $\square$

Capítulo 2

Martingalas y procesos estocásticos

En este capítulo introduciremos las definiciones de martingalas y procesos estocásticos. Las martingalas eran una estrategia de juego utilizada en Francia durante la época de la Revolución, que se aseguraba que siempre permitía ganar en algunos juegos de azar. Supongamos que estamos apostando a una moneda. La estrategia consiste en apostar el doble cada vez que perdamos. De esta manera, si la moneda es justa, después de una sucesión de pérdidas vendrá alguna ganancia; en este momento, un cálculo sencillo nos permite ver que recuperamos la cantidad perdida. El truco consistía entonces en retirarse en cuanto se hubiera empezado a ganar. Los casinos modernos tienen maneras de prohibir el uso de esta estrategia, limitando su utilización o limitando las cantidades que se apuestan.

2.1. Martingalas en conjuntos discretos

Hay toda una teoría en torno a las martingalas. Nosotros estudiaremos el caso discreto, que es relativamente simple.

Definición 2.1.1. *Un proceso estocástico finito en un espacio muestral Ω es una sucesión $X_0, X_1, \dots, X_N$ de variables aleatorias definidas en Ω .*

Definición 2.1.2. *Una sucesión $\mathbb{F} = (\mathcal{P}_0, \mathcal{P}_1, \dots, \mathcal{P}_N)$ de particiones de un conjunto $\Omega = (\omega_1, \omega_2, \dots, \omega_N)$ tales que:*

$$\mathcal{P}_0 \prec \mathcal{P}_1 \prec \dots \prec \mathcal{P}_N$$

*donde $\prec$ significa "más gruesa que" se llama **filtración**.*

Definición 2.1.3. *Dada una filtración $\mathbb{F} = (\mathcal{P}_0, \mathcal{P}_1, \dots, \mathcal{P}_N)$ decimos que un proceso estocástico $\mathbb{X} = (X_0, X_1, \dots, X_N)$ es una **$\mathbb{F}$ -martingala** si $\mathbb{X}$ está adaptado a $\mathbb{F}$, es decir, X_i es $\mathcal{P}_i$ -medible y*

$$\mathcal{E}(X_{k+1}|\mathcal{P}_k) = X_k.$$

Observación 2.1.4. *De la definición, es inmediato que*

$$\mathcal{E}(X_k|\mathcal{P}_k) = X_k.$$

Las martingalas modelan juegos justos, es decir, juegos donde la ganancia esperada de un turno a otro es cero.

Gracias a la propiedad de las torres que enunciamos arriba, tenemos que:

Proposición 2.1.5. *Si $\mathbb{X} = (X_0, X_1, \dots, X_N)$ es una $\mathbb{F}$ -martingala donde $\mathbb{F} = (\mathcal{P}_0, \mathcal{P}_1, \dots, \mathcal{P}_N)$, para toda $i > 0$, se tiene que:*

$$\mathcal{E}(X_{k+i}|\mathcal{P}_k) = X_k.$$

Demostración. Consideramos k fijo. Para $i = 1$ se sigue directamente de la definición. Supongamos que es cierto para $i = j$. Para $i = j + 1$, por definición se tiene que:

$$X_{k+j} = \mathcal{E}(X_{k+j+1}|\mathcal{P}_{k+j}).$$

Tomando la esperanza respecto a la partición $\mathcal{P}_k$ en ambos lados y aplicando la hipótesis de inducción y la propiedad de las torres, tenemos que:

$$X_k = \mathcal{E}(X_{k+j}|\mathcal{P}_k) = \mathcal{E}(\mathcal{E}(X_{k+j+1}|\mathcal{P}_{k+j})|\mathcal{P}_k) = \mathcal{E}(X_{k+j+1}|\mathcal{P}_k)$$

lo que termina nuestra prueba. □

Hay otros dos conceptos relacionados directamente con la idea de martingala.

Definición 2.1.6. *Dada una filtración $\mathbb{F} = (\mathcal{P}_0, \mathcal{P}_1, \dots, \mathcal{P}_N)$ decimos que un proceso estocástico $\mathbb{X} = (X_0, X_1, \dots, X_N)$ es una $\mathbb{F}$ -**supermartingala** si $\mathbb{X}$ está adaptado a $\mathbb{F}$, es decir, X_i es $\mathcal{P}_i$ -medible y*

$$\mathcal{E}(X_{k+1}|\mathcal{P}_k) \leq X_k.$$

Definición 2.1.7. *Dada una filtración $\mathbb{F} = (\mathcal{P}_0, \mathcal{P}_1, \dots, \mathcal{P}_N)$ decimos que un proceso estocástico $\mathbb{X} = (X_0, X_1, \dots, X_N)$ es una $\mathbb{F}$ -**submartingala** si $\mathbb{X}$ está adaptado a $\mathbb{F}$, es decir, X_i es $\mathcal{P}_i$ -medible y*

$$\mathcal{E}(X_{k+1}|\mathcal{P}_k) \geq X_k.$$

Ejemplo 2.1.8. Consideremos el caso de un apostador que gana \$1 cada vez que, al lanzar una moneda, se obtiene cara y pierde la misma cantidad si sale sello. Supongamos que la moneda no es justa y se obtiene cara con una probabilidad p . Tenemos entonces tres casos:

- (1) $p = \frac{1}{2}$. En este caso, en promedio el apostador no pierde ni gana, por lo que este proceso estocástico es una martingala.
- (2) $p \geq \frac{1}{2}$. En este caso, en promedio el apostador gana, por lo que este proceso estocástico es una submartingala.
- (3) $p \leq \frac{1}{2}$. En este caso, en promedio el apostador pierde, por lo que este proceso estocástico es una supermartingala.

2.2. Descomposición de Doob

Un proceso (X_k) adaptado a una filtración puede pensarse como una martingala *desacomodada*. Podemos fijar X_0 . Sabemos que X_1 toma valores constantes en los bloques de la partición $\mathcal{P}_1$. Si restamos a estos valores constantes otra constante apropiada, podemos hacer que el promedio de estos nuevos valores adaptados sea justamente X_0 . Esto es esencialmente la idea del teorema de descomposición de Doob.

Un proceso adaptado a una filtración puede descomponerse, bajo ciertas condiciones, como la suma de una martingala y de un **proceso predecible**.

Definición 2.2.1. Decimos que un proceso (T_k) es predecible si T_k es $\mathcal{P}_{k-1}$ -medible.

Teorema 2.2.2. (Descomposición de Doob) Sea $\mathbb{X} = (X_k)$ un proceso estocástico $\mathbb{F}$ -adaptado. Entonces existe una única martingala (M_k) y un único proceso predecible (T_k) tal que $X_k = M_k - T_k$ y $T_0 = 0$. Además, si $\mathbb{X}$ es una supermartingala, $T_{k+1} \geq T_k$.

Demostración. Sea $M_0 = X_0$, $T_0 = 0$, y para $k > 0$,

$$M_k = \sum_{i=1}^k [X_i - \mathcal{E}(X_i | \mathcal{P}_{i-1})] + X_0$$

y $T_k = M_k - X_k$. Entonces,

$$M_{k+1} = M_k + X_{k+1} - \mathcal{E}(X_{k+1} | \mathcal{P}_k),$$

de donde

$$\mathcal{E}(M_{k+1} | \mathcal{P}_k) = \mathcal{E}(M_k | \mathcal{P}_k) + \mathcal{E}(X_{k+1} | \mathcal{P}_k) - \mathcal{E}(\mathcal{E}(X_{k+1} | \mathcal{P}_k) | \mathcal{P}_k).$$

De la definición de M_k vemos que es claro que M_k es $\mathcal{P}_k$ -medible, ya que, para $i < k$, X_i es $\mathcal{P}_k$ -medible, además,

$$\mathcal{E}(\mathcal{E}(X_{k+1}|\mathcal{P}_k)|\mathcal{P}_k) = \mathcal{E}(X_{k+1}|\mathcal{P}_k),$$

de donde la igualdad de arriba se transforma en

$$\mathcal{E}(M_{k+1}|\mathcal{P}_k) = M_k,$$

hemos probado así que (M_k) es martingala.

Por otro lado,

$$T_k = M_{k-1} - \mathcal{E}(X_k|\mathcal{P}_{k-1})$$

que es $\mathcal{P}_{k-1}$ -medible.

Supongamos ahora que $\mathbb{X}$ es supermartingala. Entonces tenemos que

$$\begin{aligned} M_k - T_k &= X_k \geq \mathcal{E}(X_{k+1}|\mathcal{P}_k) \\ &= \mathcal{E}(M_{k+1}|\mathcal{P}_k) - \mathcal{E}(T_{k+1}|\mathcal{P}_k) \\ &= M_k - T_{k+1} \end{aligned}$$

de donde $T_k \leq T_{k+1}$.

Veamos por qué esta descomposición es única. Sean $X_k = M_k - T_k = N_k - S_k$ dos descomposiciones de Doob para (X_k) con $(M_k), (N_k)$ martingalas y $(S_k), (T_k)$ procesos predecibles. Es claro que $X_0 = M_0 = N_0$ y que $T_0 = S_0 = 0$. Tenemos entonces que:

$$\begin{aligned} M_0 - \mathcal{E}(T_1|\mathcal{P}_0) &= \mathcal{E}(X_1|\mathcal{P}_0) \\ &= \mathcal{E}(N_1 - S_1|\mathcal{P}_0) \\ &= N_0 - \mathcal{E}(S_1|\mathcal{P}_0) \end{aligned}$$

dado que $X_0 = M_0 = N_0$,

$$\mathcal{E}(S_1|\mathcal{P}_0) = \mathcal{E}(T_1|\mathcal{P}_0),$$

pero como S_1, T_1 son constantes en $\mathcal{P}_0$, tenemos entonces que $S_1 = T_1$. La unicidad se sigue de un argumento análogo por inducción.

□

Reflexionemos un poco sobre la definición que dimos para predictibilidad. El proceso predecible son las constantes que mencionamos al principio de la sección que nos permitieron ajustar nuestro proceso estocástico a una martingala. Al tiempo cero, la martingala

empieza en X_0 y no hay nada que predecir. Como ya conocemos X_1 , en ese mismo tiempo cero sabemos cómo tenemos que ajustar el proceso (X_k) a una martingala, es decir, conocemos la constante que tenemos que restar a X_1 para hacer que X_0 sea su promedio. En este sentido estamos considerando que un proceso es predecible: antes de hacer el ajuste correspondiente, ya sabemos cómo tenemos que ajustar.

Capítulo 3

El modelo Cox-Ross-Rubinstein

3.1. Una pequeña introducción al modelo

El modelo discreto que presentamos a continuación es una representación relativamente simple del comportamiento del precio de una acción a través del tiempo. A pesar de su simplicidad, esta es una herramienta muy poderosa, como podremos darnos cuenta más adelante.

Consideraremos un modelo con dos activos, uno de ellos libre de riesgo y el otro con riesgo. Decir que un activo no tenga riesgo se refiere a que sabemos exactamente cómo va a evolucionar durante el tiempo que vamos a modelar. El activo con riesgo es un activo cuyo precio puede variar durante este intervalo de tiempo. Supondremos que tenemos una opción americana de tipo *call*.

Queremos modelar la situación para los tiempos $t_0, t_1, \dots, t_N$. Estos tiempos se llaman **tiempos de oportunidad** y son los tiempos en los que podremos o no ejercer nuestra opción. En general, los intervalos de tiempo no tienen por qué tener la misma longitud, supondremos que sí, para evitarnos algunas constantes de ajuste que aparecerían más adelante. Hecha esta aclaración, no haremos distinción entre decir tiempo t_i y tiempo i .

Empecemos considerando un espacio muestral $\Omega_0 = \{0, 1\}$, con una función de probabilidad $\mathbb{P}_0$ tal que $\mathbb{P}_0(1) = p$, $\mathbb{P}_0(0) = 1 - p$ donde 1 representa un aumento en el valor del activo y 0 una disminución en el mismo. Sean u, d los porcentajes de aumento y disminución, respectivamente, del precio del activo. Cuando haya transcurrido un tiempo N , tendremos un nuevo espacio muestral:

$$\Omega = \Omega_0^N = \{(\omega_1, \omega_2, \dots, \omega_N) | \omega_i \in \Omega_0, i = 1, 2, \dots, N\}$$

Consideremos también las proyecciones $Z_t : \Omega \rightarrow \{0, 1\}$, tales que, para $\omega = (\omega_1, \omega_2, \dots, \omega_N)$, $Z_t(\omega) = \omega_t$, es decir, el resultado obtenido en el tiempo t , y también consideremos:

$$D_t(\omega) = \sum_{s=1}^t Z_s(\omega)$$

el total de veces que la acción ha subido de precio hasta el tiempo t . Tenemos entonces que la distribución de probabilidad para estas variables aleatorias (D_t) es una distribución de Bernoulli:

$$\mathbb{P}(\{D_t = k\}) = \binom{t}{k} p^k (1-p)^{t-k}.$$

Podemos entonces extender la función $\mathbb{P}_0$ a Ω de la siguiente manera.

$$\mathbb{P} : \Omega \rightarrow \mathbb{R}, \quad \mathbb{P}(\omega) = p^{D_N(\omega)} (1-p)^{N-D_N(\omega)}.$$

Supongamos que el activo sin riesgo tiene un valor inicial A_0 que cambia con el tiempo a una tasa de interés constante r . Entonces, en el tiempo t , el valor del activo es:

$$A_t = A_0(1+r)^t.$$

De la misma manera, si S_0 es el valor inicial del activo con riesgo, su valor al tiempo t está dado por:

$$S_t(\omega) = S_0(1+u)^{D_t(\omega)}(1+d)^{t-D_t(\omega)}.$$

Además, es fácil observar que se tiene la siguiente relación de recurrencia:

$$S_{t+1}(\omega) = S_t(1+u)^{Z_{t+1}(\omega)}(1+d)^{t-Z_{t+1}(\omega)}.$$

Esta relación nos será de mucha utilidad más adelante.

Para el modelo Cox-Ross-Rubinstein se supone además las siguientes dos condiciones para u y d :

$$(1+u)(1+d) = 1,$$

es decir, aumento y disminución en el precio del activo se anulan mutuamente.

En nuestros ejemplos, como en muchos casos reales, consideraremos el caso en que d es negativo.

Consideraremos de nuestro interés el caso $d < r < u$. Si $u < r$, en realidad no nos habría convenido entrar al mercado y hubiera sido mejor invertir todo nuestro dinero en el activo sin riesgo. Si $d > r$, significaría que nos conviene más perder dinero que invertir en el activo sin riesgo, lo cual no tiene sentido. El porqué de estas desigualdades tiene una justificación formal en términos del teorema siguiente.

Definición 3.1.1. Una **medida de martingala equivalente** asociada a un proceso (S_k) es una probabilidad $\mathbb{Q}$ tal que:

- (1) (S_k) es martingala (con esta probabilidad).
- (2) $\mathbb{Q}$ es **fuertemente positiva**, i.e., $\mathbb{Q}(\omega) > 0, \forall \omega \in \Omega$.

Teorema 3.1.2. (Teorema fundamental de valuación de activos) Un mercado es libre de oportunidades de arbitraje sí y sólo sí admite una medida de martingala equivalente.

Hay dos razones para trabajar con el activo sin riesgo como unidad: por un lado nos evitamos la complicación de tener que usar unidades monetarias específicas que también varían con el tiempo; por otra parte, el utilizar al activo sin riesgo como unidad de medida nos permite evaluar si las ganancias que esperamos obtener al arriesgarnos son en verdad mejores que las ganancias que obtendríamos invirtiendo solamente en el activo sin riesgo.

Definición 3.1.3. Los **precios descontados** de los activos sin riesgo y con riesgo, respectivamente, están dados por:

$$\tilde{A}_t = \frac{A_t}{A_t} = 1, \quad \tilde{S}_t = \frac{S_t}{A_t} = \frac{S_t}{A_0(1+r)^t}.$$

Observación 3.1.4. La relación de recurrencia para (S_t) induce de manera natural una relación de recurrencia para $(\tilde{S}_t)$, dada por:

$$\tilde{S}_{t+1} = \frac{\tilde{S}_t(1+u)^{Z_{t+1}}(1+d)^{1-Z_{t+1}}}{1+r}.$$

Impondremos la condición adicional de que el proceso $(\tilde{S}_t)$ sea una martingala respecto a la distribución de probabilidad dada por $\mathbb{P}$, para que esto nos permita obtener un valor para p .

Para el modelo Cox-Ross-Rubinstein, que denotaremos de ahora en adelante como CRR, ignoraremos la probabilidad natural de que haya un aumento en el precio de la acción. En vez de ello, impondremos una probabilidad un tanto artificial para poder suponer que el proceso descontado es una martingala. Un hecho muy interesante es que esta hipótesis, que es deseable matemáticamente, es equivalente a una condición económicamente deseable. Puntualizaremos este hecho en su momento.

Teorema 3.1.5. El proceso $(\tilde{S}_t)$ es martingala sí y sólo sí $p = \frac{r-d}{u-d}$.

Demostración. Si la partición $\mathcal{P}_t$ está formada por los bloques $B_1, \dots, B_k$, se tiene que

$$\mathcal{E}_{\mathbb{P}}(\tilde{S}_{t+1}|\mathcal{P}_t) = \mathcal{E}_{\mathbb{P}}(\tilde{S}_{t+1}|B_1)1_{B_1} + \mathcal{E}_{\mathbb{P}}(\tilde{S}_{t+1}|B_2)1_{B_2} + \dots + \mathcal{E}_{\mathbb{P}}(\tilde{S}_{t+1}|B_k)1_{B_k}.$$

Observemos que en cada uno de los bloques, la esperanza condicional podemos escribirla de la siguiente manera:

$$\mathcal{E}_{\mathbb{P}}(\tilde{S}_{t+1}|B_j) = p \frac{1+u}{1+r} \tilde{S}_t + (1-p) \frac{1+d}{1+r} \tilde{S}_t,$$

ya que, para este bloque, S_t es una constante porque conocemos como fue la historia hasta el tiempo t para todos los eventos que pertenecen a este bloque.

Podemos entonces extender nuestra definición a todos los bloques de la partición, tomando el valor constante correspondiente a S_t en cada uno de los bloques, entonces podemos escribir

$$\mathcal{E}(\tilde{S}_{t+1}|\mathcal{P}_t)(\bar{\omega}) = \frac{p(1+u) + (1-p)(1+d)}{1+r} \tilde{S}_t(\bar{\omega}).$$

De esta última ecuación podemos obtener las dos implicaciones. Si hacemos que $p = \frac{r-d}{u-d}$, obtenemos que el proceso $(\tilde{S}_t)$ es martingala; si suponemos que es martingala, podemos cancelar en ambos lados de la igualdad y encontrar el valor de p , lo que termina nuestra prueba. $\square$

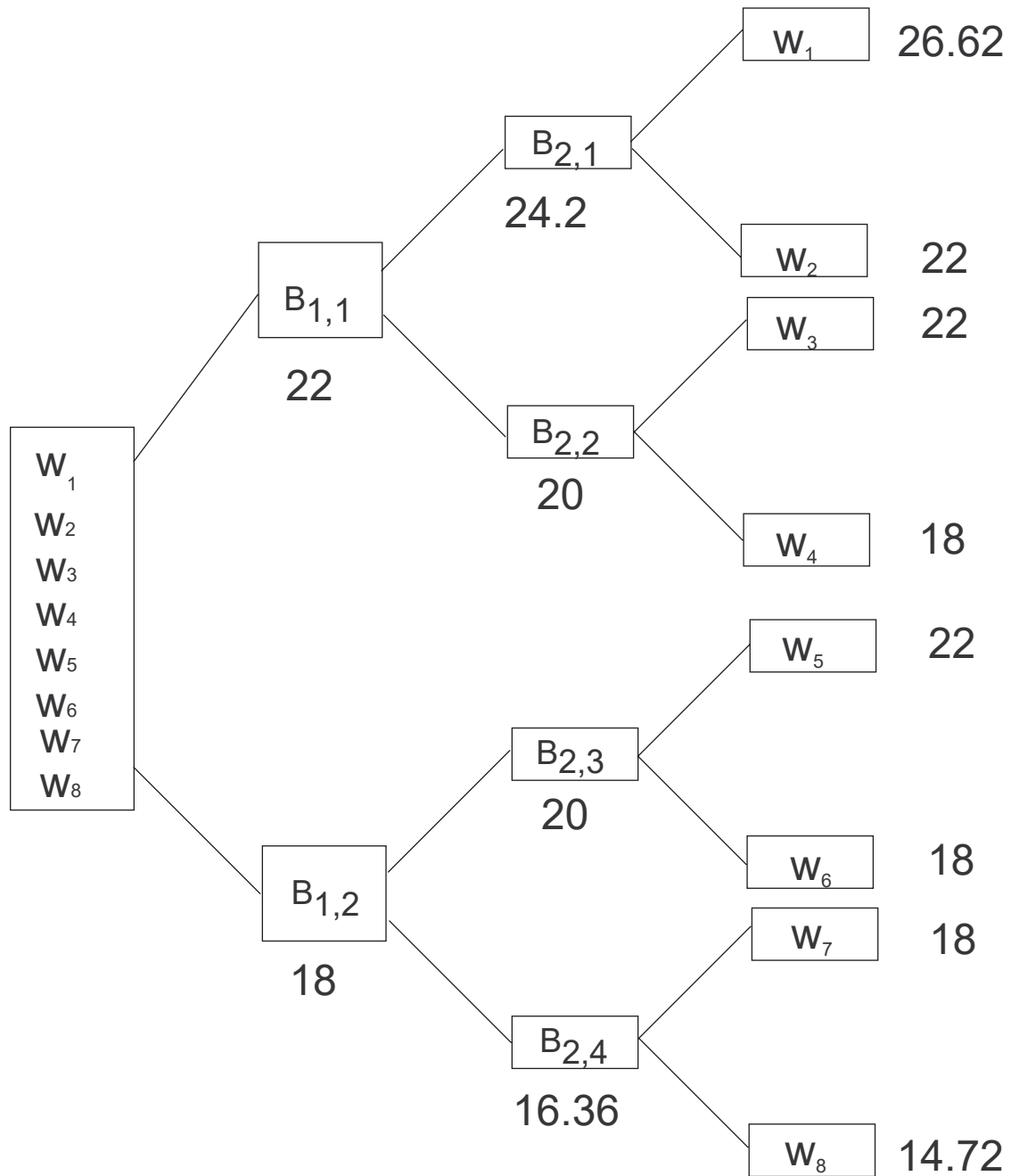
Ahora consideraremos un ejemplo de un modelo CRR para el caso $N = 3$, cuyo diagrama se muestra en la figura (??).

Los eventos ω_i son los siguientes:

$$\begin{aligned} \omega_1 &= (1, 1, 1), & \omega_2 &= (1, 1, 0) \\ \omega_3 &= (1, 0, 1), & \omega_4 &= (1, 0, 0) \\ \omega_5 &= (0, 1, 1), & \omega_6 &= (0, 1, 0) \\ \omega_7 &= (0, 0, 1), & \omega_8 &= (0, 0, 0). \end{aligned}$$

En este caso, consideraremos también $r = 0$, $u = 0,1$, $d = \frac{1}{1,1} - 1$. Consideraremos que tenemos una opción americana de tipo *call* la cual tenemos derecho de ejercerla y comprar la mercancía a precio $K = 21$. En el tiempo inicial, el precio de la mercancía es 20. Para cada uno de los tiempos, consideraremos las ganancias en ese tiempo. La ganancia en el tiempo k es entonces la siguiente variable aleatoria:

$$Y_k := \max\{S_k - 21, 0\}$$

Figura 3.1: Diagrama CRR en el caso $N=3$

donde S_k es el precio sin descontar de la acción al tiempo k . Tomamos $r = 0$ por simplicidad; la idea es la misma pero la notación se vuelve muy confusa. Esta definición de ganancias lo que quiere decir es que si en el tiempo k la acción subió de 21, nos conviene ejercerla, es decir, comprar la mercancía a 21 y venderla inmediatamente, siendo nuestra ganancia justamente dicha diferencia. En caso de que no haya subido, nos conviene esperar y ejercer la acción después (o incluso dejarla expirar). Para nuestro ejemplo tenemos que:

$$Y_3 = \begin{cases} 5,62 & \omega = \omega_1, \\ 1 & \omega = \omega_2, \omega_3, \omega_5, \\ 0 & \text{en otro caso.} \end{cases}$$

$$Y_2 = \begin{cases} 3,2 & \omega \in B_{2,1}, \\ 0 & \text{en otro caso.} \end{cases}$$

$$Y_1 = \begin{cases} 1 & \omega \in B_{1,1}, \\ 0 & \text{en otro caso.} \end{cases}$$

y además, $Y_0 = 0$.

3.2. Valuación de opciones en el modelo CRR

Una de las aplicaciones más sorprendentes del modelo CRR es la deducción de la fórmula de Black y Scholes para valorar opciones europeas. El enfoque original involucra cálculo estocástico y herramientas de matemáticas más sofisticadas. El enfoque relativamente sencillo del modelo CRR, junto con algunos teoremas fuertes de probabilidad, permite llegar a la fórmula de Black y Scholes. En varios de los textos presentados en la bibliografía puede encontrarse una deducción de esta fórmula a partir del modelo CRR. El problema de la valuación de opciones consiste en encontrar un precio justo tanto para el vendedor como el comprador de la opción, dado que ninguno de los dos tiene certeza sobre el desarrollo futuro de los valores de la acción, y por tanto, no puede asegurar nada sobre el precio de la opción. De momento nos olvidaremos de las opciones americanas para concentrarnos en el caso europeo.

Sea K el precio de ejercicio de una opción europea tipo *call* vendida por nosotros. Es claro que si el precio de la acción al tiempo T es $\tilde{S}_T$, con $\tilde{S}_T > K$, al comprador de la opción le conviene ejercerla, es decir, comprarnos la acción a K pesos y venderla inmediatamente después al precio $\tilde{S}_T$ en el mercado. Por el contrario, si $\tilde{S}_T < K$ entonces no le conviene comprarnos la acción, sino dejarla expirar. La función de pago (la cantidad

de dinero que perderemos como consecuencia de que el comprador ejerza su opción) de la opción al tiempo T está dada por:

$$C_T = (S_T - K)_+ = \max\{S_T - K, 0\}.$$

Observemos que en esta última ecuación ponemos S_T y no $\tilde{S}_T$. Esto es porque nosotros como vendedores tenemos que considerar el precio sin el descuento correspondiente. El comprador necesita hacer esta consideración porque él está interesado en obtener el mayor beneficio posible de ganancias *relativas*, es decir, si está arriesgándose a comprar una opción es porque espera ganar más de lo que ganaría invirtiendo simplemente en el activo sin riesgo. Para el vendedor esto no importa: una pérdida es simplemente una pérdida, no le interesa saber si va a perder mucho o poco en relación al activo sin riesgo, sino simplemente evitar perder.

Para ser consistente con la notación usual, escribiremos ξ_t en lugar de $\tilde{S}_t$.

Ahora construimos una sucesión de variables aleatorias (η_j) tales que:

$$\eta_N := e^{-rN}(S_N - K), \quad \eta_j = \mathcal{E}(\eta_{j+1} | \mathcal{P}_j).$$

La sucesión (η_j) es simplemente trasladar el proceso $(\tilde{S}_t)$ mediante constantes y medir a partir del tiempo cero. Estamos considerando el caso en que el comprador decida ejercer su opción y estamos previniéndonos para eso.

Observemos que el valor de la función de pago al tiempo j , medido al tiempo 0 está representado por η_j , mientras que $B_j := e^{rj}\eta_j$ es el valor medido, al tiempo j , de la función de pago al tiempo j . Observemos que $B_0 = \eta_0$.

El siguiente lema es un resultado clave que nos permitirá encontrar el precio justo inicial de la opción.

Lema 3.2.1. *Consideremos la sucesión (Φ_j) dada por:*

$$\Phi_{j+1} = \frac{\eta_{j+1} - \eta_j}{\xi_{j+1} - \xi_j}.$$

Esto es una cantidad conocida al tiempo j .

Demostración. A partir del estado del mercado hasta el día j , podemos pensar en dos estados posibles para el día $j + 1$. Sean $C < B$ los dos valores posibles de ξ_{j+1} y sea $A = pB + (1 - p)C$. Análogamente, sean $\gamma < \beta$ los dos valores posibles de η_{j+1} y sea $\alpha = p\beta + (1 - p)\gamma$. De la condición que los procesos (η_j) y (ξ_j) sean martingalas, tenemos entonces que $\alpha = \eta_j$ y $A = \xi_j$.

Si probamos que:

$$\frac{\beta - \alpha}{B - A} = \frac{\gamma - \alpha}{C - A}$$

es decir, que los dos valores posibles de Φ_{j+1} coinciden. Pero observemos que:

$$\frac{\beta - \alpha}{\gamma - \alpha} = -\frac{1 - p}{p} = \frac{B - A}{C - A}$$

lo que concluye nuestra prueba. $\square$

Mostraremos que este lema implica que el precio B_0 es un precio justo para la opción, ya que utilizando (Φ_j) construiremos un portafolio que replique la función de pago de quien nos compra la opción.

Consideremos el portafolio dado por $(B_j - \Phi_{j+1}S_j, \Phi_{j+1})$ donde Φ_{j+1} representa la cantidad de acciones que un inversionista adquiere al tiempo j y la primera coordenada representa la cantidad de dinero que el inversionista deposita en (o toma prestado de) el banco al tiempo j .

Observemos que al tiempo j , el costo total del portafolio es B_j . Ahora, al tiempo $j + 1$ el valor del portafolio es:

$$S_{j+1}\Phi_{j+1} + (B_j - \Phi_{j+1}S_j)e^r.$$

medido al tiempo cero, es decir, multiplicando por un factor de descuento $e^{-(j+1)r}$, obtenemos que

$$e^{-(j+1)r}S_{j+1}\Phi_{j+1} + (B_j - \Phi_{j+1}S_j)e^{-jr}.$$

Reescribiendo y utilizando el lema anterior,

$$\eta_j + \Phi_{j+1}(\xi_{j+1} - \xi_j) = \eta_{j+1} = e^{-r(j+1)}B_{j+1}.$$

Pero η_{j+1} representa B_{j+1} pesos al tiempo $j + 1$. Regresando al valor del portafolio al tiempo $j + 1$, hemos probado que este valor es justamente B_{j+1} . Este dinero se puede reinvertir de acuerdo al portafolio para el tiempo $j + 1$.

Si repetimos este proceso N veces, observamos que el vendedor de la opción tiene la posibilidad de terminar con la misma cantidad $B_N = C_N$ que el comprador de la opción. Esto justifica el hecho de que la opción se venda en B_0 pesos.

Capítulo 4

La envolvente de Snell

4.1. Tiempos de paro

Partiendo de un modelo CRR para el mercado, nuestro interés es construir estrategias óptimas para maximizar nuestras ganancias, basándonos en la información que dispongamos sobre el estado del mercado en los días anteriores. Para ello, introduciremos una definición formal de un tiempo de paro óptimo y construiremos, utilizando la teoría de los capítulos anteriores, una herramienta auxiliar para calcular estrategias óptimas (pensando en que una estrategia óptima es hacer una elección de un tiempo de paro óptimo). Esta herramienta auxiliar, la envolvente de Snell, nos permite calcular solamente el primero y el último tiempo de paro óptimo, pero desafortunadamente no nos permite calcular tiempos intermedios. El cálculo de tiempos de paro óptimos intermedios es un problema computacionalmente difícil y hasta el momento no hay una estrategia eficaz para los casos reales en que se consideran más activos con riesgo.

Definición 4.1.1. Sea Ω un conjunto no vacío. Una colección $\mathcal{A}$ de subconjuntos de Ω se llama un **álgebra de conjuntos** (o simplemente, un **álgebra**) si satisface las siguientes propiedades:

- (1) $\emptyset \in \mathcal{A}$.
- (2) $A \in \mathcal{A} \rightarrow A^c \in \mathcal{A}$.
- (3) $A, B \in \mathcal{A} \rightarrow A \cup B \in \mathcal{A}$.

Se puede demostrar utilizando las leyes de De Morgan el siguiente lema:

Lema 4.1.2. Sean $A, B \in \mathcal{A}$. Entonces:

- (1) $A \cap B \in \mathcal{A}$.

(2) $A \setminus B \in \mathcal{A}$.

Demostración. Para (1), tenemos que:

$$A \cap B = (A^c \cup B^c)^c \in \mathcal{A}.$$

Para (2),

$$A \setminus B = A \cap B^c \in \mathcal{A}$$

lo que concluye la prueba. $\square$

Es claro que cualquier partición para Ω genera un álgebra. La intención de introducir la definición de álgebra es dar formalidad a la definición de tiempo de paro óptimo, pero no es un concepto que utilizaremos mucho en realidad.

Definición 4.1.3. *Un **tiempo de paro** es una función*

$$\tau : \Omega \rightarrow \{0, 1, \dots, T\}$$

tal que el conjunto $\{\tau \leq k\}$ pertenece al álgebra generada por la partición $\mathcal{P}_k$. Si $\tau(\omega) = k$, decimos que el evento ω fue interrumpido en el tiempo k .

Veamos la idea detrás de la definición. Pedimos que $\{\tau \leq k\}$ pertenezca al álgebra generada por la partición $\mathcal{P}_k$ por que si un evento $\omega \in \Omega$ es interrumpido en el tiempo k , entonces todos los que coinciden con él hasta la k -ésima entrada también tienen que haber sido interrumpidos en el mismo tiempo, (i.e. si algún elemento es interrumpido en el tiempo k , todos los que pertenezcan al mismo bloque en $\mathcal{P}_k$ también tienen que ser interrumpidos en el tiempo k). pues en el tiempo k no podemos distinguir dos elementos que difieran en la coordenada $k + 1$. Sería tan absurdo como decir "si mañana va a subir la acción, entonces ejerce la opción mañana".

Consideremos ahora qué pasa cuando tenemos un tiempo de paro. Entonces, para cada ω la acción será ejercida en el tiempo $\tau(\omega)$, y la ganancia en dicho tiempo es entonces:

$$Y_{\tau(\omega)}(\omega)$$

donde

$$Y_k := \max\{S_k - 21, 0\}.$$

A esta función sobre Ω la denotaremos como Y_τ . Decimos que Y_τ es el **valor final del proceso** según el tiempo τ .

Ejemplo 4.1.4. En el ejemplo de modelo CRR que estamos considerando, sea τ el tiempo de paro dado por:

$$\tau(\omega) = \begin{cases} \min\{\{k|S_k > 20\}, \{k|S_k < 17\}\} & \text{si este mínimo es menor a 3} \\ 3 & \text{en otro caso.} \end{cases}$$

más explícitamente,

$$\tau(\omega) = \begin{cases} 1 & \omega \in B_{1,1}, \\ 3 & \omega = \omega_5, \omega_6, \\ 2 & \omega = \omega_7, \omega_8. \end{cases}$$

Entonces, el valor final Y_τ para este tiempo de paro es:

$$Y_\tau(\omega) = \begin{cases} 1 & \omega \in B_{1,1}, \\ 0,78 & \omega = \omega_5, \\ 0 & \text{en otro caso.} \end{cases}$$

4.2. La envolvente de Snell y tiempos de paro óptimo

Supongamos que estamos en el tiempo cero. En este momento, no tenemos ninguna información sobre el mercado, así que lo más sensato que podemos hacer es escoger una estrategia o tiempo de paro que maximice nuestra ganancia esperada, es decir, buscamos τ^* tal que:

$$\mathcal{E}(Y_{\tau^*}) = \max_{\tau \geq 0} \mathcal{E}(Y_\tau).$$

Como estamos considerando que τ toma un conjunto finito de valores y está definida sobre un conjunto finito, este máximo existe. Decimos entonces que τ^* es un **tiempo de paro óptimo**.

Para nuestro caso, en total disponemos de $3^8 = 6561$ funciones. De ahí, hay que escoger las funciones que correspondan a un tiempo de paro. Después, evaluarlas y encontrar los máximos. Es muy fácil convencerse de que es posible encontrar tiempos de paro óptimo, pero, como ya comentamos en la introducción, no necesariamente es muy fácil encontrarlos.

Para el caso general, en el que disponemos de información sobre los días anteriores, consideraremos la siguiente definición.

Definición 4.2.1. Sea $\mathbb{Y} = (Y_0, Y_1, \dots, Y_T)$ un proceso estocástico. Consideramos entonces el proceso $\mathbb{U}$ dado por:

$$U_k = \max_{\tau \geq k} \mathcal{E}(Y_\tau | \mathcal{P}_k).$$

El proceso $\mathbb{U}$ se llama **envolvente de Snell** para $\mathbb{Y}$. La función de tiempo de paro donde se alcanza dicho máximo se llama **tiempo de paro óptimo** para $\mathbb{Y}$ sobre $\{k, \dots, T\}$.

Esto es, en cada tiempo, la envolvente de Snell representa las ganancias máximas esperadas a partir de ese tiempo. Trabajaremos más con la envolvente de Snell más adelante, antes demostraremos el siguiente teorema.

Teorema 4.2.2. *La envolvente de Snell satisface:*

$$U_k = \max\{Y_k, \max_{\tau \geq k+1} \mathcal{E}(Y_\tau | \mathcal{P}_k)\}$$

para todo $0 \leq k \leq T$.

Demostración. Por definición, tenemos que:

$$\begin{aligned} U_k &= \max_{\tau \geq k} \mathcal{E}(Y_\tau | \mathcal{P}_k) \\ &\geq \max\{\mathcal{E}(Y_k | \mathcal{P}_k), \max_{\tau \geq k+1} \mathcal{E}(Y_\tau | \mathcal{P}_k)\} \\ &= \max\{Y_k, \max_{\tau \geq k+1} \mathcal{E}(Y_\tau | \mathcal{P}_k)\}. \end{aligned}$$

Para la desigualdad contraria, consideraremos $\tau'(\omega)$ de la siguiente manera:

$$\tau'(\omega) = \max\{\tau, k+1\}.$$

Lo único que estamos haciendo es posponer los eventos que se detenían en el tiempo k al tiempo $k+1$. Para verificar que esto es también un tiempo de paro válido, notemos que:

$$\{\tau = k\} \cup \{\tau = k+1\} = \{\tau' = k+1\}.$$

Ahora observemos que el conjunto

$$\{\tau = k\} \cup \{\tau = k+1\}$$

también está en el álgebra generada por $\mathcal{P}_{k+1}$, pues $\{\tau = k\}$ pertenece al álgebra generada por $\mathcal{P}_k$, y $\mathcal{P}_{k+1}$ es un refinamiento de $\mathcal{P}_k$, se sigue que τ' es un tiempo de paro válido.

Además, tenemos que, como $\{\tau > k\} = \{\tau = k\}^c$, entonces

$$\begin{aligned}
\mathcal{E}(Y_\tau|\mathcal{P}_k) &= \mathcal{E}(Y_\tau 1_{\{\tau=k\}}|\mathcal{P}_k) + \mathcal{E}(Y_\tau 1_{\{\tau>k\}}|\mathcal{P}_k) \\
&= \mathcal{E}(Y_\tau 1_{\{\tau=k\}}|\mathcal{P}_k) + \mathcal{E}(Y_{\tau'} 1_{\{\tau>k\}}|\mathcal{P}_k) \\
&\leq \max\{Y_k, \mathcal{E}(Y_{\tau'}|\mathcal{P}_k)\} \\
&\leq \max\{Y_k, \max_{\sigma \geq k+1} \mathcal{E}(Y_\sigma|\mathcal{P}_k)\}.
\end{aligned}$$

Como la desigualdad es válida para todo τ , tenemos que:

$$U_k \leq \max\{Y_k, \max_{\sigma \geq k+1} \mathcal{E}(Y_\sigma|\mathcal{P}_k)\}$$

que es la desigualdad que queríamos. □

Observemos que, de acuerdo al teorema anterior:

$$U_N = \max\{Y_N, \max_{\tau \geq N} \mathcal{E}(Y_\tau|\mathcal{P}_N)\}$$

de donde

$$U_N = Y_N.$$

Ahora consideremos

$$\max_{\tau \geq k+1} \mathcal{E}(Y_\tau|\mathcal{P}_k).$$

Observemos que si la esperanza condicional fuera con respecto a $\mathcal{P}_{k+1}$, ya no tendríamos mucho que hacer, puesto que este máximo sería simplemente U_{k+1} . Esto nos sugiere utilizar la propiedad de las torres para convertirlo en algo útil.

La siguiente relación de recurrencia en la envolvente de Snell es quizá la propiedad más importante. El tener una manera de calcular recursivamente la envolvente de Snell nos permitirá calcular los dos tiempos de paro óptimo que ya hemos mencionado antes.

Teorema 4.2.3. *La envolvente de Snell satisface:*

$$(1) \ U_N = Y_N.$$

$$(2) \ U_k = \max\{Y_k, \mathcal{E}(U_{k+1}|\mathcal{P}_k)\}.$$

para todo $k = 0, 1, \dots, N-1$.

Demostración. Lo único que tenemos que demostrar en realidad es:

$$\mathcal{E}(U_{k+1}|\mathcal{P}_k) = \max_{\tau \geq k+1} \mathcal{E}(Y_\tau|\mathcal{P}_k).$$

Para esto tenemos, por la propiedad de las torres:

$$\begin{aligned} \max_{\tau \geq k+1} \mathcal{E}(Y_\tau|\mathcal{P}_k) &= \max_{\tau \geq k+1} \mathcal{E}(\mathcal{E}(Y_\tau|\mathcal{P}_{k+1})|\mathcal{P}_k) \\ &\leq \mathcal{E}(\max_{\tau \geq k+1} \mathcal{E}(Y_\tau|\mathcal{P}_{k+1})|\mathcal{P}_k) \\ &= \mathcal{E}(U_{k+1}|\mathcal{P}_k) \end{aligned}$$

La segunda desigualdad es gracias al lema (??).

Para probar la desigualdad contraria, definamos τ^* como el tiempo de paro óptimo que satisface:

$$U_{k+1} = \max_{\tau \geq k+1} \mathcal{E}(Y_\tau|\mathcal{P}_{k+1}) = \mathcal{E}(Y_{\tau^*}|\mathcal{P}_{k+1}).$$

Se tiene que:

$$\begin{aligned} \mathcal{E}(U_{k+1}|\mathcal{P}_k) &= \mathcal{E}(\mathcal{E}(Y_{\tau^*}|\mathcal{P}_{k+1})|\mathcal{P}_k) \\ &= \mathcal{E}(Y_{\tau^*}|\mathcal{P}_k) \\ &\leq \max_{\tau \geq k+1} \mathcal{E}(Y_\tau|\mathcal{P}_k). \end{aligned}$$

□

La envolvente de Snell tiene la siguiente interpretación en términos económicos: representa la ganancia máxima esperada si a partir del tiempo k empezamos a seguir una estrategia (i.e. tiempo de paro) óptima.

Observación 4.2.4. *La envolvente de Snell de (Y_k) es un proceso $\mathcal{P}_k$ -medible para todo k .*

Demostración. Observemos que U_k está definido como el máximo sobre un número finito de variables aleatorias $\mathcal{P}_k$ -medibles. □

Ejemplo 4.2.5. *Ya hemos calculado Y_k en base al diagrama de la figura (??). Para este ejemplo, haremos referencia a esos valores. Tenemos que $U_3 = Y_3$ por la relación de recurrencia que encontramos.*

Ahora, necesitamos calcular:

$$\mathcal{E}(U_3|\mathcal{P}_2) = \mathcal{E}(Y_3|\mathcal{P}_2) = \begin{cases} \frac{1}{2}(5,26 + 0,78) = 3,2 & \omega \in B_{2,1} \\ \frac{1}{2}(0,78 + 0) = 0,39 & \omega \in B_{2,2} \\ \frac{1}{2}(0,78 + 0) = 0,39 & \omega \in B_{2,3} \\ 0 & \omega \in B_{2,4} \end{cases}$$

de donde obtenemos que:

$$U_2 = \max\{Y_2, \mathcal{E}(U_3|\mathcal{P}_2)\} = \begin{cases} 3,2 & \omega \in B_{2,1} \\ 0,39 & \omega \in B_{2,2} \\ 0,39 & \omega \in B_{2,3} \\ 0 & \omega \in B_{2,4} \end{cases}$$

Ahora necesitamos:

$$\mathcal{E}(U_2|\mathcal{P}_1) = \begin{cases} \frac{1}{2}(3,2 + 0,39) = 1,795 & \omega = \omega_1, \omega_2, \omega_3, \omega_4 \\ \frac{1}{2}(0,39 + 0) = 0,195 & \omega = \omega_5, \omega_6, \omega_7, \omega_8 \end{cases}$$

entonces,

$$U_1 = \max\{Y_1, \mathcal{E}(U_2|\mathcal{P}_1)\} = \begin{cases} \frac{1}{2}(3,2 + 0,39) = 1,795 & \omega = \omega_1, \omega_2, \omega_3, \omega_4 \\ \frac{1}{2}(0,39 + 0) = 0,195 & \omega = \omega_5, \omega_6, \omega_7, \omega_8 \end{cases}$$

Por último,

$$U_0 = \max\{Y_0, \mathcal{E}(U_1|\mathcal{P}_0)\} = \max\{0, \frac{1}{2}(1,795 + 0,195)\} = 0,995$$

Corolario 4.2.6. *La envolvente de Snell es la menor supermartingala que domina a $\mathbb{Y}$.*

Demostración. De la relación de recurrencia anterior, tenemos claramente que:

$$U_k = \max\{Y_k, \mathcal{E}(U_{k+1}|\mathcal{P}_k)\} \geq \mathcal{E}(U_{k+1}|\mathcal{P}_k),$$

que es la condición para que (U_k) sea supermartingala.

Sea V_k una supermartingala que domine a Y_k , esto es:

$$V_k \geq Y_k.$$

De la definición de supermartingala, tenemos que:

$$V_k \geq \mathcal{E}(V_{k+1}|\mathcal{P}_k),$$

esto es:

$$V_k \geq \max\{Y_k, \mathcal{E}(V_{k+1}|\mathcal{P}_k)\}.$$

Ya que V_N domina a Y_N y $Y_N = U_N$, tenemos que $V_N \geq U_N$. Hagamos inducción al revés. Entonces, nuestra hipótesis de inducción es que $V_{k+1} \geq U_{k+1}$. Basta ver que $V_k \geq U_k$. Esto se sigue de lo siguiente:

$$\begin{aligned} V_k &\geq \max\{Y_k, \mathcal{E}(V_{k+1}|\mathcal{P}_k)\} \\ &\geq \max\{Y_k, \mathcal{E}(U_{k+1}|\mathcal{P}_k)\} \\ &= U_k. \end{aligned}$$

Entonces, la supermartingala (V_k) domina a (U_k) , de donde la envolvente de Snell es justamente la menor supermartingala que domina al proceso $\mathbb{Y}$ y se termina nuestra prueba. $\square$

4.3. Procesos interrumpidos

Aún no es claro cómo utilizaremos la envolvente de Snell para escoger tiempos de paro óptimos. Para cada elección posible de un τ , nuestras ganancias finales tienen un cierto valor. Nuestra elección de dicho τ puede basarse naturalmente en la información hasta un cierto tiempo, lo cual está indicado por la presencia de un k en los resultados obtenidos hasta ahora. Por otro lado, una vez que un evento ha sido interrumpido, si fue interrumpido en un tiempo menor que el final, su ganancia no cambiará, i.e. si ya ejercimos la opción hoy, mañana no ganamos ni perdemos nada por ello. Consideremos la siguiente familia de funciones: Denotaremos $T := t_N$ a partir de ahora.

$$Y_k^\tau(\omega) := \begin{cases} Y_k(\omega) & k \leq \tau(\omega), \\ Y_{\tau(\omega)}(\omega) & k \geq \tau(\omega) \end{cases}$$

donde $\tau(\omega)$ es un tiempo de paro. Las funciones Y_k^τ se llaman **proceso detenido**.

Observemos que podemos escribir:

$$Y_k^\tau = Y_k 1_{\{\tau \geq k\}} + \sum_{s=0}^{k-1} Y_s 1_{\{\tau=s\}}$$

porque si un evento se interrumpe en el tiempo $s < k$, de acuerdo a la definición lo que hacemos es tomar $Y_s(\omega)$. Esto está representado en el término $\sum_{s=0}^{k-1} Y_s 1_{\{\tau=s\}}$ para cada uno de los s posibles. Si el tiempo de interrupción es más grande que k , simplemente consideraremos evaluar a este elemento en $Y_k(\omega)$.

Lema 4.3.1. *Si el proceso (Y_k) es adaptado, entonces el proceso interrumpido (Y_k^τ) también lo es.*

Demostración. Sea $k \in \{0, 1, \dots, T\}$. Entonces,

$$\{\tau \geq k\} = \{\tau \leq k-1\}^c$$

que está en el álgebra generada por $\mathcal{P}_{k-1}$.

Además, para $s \leq k-1$, $\{\tau = s\}$ también pertenece al álgebra generada por $\mathcal{P}_{k-1}$, pues esta partición es un refinamiento de la partición $\mathcal{P}_s$. Esto significa que las funciones $1_{\{\tau \geq k\}}, 1_{\{\tau = s\}}$ son $\mathcal{P}_{k-1}$ -medibles y por tanto $\mathcal{P}_k$ -medibles. Como Y_k también es $\mathcal{P}_k$ -medible, concluimos que el proceso interrumpido también lo es. $\square$

Presentamos ahora un resultado muy importante, que será clave en la caracterización de tiempos de paro óptimos.

Teorema 4.3.2. *(Teorema de interrupción de Doob) Sea $\mathbb{U}$ supermartingala y τ un tiempo de paro óptimo. El proceso interrumpido $\mathbb{U}^\tau := (U_k^\tau)$ también es supermartingala. El teorema también es cierto para martingalas y submartingalas.*

Demostración. Por el lema anterior sabemos que U_k^τ es $\mathcal{P}_k$ -medible. Enseguida probaremos la condición para supermartingalas. La prueba para los otros dos casos es análoga.

Tenemos que:

$$\mathcal{E}(U_{k+1}^\tau | \mathcal{P}_k) = \mathcal{E}(U_{k+1} 1_{\{\tau \geq k+1\}} | \mathcal{P}_k) + \sum_{s=0}^k \mathcal{E}(U_s 1_{\{\tau = s\}} | \mathcal{P}_k).$$

Como $U_s 1_{\{\tau = s\}}$ es $\mathcal{P}_s$ -medible para $s \leq k$, tenemos que

$$\mathcal{E}(U_s 1_{\{\tau = s\}} | \mathcal{P}_k) = U_s 1_{\{\tau = s\}}.$$

Además, ya que $1_{\{\tau \geq k+1\}}$ es $\mathcal{P}_k$ -medible,

$$\mathcal{E}(U_{k+1} 1_{\{\tau \geq k+1\}} | \mathcal{P}_k) = \mathcal{E}(U_{k+1} | \mathcal{P}_k) 1_{\{\tau \geq k+1\}}$$

esto puede verse fácilmente por bloques, ya que la función $1_{\{\tau \geq k+1\}}$ es constante en cada uno de los bloques y por linealidad de $\mathcal{E}(\cdot | \mathcal{P})$ estas constantes pueden salir.

Por tanto,

$$\begin{aligned}
\mathcal{E}(U_{k+1}^\tau | \mathcal{P}_k) &= \mathcal{E}(U_{k+1} | \mathcal{P}_k) 1_{\{\tau \geq k+1\}} + \sum_{s=0}^k U_s 1_{\{\tau=s\}} \\
&\leq U_k 1_{\{\tau \geq k+1\}} + \sum_{s=0}^k U_s 1_{\{\tau=s\}} \\
&= U_k 1_{\{\tau \geq k+1\}} + U_k 1_{\{\tau=k\}} + \sum_{s=0}^{k-1} U_s 1_{\{\tau=s\}} \\
&= U_k 1_{\{\tau \geq k\}} + \sum_{s=0}^{k-1} U_s 1_{\{\tau=s\}} \\
&= U_k^\tau
\end{aligned}$$

utilizando la propiedad de que (U_k) es supermartingala para obtener la desigualdad. $\square$

Este resultado tiene una interpretación muy interesante. No importa el tiempo de paro óptimo que escojamos (i.e. la estrategia que usemos para decidir si jugar o no seguir jugando), un juego justo será siempre un juego justo, en el sentido de juegos justos que comentamos anteriormente. De la misma manera, un juego desfavorable no se volverá repentinamente favorable bajo ninguna estrategia. Lo que sí podemos hacer es convertir un juego desfavorable en uno justo. Consideremos por ejemplo el caso del juego en donde primero se juega una moneda justa y después una moneda injusta. Este juego es desfavorable si se juega durante un tiempo grande. Pero una estrategia de paro es abandonar el juego en cuanto se cambia de moneda, volviéndose de esta manera un juego justo (aunque posiblemente muy corto).

Como hemos probado antes, la envolvente de Snell es supermartingala. Por el teorema anterior, no podemos aumentar nuestras esperanzas de ganancia con ninguna estrategia. Lo que sí podemos hacer es escoger una estrategia para protegerlas y evitar que disminuyan.

Utilizando toda la información de la que disponemos, ya estamos listos para caracterizar los tiempos de paro óptimos.

4.4. Caracterización de tiempos óptimos

Por el teorema de paro de Doob, es natural pensar que, ya que ningún tiempo de paro puede convertir la envolvente de Snell en submartingala, lo mejor que puede pasar es que se vuelva martingala (i.e. que la igualdad se alcance). También es natural pensar que si empleamos un tiempo de paro óptimo, las ganancias obtenidas serán mayores o iguales

a las ganancias máximas esperadas. Estas dos afirmaciones caracterizan a los tiempos de paro óptimos.

Teorema 4.4.1. *Sea (Y_k) la sucesión de variables aleatorias dada por:*

$$Y_k := \max\{S_k - 21, 0\}.$$

Un tiempo de paro τ es óptimo para el intervalo $[0, T]$ sí y sólo sí:

(i) (U_k^τ) es martingala

(ii) $U_\tau = Y_\tau$.

Por otra parte, si τ es un tiempo de paro óptimo, tenemos que:

$$U_0 = \mathcal{E}(Y_\tau | \mathcal{P}_0).$$

Demostración. Supongamos (i) y (ii) para demostrar que τ es tiempo de paro óptimo. Tenemos que:

$$U_0 = U_0^\tau = \mathcal{E}(U_T^\tau | \mathcal{P}_0) = \mathcal{E}(U_\tau | \mathcal{P}_0) = \mathcal{E}(Y_\tau | \mathcal{P}_0).$$

La primera igualdad se tiene por la definición de proceso interrumpido (en este caso, todos los eventos están obligados a interrumpirse al tiempo cero). La segunda igualdad se da gracias a que el proceso interrumpido es también martingala, mientras que la tercera igualdad la tenemos otra vez por definición de proceso interrumpido (todos los procesos se interrumpen a su propio tiempo de interrupción). La cuarta igualdad se da por hipótesis.

Ya demostramos que U_k es la menor supermartingala que domina a Y_k . Entonces tenemos que para un tiempo de paro σ :

$$\mathcal{E}(Y_\sigma | \mathcal{P}_0) = \mathcal{E}(Y_T^\sigma | \mathcal{P}_0) = \mathcal{E}(U_T^\sigma | \mathcal{P}_0) \leq U_0^\sigma = U_0.$$

Se sigue entonces que para cualquier σ ,

$$\mathcal{E}(Y_\sigma) = \mathcal{E}(Y_\sigma | \mathcal{P}_0) \leq U_0 = \mathcal{E}(Y_\tau | \mathcal{P}_0) = \mathcal{E}(Y_\tau)$$

lo que prueba la condición de optimalidad para τ .

Ahora supongamos que tenemos τ tiempo de paro óptimo. Primero probaremos que $U_\tau = Y_\tau$. Observemos primero que, como (U_k) domina a (Y_k) , $Y_\tau \leq U_\tau$. De la optimalidad de τ y la definición de U_0 obtenemos que:

$$U_0 = \mathcal{E}(Y_\tau | \mathcal{P}_0) \leq \mathcal{E}(U_\tau | \mathcal{P}_0).$$

Ya que (U_k^τ) es supermartingala obtenemos que:

$$U_0 = U_0^\tau \geq \mathcal{E}(U_T^\tau | \mathcal{P}_0) = \mathcal{E}(U_\tau | \mathcal{P}_0).$$

Estas dos desigualdades implican

$$\mathcal{E}(U_\tau | \mathcal{P}_0) = U_0 = \mathcal{E}(Y_\tau | \mathcal{P}_0).$$

Pero ya que $U_\tau \geq Y_\tau$ se sigue que

$$U_\tau = Y_\tau.$$

Esto es gracias al lema (??).

Por el teorema de paro de Doob, sabemos que el proceso interrumpido (U_k^τ) es supermartingala. Entonces tenemos que:

$$\begin{aligned} U_0 &= U_0^\tau \geq \mathcal{E}(U_k^\tau | \mathcal{P}_0) \\ &\geq \mathcal{E}(\mathcal{E}(U_{k+1}^\tau | \mathcal{P}_k) | \mathcal{P}_0) \\ &= \mathcal{E}(U_{k+1}^\tau | \mathcal{P}_0) \\ &\geq \mathcal{E}(\mathcal{E}(U_T^\tau | \mathcal{P}_{k+1}) | \mathcal{P}_0) \\ &= \mathcal{E}(U_T^\tau | \mathcal{P}_0) = \mathcal{E}(U_\tau | \mathcal{P}_0) \\ &= \mathcal{E}(Y_\tau | \mathcal{P}_0) \\ &= U_0. \end{aligned}$$

De eso deducimos que:

$$\mathcal{E}(U_k^\tau | \mathcal{P}_0) = \mathcal{E}(\mathcal{E}(U_{k+1}^\tau | \mathcal{P}_k) | \mathcal{P}_0).$$

Pero ya que $U_k^\tau \geq \mathcal{E}(U_{k+1}^\tau | \mathcal{P}_k)$ se sigue de la igualdad anterior que

$$U_k^\tau = \mathcal{E}(U_{k+1}^\tau | \mathcal{P}_k)$$

lo que prueba que (U_k^τ) es martingala. □

El término cero de la demostración anterior no se refiere a que sólo puedan caracterizarse los tiempos óptimos a partir del momento inicial (cuando se compra la opción). Aparece simplemente por razones técnicas. En la práctica, podemos simplemente llamar tiempo cero al tiempo t a partir de que compramos la opción y tomar el valor de la acción en ese momento como el valor inicial, no el verdadero valor inicial.

Dada esta caracterización, es natural tratar de encontrar tiempos de paro óptimo. Esto es lo que haremos enseguida.

El menor tiempo de paro óptimo para el que $U_\tau = Y_\tau$ está definido por:

$$\tau_{min}(\omega) := \begin{cases} \min\{k; Y_k(\omega) = U_k(\omega)\} & \text{si } \{k; Y_k(\omega) = U_k(\omega)\} \neq \emptyset \\ T & \text{en otro caso.} \end{cases}$$

Proposición 4.4.2. *El tiempo de paro τ_{min} es el menor tiempo de paro óptimo para (Y_k) .*

Demostración. Utilizaremos los criterios que acabámos de encontrar para demostrar la optimalidad de $\tau := \tau_{min}$. De la definición de τ , es inmediato que $Y_\tau = U_\tau$. Para probar que (U_k^τ) es martingala, observemos que

$$U_k^\tau = U_k 1_{\{\tau \geq k\}} + \sum_{s=0}^{k-1} U_s 1_{\{\tau=s\}}.$$

Entonces,

$$\mathcal{E}(U_{k+1}^\tau | \mathcal{P}_k) = \mathcal{E}(U_{k+1} | \mathcal{P}_k) 1_{\{\tau \geq k+1\}} + \sum_{s=0}^k U_s 1_{\{\tau=s\}}$$

donde hemos utilizado otra vez el hecho de que para $0 \leq s \leq k$, las variables aleatorias $U_s 1_{\{\tau=s\}}$ y $1_{\{\tau \geq k+1\}}$ son $\mathcal{P}_k$ -medibles.

En el conjunto $\{\tau \leq k\} = \{\tau \geq k+1\}^c$ es claro que

$$\begin{aligned} \mathcal{E}(U_{k+1}^\tau | \mathcal{P}_k) &= \mathcal{E}(U_{k+1} | \mathcal{P}_k) 1_{\{\tau \geq k+1\}} + \sum_{s=0}^k U_s 1_{\{\tau=s\}} \\ &= \sum_{s=0}^k U_s 1_{\{\tau=s\}} \\ &= U_k 1_{\{\tau=k\}} + \sum_{s=0}^{k-1} U_s 1_{\{\tau=s\}} \\ &= U_k 1_{\{\tau \geq k\}} + \sum_{s=0}^{k-1} U_s 1_{\{\tau=s\}} \\ &= U_k^\tau. \end{aligned}$$

Por otro lado, de la definición de τ tenemos que en el conjunto $\{\tau \geq k+1\}$ se tiene que:

$$Y_k \neq U_k.$$

Pero

$$U_k = \max\{Y_k, \mathcal{E}(U_{k+1}|\mathcal{P}_k)\},$$

por tanto,

$$U_k = \mathcal{E}(U_{k+1}|\mathcal{P}_k).$$

Utilizando esto y la segunda igualdad que escribimos, concluimos que en este conjunto:

$$\begin{aligned} \mathcal{E}(U_{k+1}^\tau|\mathcal{P}_k) &= U_k 1_{\{\tau \geq k+1\}} + \sum_{s=0}^k U_s 1_{\{\tau=s\}} \\ &= U_k 1_{\{\tau \geq k\}} + \sum_{s=0}^{k-1} U_s 1_{\{\tau=s\}} \\ &= U_k^\tau. \end{aligned}$$

Esto prueba entonces que $\tau = \tau_{min}$ es tiempo de paro óptimo. De la manera en que lo definimos, se sigue que es entonces el menor tiempo de paro óptimo. $\square$

Hay que señalar un detalle. Como tenemos que $U_T = Z_T$, sería natural pensar que T pudiera ser el tiempo de paro óptimo máximo. Sin embargo, no es así. Hay que hacer una consideración más respecto a la envolvente de Snell para poder garantizar la segunda condición de nuestra caracterización para tiempos de paro óptimos. De la misma manera, tampoco podemos pedir una simetría en las condiciones, como pensar en que el tiempo de paro óptimo es el último en el que $U_\tau = Y_\tau$, pues esto se alcanza en $U_T = Y_T$.

Para encontrar el tiempo de paro óptimo máximo, necesitamos descomponer la envolvente de Snell en el sentido de Doob:

$$U_k = M_k - T_k$$

donde M_k es martingala y T_k es un proceso predecible y no decreciente.

Lema 4.4.3. *Sea $\mathbb{U} = (U_k)$ la envolvente de Snell para $\mathbb{Y} = (Y_k)$. Entonces, para un tiempo de paro óptimo τ , el proceso interrumpido $\mathbb{U}_k^\tau = (U_k^\tau)$ es martingala sí y sólo sí $\mathbb{T}_\tau = 0$, en donde $\mathbb{T}_\tau$ es el proceso predecible en la descomposición de Doob para $\mathbb{U}$.*

Demostración. Si escribimos

$$U_k = M_k - T_k$$

tenemos, al interrumpir en el tiempo τ que

$$U_k^\tau = M_k^\tau - T_k^\tau.$$

Por la unicidad de la descomposición de Doob, tenemos que $\mathbb{U}^\tau$ es martingala sí y sólo sí $\mathbb{T}^\tau$ es el proceso cero.

Como $T_0 = 0$ y (T_k) es no decreciente, el proceso interrumpido (T_k^τ) es cero sí y sólo sí $\mathbb{T}_\tau = 0$, donde $[T_\tau](\omega) = T_{\tau(\omega)}(\omega)$.

□

Teorema 4.4.4. *El tiempo de paro óptimo máximo está dado por:*

$$\tau_{max}(\omega) = \begin{cases} \min\{k | T_{k+1}(\omega) > 0\} & \text{si } \{k | T_{k+1}(\omega) > 0\} \neq \emptyset \\ T & \text{en otro caso.} \end{cases}$$

donde $U_k = M_k - T_k$ es la descomposición de Doob.

Demostración. Es claro que esta función es un tiempo de paro, ya que el proceso $\mathbb{T}$ es tal que T_k es $\mathcal{P}_{k-1}$ -medible, por tanto es $\mathcal{P}_k$ -medible, por lo que si algún ω es interrumpido en el tiempo k , todos los elementos que pertenecen al mismo bloque en $\mathcal{P}_{k-1}$ también se interrumpen en el tiempo k . Ya que cualquier tiempo óptimo σ satisface $T_\sigma = 0$ entonces, si τ es tiempo de paro óptimo se sigue que por la manera en que lo hemos construido tiene que ser el máximo.

Para ver que es un tiempo de paro óptimo basta entonces demostrar que $U_\tau = Y_\tau$. Por la relación de recurrencia que teníamos,

$$U_k = \max\{Y_k, \mathcal{E}(U_{k+1} | \mathcal{P}_k)\}.$$

Por otra parte, utilizando la descomposición de Doob,

$$\mathcal{E}(U_{k+1} | \mathcal{P}_k) = M_k - T_{k+1} = U_k + T_k - T_{k+1}.$$

Entonces la relación de recurrencia puede escribirse como

$$U_k = \max\{Y_k, U_k - (T_{k+1} - T_k)\}.$$

Al tomar $k = \tau(\omega)$

$$T_{\tau+1}(\omega) - T_{\tau}(\omega) = T_{\tau+1}(\omega) > 0.$$

Como esto es independiente de ω ,

$$U_{\tau} = Y_{\tau}$$

como queríamos demostrar. □

Ejemplo 4.4.5. *Calculemos un tiempo de paro óptimo para nuestro modelo CRR. Ya hemos calculado la envolvente de Snell para este proceso. Observemos que el tiempo de paro óptimo σ dado por:*

$$\sigma = \begin{cases} 2 & \omega = \omega_1, \omega_2, \omega_7, \omega_8 \\ 3 & \text{en otro caso.} \end{cases}$$

es un tiempo de paro óptimo mínimo, ya que es la primera vez que $U_{\tau} = Y_{\tau}$. Podemos construir este ejemplo de tiempo de paro óptimo mínimo, a partir de la envolvente de Snell que ya construimos en el ejemplo anterior. Es fácil observar, a partir del ejemplo anterior, que σ tiene la propiedad siguiente

$$\sigma(\omega) \in \min\{k | Y_k(\omega) = U_k(\omega)\}$$

por lo que es el menor tiempo de paro óptimo.

Ejemplo 4.4.6. *Modificando un poco nuestro ejemplo anterior, mostraremos un tiempo de paro máximo no trivial. El diagrama (??) muestra entre paréntesis los precios descontados. Consideremos ahora $K = 20,5$, $r = 0,1$ La probabilidad p está dada por*

$$p = \frac{r - d}{u - d} \approx 0,35$$

Los precios del activo sin riesgo son:

$$A_0 = 1, A_1 = 1,05, A_2 = 1,10, A_3 = 1,16.$$

Por simplicidad, no colocaremos tilde a los valores descontados correspondientes. Tenemos entonces que:

$$Y_3 = \begin{cases} 2,45 & \omega = \omega_1, \\ 0 & \text{en otro caso.} \end{cases}$$

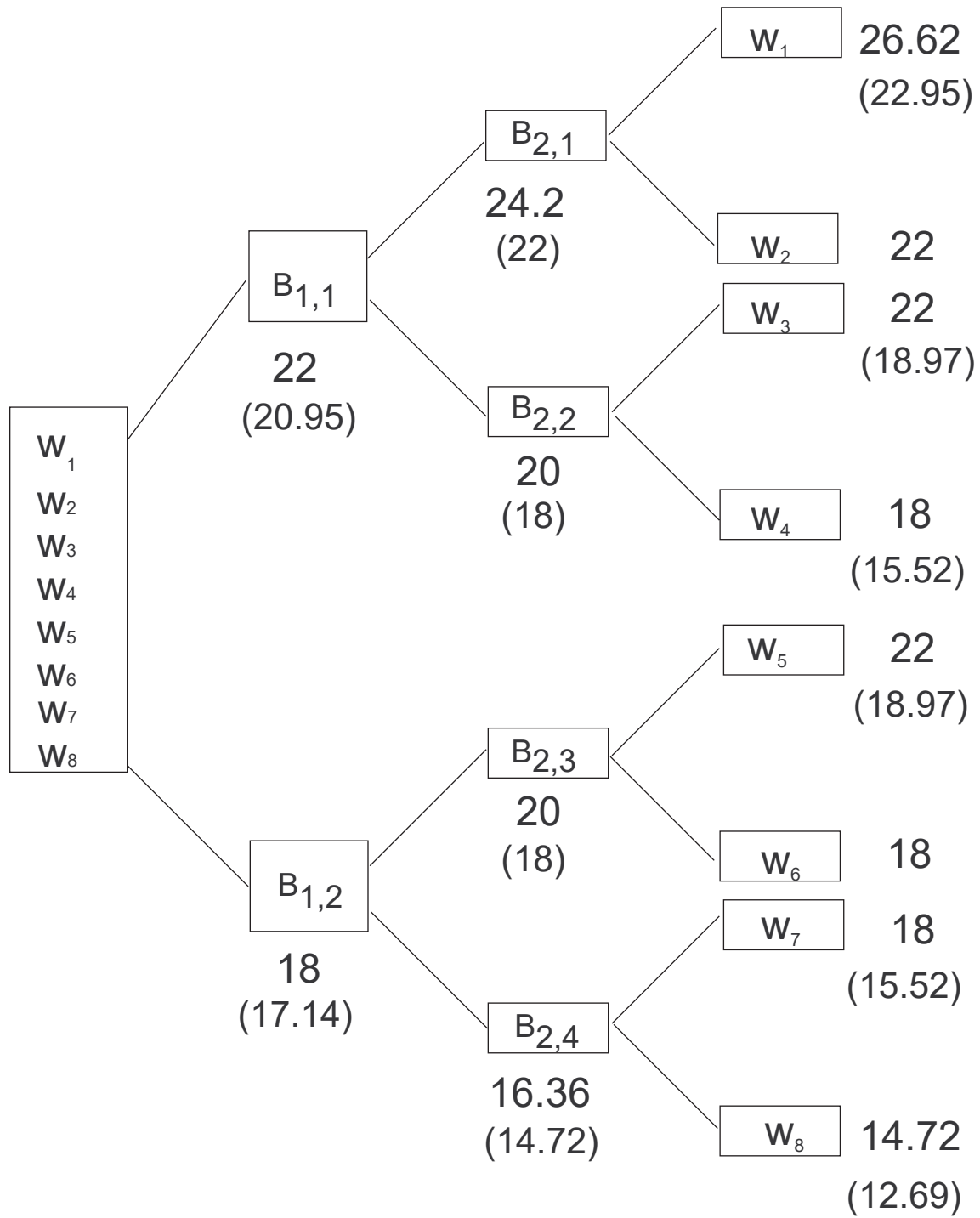


Figura 4.1: Diagrama CRR en el caso N=3 con descuentos

$$Y_2 = \begin{cases} 1,5 & \omega \in B_{2,1}, \\ 0 & \text{en otro caso.} \end{cases}$$

y además, $Y_1 = Y_0 = 0$.

Ahora, para construir la envolvente de Snell, consideramos $U_3 = Y_3$ y observemos que:

$$\mathcal{E}(U_3|\mathcal{P}_2) = \mathcal{E}(Y_3|\mathcal{P}_2) = \begin{cases} 0,35 \cdot 2,45 = 0,86 & \omega \in B_{2,1} \\ 0 & \text{en otro caso} \end{cases}$$

de donde obtenemos que $U_2 = Y_2$. Ahora necesitamos:

$$\mathcal{E}(U_2|\mathcal{P}_1) = \begin{cases} 0,35 \cdot 1,5 = 0,53 & \omega = \omega_1, \omega_2, \omega_3, \omega_4 \\ 0 & \text{en otro caso} \end{cases}$$

entonces, como $Y_1 = 0$,

$$U_1 = \mathcal{E}(U_2|\mathcal{P}_1).$$

A partir de este momento, la envolvente de Snell se ha convertido en martingala. Entonces, el tiempo de paro óptimo máximo es ahora, para los eventos $\omega_1, \omega_2, \omega_3, \omega_4$. Por último,

$$U_0 = \max\{Y_0, \mathcal{E}(U_1|\mathcal{P}_0)\} = \max\{0, 0,35 \cdot 0,45\} = 0,16$$

4.5. La envolvente de Snell y los procesos de Markov

Consideremos un espacio de probabilidad finito $(\Omega, \mathbb{P})$ y un proceso estocástico (X_k) , con $0 \leq k \leq T$ y

$$X_k : \Omega \rightarrow E,$$

donde $E \subset \mathbb{R}$ es un conjunto finito. El proceso estocástico (X_k) es un **proceso de Markov** o **cadena de Markov** si

$$\mathbb{P}(X_{k+1} = a_{k+1} | X_0 = a_0, \dots, X_k = a_k) = \mathbb{P}(X_{k+1} = a_{k+1} | X_k = a_k)$$

para cualesquiera $a_0, a_1, \dots, a_{k+1} \in E$ tales que

$$\mathbb{P}(X_0 = a_0, \dots, X_k = a_k) > 0.$$

Observemos que esta última condición se impone para que las probabilidades condicionales de arriba tengan sentido (si no, la parte de la izquierda no tiene sentido). Si interpretamos X_k como una representación del estado de un sistema en el tiempo k , la propiedad de Markov requiere que las probabilidades de los estados posibles de X_{k+1} dependan exclusivamente del estado X_k y no de los estados previos.

Ejemplo 4.5.1. Consideremos el tablero de Monopoly (Turista). El estado del tablero en un momento dado solamente depende del estado anterior y del lanzamiento de los dados, y no depende en cómo los estados anteriores influyeron para llegar a este estado. Por lo tanto, este juego es un proceso de Markov. Ahora pensemos en un juego de pókar o de dominó. Un jugador atento puede recordar las cartas o las fichas que ya pasaron, y por tanto saber qué cartas o fichas ya no están disponibles para armar jugadas. En este caso, el estado del juego en el siguiente movimiento sí depende de todos los estados anteriores. Este último tipo de juegos no son un proceso de Markov.

Proposición 4.5.2. Sea (X_k) un proceso de Markov sobre $(\Omega, \mathbb{P})$. Entonces, para cualquier elección $a_0, a_1, \dots, a_k \in E$ se tiene la siguiente fórmula:

$$\mathbb{P}(X_0 = a_0, \dots, X_k = a_k) = \mathbb{P}(X_0 = a_0) \dots \mathbb{P}(X_k = a_k | X_{k-1} = a_{k-1}).$$

De acuerdo a esta fórmula, es claro que cualquier proceso estocástico que consista de variables aleatorias independientes es también un proceso de Markov.

Definición 4.5.3. Una cadena de Markov es **homogénea** si para cualesquiera $a, b \in E$ las probabilidades de transición

$$\mathbb{P}(X_{t+1} = b | X_t = a)$$

dependen únicamente de los valores de a, b pero no de t . Escribimos

$$\mathbb{P}(a, b) = \mathbb{P}(X_{t+1} = b | X_t = a)$$

en este caso.

A partir de ahora, consideraremos una cadena de Markov homogénea (X_t) , $t = 0, \dots, T$ y la sucesión de particiones $(\mathcal{P}_t)$ que induce de manera natural. Una caracterización para $\mathcal{P}_t$ que nos será útil es:

$$\mathcal{P}_t = \{A_{t,\mathbf{a}} | \mathbf{a} = (a_1, a_2, \dots, a_t) \in E^t\}$$

donde

$$A_{t,\mathbf{a}} = \{\omega \in \Omega | X_1 = a_1, X_2 = a_2, \dots, X_t = a_t\}.$$

Tenemos entonces el siguiente resultado importante.

Proposición 4.5.4. *Sea (X_t) una cadena de Markov homogénea. Entonces, para cualquier función $f : \mathbb{R} \rightarrow \mathbb{R}$ tenemos que*

$$\mathcal{E}[f(X_{t+1})|\mathcal{P}_t] = \sum_{b \in E} \mathbb{P}(X_t(\omega), b) f(b).$$

Demostración. Es claro que basta considerar la esperanza condicional con respecto a los bloques de la partición $\mathcal{P}_t$. Por la propiedad de las cadenas de Markov, tenemos que:

$$\mathbb{P}(X_{t+1} = b | A_{t,\mathbf{a}}) = \mathbb{P}(X_{t+1} = b | X_t = a_t) = \mathbb{P}(a_t, b).$$

Por otro lado, tenemos que

$$\mathcal{E}[f(X_{t+1})|A_{t,\mathbf{a}}] = \sum_{b \in E} \mathbb{P}(X_{t+1} = b | A_{t,\mathbf{a}}) f(b)$$

pero esta última expresión es

$$\sum_{b \in E} f(b) \mathbb{P}(a_t, b).$$

A partir de aquí, si cambiamos a_t por cada uno de los posibles valores de $X_t(\omega)$ podemos hacer esta misma construcción en cada uno de los bloques de la partición y de esta manera obtener el resultado que queríamos probar. $\square$

Ahora aplicaremos algo de lo que hemos visto sobre procesos de Markov a la envolvente de Snell.

Hemos visto que el valor de una acción depende del estado del mercado en ese día, es decir, existen funciones (f_t) tales que

$$Y_t = f_t(S_t),$$

con $f_t : \mathbb{R} \rightarrow \mathbb{R}$.

Probaremos un poco más con el siguiente teorema.

Teorema 4.5.5. *Existen funciones $h_t : E \rightarrow \mathbb{R}$ tales que:*

$$U_t = h_t(S_t)$$

donde (h_t) es una familia de funciones definida recursivamente por

$$\begin{aligned} h_T(s) &:= f_T(s), \\ h_t(s) &:= \max\{f_t(s), \sum_{b \in E} \mathbb{P}(s, b) h_{t+1}(b)\}. \end{aligned}$$

Demostración. Probaremos esto por inducción al revés. Para $t = T$ es obvio. Supongamos ahora que $U_{t+1} = h_{t+1}(S_{t+1})$ para algún $t \in \{0, 1, \dots, T-1\}$. El paso de inducción consiste entonces en encontrar h_t tal que $U_t = h_t(S_t)$. Por la recursividad de la envolvente de Snell tenemos que:

$$U_t(\omega) = \max\{Y_t(\omega), \mathcal{E}(U_{t+1}|\mathcal{P}_t)(\omega)\} = \max\{f_t(S_t(\omega)), \mathcal{E}(h_{t+1}(S_{t+1})|\mathcal{P}_t)(\omega)\}.$$

Por otra parte, probamos que:

$$\mathcal{E}[h_{t+1}(S_{t+1})|\mathcal{P}_t] = \sum_{b \in E} \mathbb{P}(S_t(\omega), b) h_{t+1}(b)$$

de donde se sigue

$$U_t(\omega) = \max\{f_t(S_t(\omega)), \sum_{b \in E} \mathbb{P}(S_t(\omega), b) h_{t+1}(b)\}.$$

Definimos entonces

$$h_t(s) = \max\{f_t(s), \sum_{b \in E} \mathbb{P}(b, s) h_{t+1}(b)\}.$$

Esto es una función de E en $\mathbb{R}$ que satisface $U_t = h_t(X_t)$ y esto concluye nuestro argumento de inducción. $\square$

Desarrollando el último término de nuestra expresión para $h_t(s)$ obtenemos:

$$h_t(s) = \max\{f_t(s), p \cdot h_{t+1}(s(1+u)) + (1-p) \cdot h_{t+1}(s(1+d))\}.$$

Esencialmente, el resultado anterior nos dice que U_t depende únicamente de S_t y no de su historia previa, ya que los dos valores posibles para U_t dependen de h_{t+1} y de f_t .

Capítulo 5

Reclamos contingentes

El problema de la valuación de derivados financieros consiste en fijar un precio justo para un derivado financiero. La dificultad principal consiste en que en el tiempo $t = t_0$ no tenemos información sobre el precio final del activo. Supondremos que el precio final del activo es una variable aleatoria y que por tanto sólo conocemos el conjunto de valores finales del activo, lo que nos permite conocer el conjunto de valores finales del derivado. El conocimiento de este conjunto de valores junto con el principio de no arbitraje son la clave para la valuación de derivados financieros.

El que un mercado esté libre de arbitraje se refiere a que no haya oportunidad de enriquecerse sin hacer nada. Por ejemplo, pensemos que un dólar y un euro cuestan lo mismo tanto en Nueva York como en París. ¿Qué pasaría si en París un dólar costara 0.9 euros, mientras que en Nueva York un dólar y un euro siguieran costando lo mismo? Un inversionista compraría 10 dólares en París, gastando 9 euros en ello y luego puede transferir sus 10 dólares a Nueva York y cambiarlos al momento por 10 euros. Sin hacer prácticamente nada, nuestro especulador acaba de ganarse un euro. Entonces lo que intentaría hacer es mover de la misma manera cantidades de dinero cada vez más grandes, para poder enriquecerse. Esto desequilibraría el mercado y eventualmente movería los precios del dólar y del euro. Las oportunidades de arbitraje desaparecen en cuanto ambos precios se igualan. Un mercado con arbitraje es entonces un mercado en donde pueden generarse ganancias sin inversión neta y sin riesgo. El principio de no arbitraje, que mencionaremos más adelante, dice esencialmente que un modelo matemático de un mercado financiero no debe admitir oportunidades de arbitraje.

5.1. Portafolios y estrategias de mercado

Supondremos que tenemos dos activos, uno de ellos que lo consideraremos sin riesgo (un activo cuyo precio puede subir a través del tiempo, pero nunca bajar, además, conocemos esta tasa de crecimiento) y otro de ellos será considerado con riesgo (un activo cuyo precio está sujeto a cambios aleatorios a través del tiempo). Un ejemplo del primer tipo de activos son los bonos del tesoro. También supondremos que lo que cambia a través del tiempo no es la cantidad física del activo (sin riesgo o con riesgo) que estamos tratando, sino que únicamente cambia el precio (los cambios en la cantidad física del activo sólo se referirán a que nosotros decidamos comprar o vender). Por otra parte, también supondremos que el mercado es infinitamente divisible (tiene sentido hablar de cantidades como π), que no hay comisiones, que las tasas de préstamos que nos hacen y préstamos que hacemos son iguales, todas las transacciones toman lugar inmediatamente y sin retrasos externos. Además, también supondremos que no hay oportunidades de arbitraje.

Consideraremos dos activos, uno sin riesgo y otro con riesgo, a saber, α_1 y α_2 . Empezaremos con introducir algunas definiciones.

Definición 5.1.1. *Un **portafolio** para el intervalo de tiempo $[t_{i-1}, t_i)$ es un par ordenado*

$$\Theta_i = (\Theta_{i,1}, \Theta_{i,2})$$

donde $\Theta_{i,j}(\omega_k)$ es la cantidad del activo α_j adquirido en el tiempo t_{i-1} y que se conservó durante el intervalo de tiempo $[t_{i-1}, t_i)$ conociendo que el estado del mercado es ω_k . Más aún, pedimos que $\Theta_{i,j}$ sea $\mathcal{P}_{i-1}$ -medible. Esto corresponde al hecho obvio de que las cantidades $\Theta_{i,j}$ deben ser conocidas al tiempo t_{i-1} en el cual se adquieren los activos.

Un portafolio es entonces una herramienta que nos permite conocer las cantidades de cada activo que se van comprando o vendiendo a través del tiempo.

Pensemos en un portafolio de la siguiente manera:

- En el tiempo cero, se adquiere el portafolio Θ_1 a los precios a ese tiempo.
- En el tiempo $t = 1$, se liquida el portafolio Θ_1 y se adquiere un nuevo portafolio. Ambas transacciones se hacen simultáneamente y a los precios correspondientes de compra y venta para $t = 1$.

Es necesario que tengamos presente que el portafolio Θ_i es adquirido en el tiempo t_{i-1} y se mantiene hasta el tiempo t_i .

Al tiempo t_0 , el inversionista adquiere su primer portafolio Θ_1 . Al tiempo t_1 , debe adquirir un nuevo portafolio Θ_2 , que bien podría ser el mismo, pero en general supondremos que es diferente.

En general, al tiempo t_{i-1} el portafolio Θ_{i-1} es liquidado y se adquiere un nuevo portafolio Θ_i . Este proceso se conoce como **ajuste del portafolio**.

Definición 5.1.2. Una sucesión de portafolios $\Phi = (\Theta_1, \Theta_2, \dots, \Theta_T)$ se llama **estrategia de mercado**.

También consideraremos las variables aleatorias S_i y A_i , donde S_i es el valor del activo con riesgo en el tiempo t_i y A_i el valor del activo sin riesgo en el tiempo t_i . Se acostumbra tomar los valores descontados de estos dos procesos respecto al valor del activo sin riesgo. Por simplicidad, supondremos que la tasa de interés del activo sin riesgo es cero, a fin de mostrar la idea general sin complicar innecesariamente la notación.

Definición 5.1.3. El **valor de adquisición** de un portafolio Θ_i está dado por

$$\mathcal{V}_{i-1}(\Theta_i) = \Theta_{i,1}A_{i-1} + \Theta_{i,2}S_{i-1}.$$

El **valor de liquidación** de un portafolio Θ_i está dado por

$$\mathcal{V}_i(\Theta_i) = \Theta_{i,1}A_i + \Theta_{i,2}S_i.$$

Definición 5.1.4. Decimos que una estrategia de mercado $\Phi = (\Theta_1, \dots, \Theta_T)$ es **autofinanciada** si $\mathcal{V}_i(\Theta_{i+1}) = \mathcal{V}_i(\Theta_i)$, para cada $i \neq 0$.

Una estrategia de mercado autofinanciada es aquella que nos permite adquirir un nuevo portafolio con lo que obtengamos de liquidar el anterior, sin necesidad de hacer que entre o salga dinero durante este proceso.

Observación 5.1.5. La igualdad $\mathcal{V}_i(\Theta_{i+1}) = \mathcal{V}_i(\Theta_i)$ es obviamente equivalente a:

$$\mathcal{V}_i(\Theta_{i+1}) - \mathcal{V}_i(\Theta_i) = \Theta_{i+1,1}((A_{i+1} + S_{i+1}) - (A_i + S_i))$$

esto es, la pérdida o ganancia obtenidas al seguir una estrategia autofinanciante se deben solamente a los cambios en el precio.

5.2. Mercados completos y valuación de opciones

En el contexto de valuación de derivados financieros, se acostumbra llamar **alternativa** o **reclamo contingente** a las variables aleatorias que representan los valores finales posibles del precio del derivado financiero.

El problema de valuación de derivados está ligado al problema de valuación de alternativas. Una manera simple de valorar una alternativa X es construyendo una estrategia autofinanciante Φ tal que el valor final de dicha estrategia coincida con X . Esto motiva la siguiente definición.

Definición 5.2.1. Sea $X : \Omega \rightarrow \mathbb{R}$ una alternativa. Una **estrategia replicante** para X es una estrategia autofinanciante $\Phi = (\Theta_1, \dots, \Theta_T)$ cuyo pago es igual a X , esto es,

$$\mathcal{V}_T(\Phi) := \mathcal{V}_T(\Theta_T) = X.$$

Una alternativa X que tiene al menos una estrategia replicante se dice **accesible**. Un modelo es **completo** si toda alternativa es accesible.

La estrategia replicante Φ se construye utilizando las martingalas que aparecen en la sección 3.2: la martingala (ξ_k) que definimos en dicha sección corresponde al proceso de precios del activo y la martingala (η_k) se construye partiendo de que en el tiempo T los valores finales del activo están dados por X . En el lema 3.2.1 demostramos que era posible ir construyendo esta estrategia replicante Φ pues en el tiempo j ya disponemos de la información suficiente para ver cómo actuar en el tiempo $j + 1$.

Observación 5.2.2. El conjunto de alternativas accesibles es un espacio vectorial sobre $\mathbb{R}$.

Ejemplo 5.2.3. En este ejemplo calcularemos la estrategia replicante para una alternativa accesible. No es difícil hacer esto, pero involucra resolver muchos sistemas de ecuaciones lineales. Observemos la figura (??). En nuestro caso, tenemos que

$$\omega_1 = (1, 1), \quad \omega_2 = (1, 0), \quad \omega_3 = (0, 1), \quad \omega_4 = (0, 0)$$

$$\text{y } \Omega = \{\omega_1, \omega_2, \omega_3, \omega_4\}$$

donde 1 representa un aumento en el precio y 0 una disminución. Consideraremos una acción cuyo precio inicial es 80. Sabemos que su valor después de un cierto tiempo puede ser 85 o 78. En este momento, no sabemos exactamente cual de los eventos de Ω es el que sucederá, es decir, si la acción subió de precio, nosotros no podemos hacer una distinción

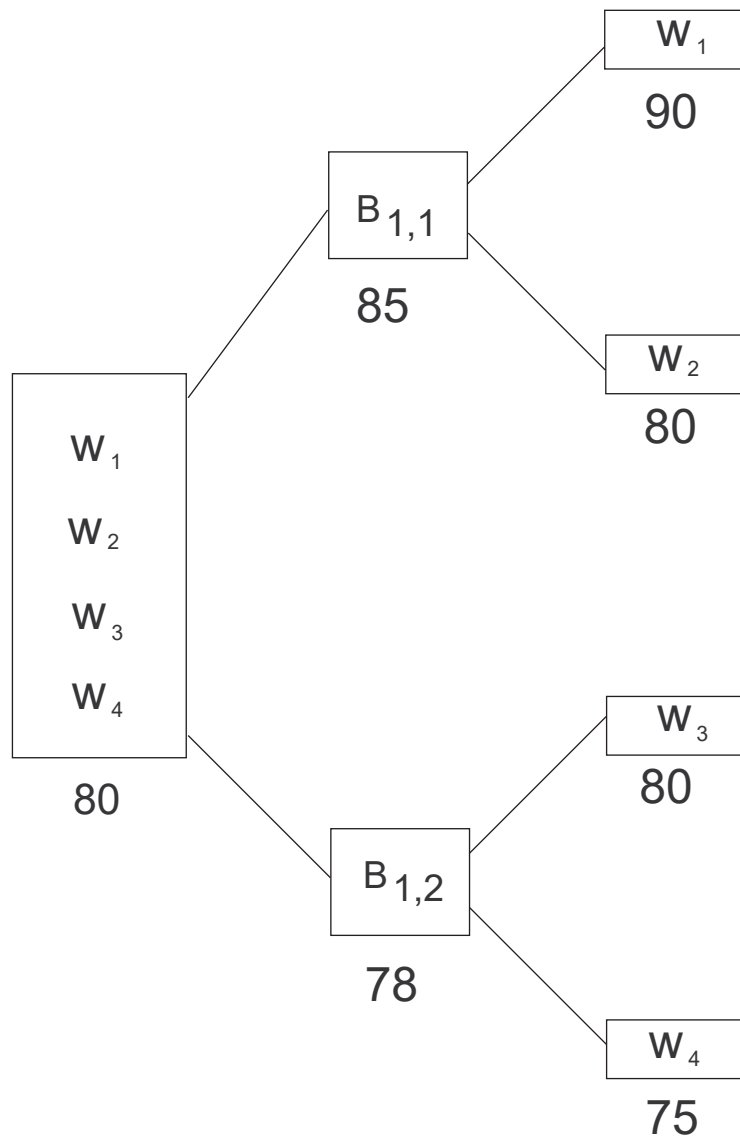


Figura 5.1: Diagrama o árbol de estado

entre ω_1 y ω_2 , aunque ya sabemos que el precio final tendrá que ser 80 o 90. Tenemos entonces que $B_{1,1} = \{\omega_1, \omega_2\}$ y también que $B_{1,2} = \{\omega_3, \omega_4\}$. Representamos $B_{1,1}$ y $B_{1,2}$ como un solo elemento, para hacer notar el hecho de que en ese momento no podemos distinguir entre los dos eventos que corresponden a un alza en el precio de la acción ni entre los dos eventos que corresponden a una baja entre el precio de la acción. Al momento cero no podemos distinguir entre nada, por lo que $B_{0,1} = \Omega$. También supondremos que $A_1 = A_2 = 1$.

Calcularemos una estrategia autofinanciante $\Phi = (\Theta_1, \Theta_2)$ que replique la alternativa

$$X(\omega_1) = 100, X(\omega_2) = 90, X(\omega_3) = 80, X(\omega_4) = 70$$

esto es, para la cual $\mathcal{V}(\Theta_2) = X$, o equivalentemente:

$$\begin{aligned}\mathcal{V}_2(\Theta_2)(\omega_1) &= 100 \\ \mathcal{V}_2(\Theta_2)(\omega_2) &= 90 \\ \mathcal{V}_2(\Theta_2)(\omega_3) &= 80 \\ \mathcal{V}_2(\Theta_2)(\omega_4) &= 70\end{aligned}$$

desarrollando estas ecuaciones nos queda

$$\begin{aligned}A_2(\omega_1)\Theta_{2,1}(\omega_1) + S_2(\omega_1)\Theta_{2,2}(\omega_1) &= 100 \\ A_2(\omega_2)\Theta_{2,1}(\omega_2) + S_2(\omega_2)\Theta_{2,2}(\omega_2) &= 90 \\ A_2(\omega_3)\Theta_{2,1}(\omega_3) + S_2(\omega_3)\Theta_{2,2}(\omega_3) &= 80 \\ A_2(\omega_4)\Theta_{2,1}(\omega_4) + S_2(\omega_4)\Theta_{2,2}(\omega_4) &= 70\end{aligned}$$

sustituyendo los precios al tiempo final nos queda

$$\begin{aligned}\Theta_{2,1}(\omega_1) + 90\Theta_{2,2}(\omega_1) &= 100 \\ \Theta_{2,1}(\omega_2) + 80\Theta_{2,2}(\omega_2) &= 90 \\ \Theta_{2,1}(\omega_3) + 80\Theta_{2,2}(\omega_3) &= 80 \\ \Theta_{2,1}(\omega_4) + 75\Theta_{2,2}(\omega_4) &= 70\end{aligned}$$

por otro lado, como mencionábamos al principio, en el primer momento no podemos hacer distinción entre ω_1 y ω_2 , ni entre ω_3 y ω_4 , por lo que tanto $\Theta_{2,1}$ como $\Theta_{2,2}$ deben ser constantes en los bloques que $B_{1,1}$ y $B_{1,2}$ de donde obtenemos

$$\begin{aligned}\Theta_{2,1}(\omega_1) &= \Theta_{2,1}(\omega_2) \\ \Theta_{2,1}(\omega_3) &= \Theta_{2,1}(\omega_4) \\ \Theta_{2,2}(\omega_1) &= \Theta_{2,2}(\omega_2) \\ \Theta_{2,2}(\omega_3) &= \Theta_{2,2}(\omega_4)\end{aligned}$$

así que nuestro sistema de ecuaciones puede escribirse en términos únicamente de ω_1 y ω_3

$$\begin{aligned}\Theta_{2,1}(\omega_1) + 90\Theta_{2,2}(\omega_1) &= 100 \\ \Theta_{2,1}(\omega_1) + 80\Theta_{2,2}(\omega_1) &= 90 \\ \Theta_{2,1}(\omega_3) + 80\Theta_{2,2}(\omega_3) &= 80 \\ \Theta_{2,1}(\omega_3) + 75\Theta_{2,2}(\omega_3) &= 70\end{aligned}$$

las primeras dos ecuaciones tiene una única solución y las dos segundas también, de donde obtenemos

$$\begin{aligned}\Theta_2(\omega_1) &= \Theta_2(\omega_2) = (10, 1) \\ \Theta_2(\omega_3) &= \Theta_2(\omega_4) = (-80, 2)\end{aligned}$$

ahora calcularemos los valores de adquisición para Θ_2

$$\begin{aligned}\mathcal{V}_1(\Theta_2)(\omega_1) &= 10 + 85 = 95 \\ \mathcal{V}_1(\Theta_2)(\omega_3) &= -80 + 78 \cdot 2 = 76\end{aligned}$$

la condición de que la estrategia sea autofinanciante requiere que estos valores también sean los valores de liquidación de Θ_1

$$\begin{aligned}\mathcal{V}_1(\Theta_1)(\omega_1) &= 95 \\ \mathcal{V}_1(\Theta_1)(\omega_3) &= 76\end{aligned}$$

desarrollando estas ecuaciones y sustituyendo los precios dados

$$\begin{aligned}\Theta_{1,1}(\omega_1) + 85\Theta_{1,2}(\omega_1) &= 95 \\ \Theta_{1,1}(\omega_3) + 78\Theta_{1,2}(\omega_3) &= 76\end{aligned}$$

pero Θ_1 es P_0 -medible, así que para toda $\omega \in \Omega$

$$\begin{aligned}\Theta_{1,1}(\omega) + 85\Theta_{1,2}(\omega) &= 95 \\ \Theta_{1,1}(\omega) + 78\Theta_{1,2}(\omega) &= 76\end{aligned}$$

el sistema tiene solución

$$\Theta_1(\omega) = \left(\frac{-950}{7}, \frac{19}{7} \right)$$

este es un portafolio consistente en vender $\frac{950}{7} \approx 135,71$ bonos del Tesoro y comprar $\frac{19}{7} \approx 2,71$ bienes del activo, por un costo inicial de

$$\frac{-950}{7} + 80 \cdot \frac{19}{7} = \frac{570}{7}$$

esto es, por un costo de $\frac{570}{7} \approx 81,43$ podemos adquirir un portafolio que nos garantiza que podemos pagar lo siguiente:

$$X(\omega_1) = 100, X(\omega_2) = 90, X(\omega_3) = 80, X(\omega_4) = 70$$

observemos que en algunos casos tendremos una ganancia, en otros una pérdida. Esto es lo que uno esperarí en un modelo sin arbitraje.

Es claro que el procedimiento anterior solamente nos sirve para valuar alternativas accesibles. De cualquier manera, todavía tenemos un pequeño problema. Podría ser que haya distintas estrategias replicantes para una alternativa dada pero que tengan diferentes valores iniciales. La solución es requerir la ley del precio único.

Teorema 5.2.4. *(Ley del precio único) Son equivalentes:*

1) Para cualesquiera dos estrategias Φ_1, Φ_2

$$\mathcal{V}_T(\Phi_1) = \mathcal{V}_T(\Phi_2) \Rightarrow \mathcal{V}_0(\Phi_1) = \mathcal{V}_0(\Phi_2).$$

2) Para toda estrategia Φ

$$\mathcal{V}_T(\Phi) = 0 \Rightarrow \mathcal{V}_0(\Phi) = 0.$$

Demostración. Es claro que 1) implica 2) ya que la estrategia $\Phi = 0$ cubre al reclamo 0 y su valor inicial es 0.

Supongamos que tenemos dos estrategias, Φ y Ψ tales que $\mathcal{V}_T(\Phi) = \mathcal{V}_T(\Psi)$ pero $\mathcal{V}_0(\Phi) \neq \mathcal{V}_0(\Psi)$. Entonces tendríamos que la estrategia $\Phi - \Psi$ cubre al reclamo cero pero su valor inicial es distinto de cero, con lo que llegamos a una contradicción. $\square$

La ley del único precio nos asegura que el siguiente funcional de valuación inicial está bien definido.

Definición 5.2.5. *El funcional de valuación inicial $\mathcal{I} : \mathcal{M} \rightarrow \mathbb{R}$ está definido sobre el espacio vectorial $\mathcal{M}$ de las alternativas accesibles y está dado por:*

$$\mathcal{I}(X) = \mathcal{V}_0(\Phi)$$

para cualquier Φ estrategia replicante de la alternativa X .

5.3. Opciones americanas

Como una opción americana puede ser ejercida en cualquier momento antes del tiempo de maduración, pensaremos en una inducción al revés. Sabemos que S_T son los valores finales posibles de la acción al tiempo de maduración T , mientras que Y_T son las ganancias posibles de ejercer dicha opción. ¿Qué valor debería tener la opción en el tiempo $T - 1$? Si el poseedor de la opción la ejerce en el tiempo $T - 1$ entonces debemos pagarle Y_{T-1} , si decide ejercer en el tiempo k debemos entonces pagarle Y_k . Por tanto, nosotros debemos de tener el máximo entre el valor de Y_{T-1} y el valor al tiempo $T - 1$ de una estrategia admisible que pague Y_T al tiempo T , es decir, $\mathcal{E}(Y_T|P_{T-1})$.

Entonces, tiene sentido valuar la opción al tiempo $T - 1$ como

$$U_{T-1} = \text{máx}\{Y_{T-1}, \mathcal{E}(Y_T|P_{T-1})\}$$

definimos entonces recursivamente el valor de la opción americana al tiempo k como

$$U_{k-1} = \text{máx}\{Y_{k-1}, \mathcal{E}(U_k|P_{k-1})\}$$

que es justamente la envolvente de Snell.

Esto nos da un valor inicial justo para la opción americana.

Definición 5.3.1. *El valor inicial justo para una opción americana es U_0 .*

En la sección 3.2, pudimos construir una estrategia replicante Φ gracias a la existencia de las martingalas (ξ_k) y (η_k) .

Ahora escribamos $U_k = M_k - T_k$ de acuerdo a la descomposición de Doob. Consideremos la estrategia Φ que replica a la alternativa M_T . Entonces, la martingala (M_k) corresponde a la martingala (η_k) de la sección 3.2, mientras que (ξ_k) es la martingala que corresponde al proceso de precios. Como la sucesión $(\mathcal{V}_k(\Phi))$ es una martingala, tenemos entonces que $\mathcal{V}_k(\Phi) = M_k$ y por tanto,

$$U_k = \mathcal{V}_k(\Phi) - T_k$$

a partir de esta última igualdad, observamos que el vendedor de la opción puede cubrirse perfectamente; una vez que recibió la cantidad $U_0 = \mathcal{V}_0(\Phi)$ puede generar una cantidad $\mathcal{V}_k(\Phi)$ al tiempo k , que es mayor que U_k y que Y_k . Si el vendedor de la opción se cubre utilizando la estrategia que hemos definido, asegura obtener una ganancia igual a $\mathcal{V}_\tau(\Phi) - Y_\tau = U_\tau + T_\tau - Y_\tau$ que es positiva cuando el tiempo τ no es un tiempo óptimo.

¿Cuál es la moraleja de esta tesis? Arriesga si tu ganancia esperada es mejor que tu ganancia actual. De lo contrario, abstente.

Ejemplo 5.3.2. *Una cita promedio con una mujer tiene un costo de C pesos. Sea p la probabilidad de que acepte ir a casa por un café y $1 - p$ la probabilidad de que le duela la cabeza. Sea V el valor de salir con esa mujer, convertido en pesos. Supongamos (sin pérdida de generalidad) que el valor que su dolor de cabeza es 0. La condición para invitarla a salir es entonces $pV > C$.*

Bibliografía

- Balanzario, E. 'El modelo CRR para la fórmula de Black y Scholes', apuntes, Instituto de Matemáticas UNAM-Morelia, 2005.
- Cox, J., Ross, S., Rubinstein, M. 'Option pricing: a simplified approach', *Journal of Financial Economics* , 1979.
- Delbaen, F., Schachermayer, W. 'What is a Free lunch?', *Notices of the AMS* Vol. 51 No.5, pp. 526-528 , 2004.
- Hull, John C. 'Options, futures, and other derivatives', Pearson-Prentice Hall New Jersey, 2006.
- Koch-Medina, P., Merino, S. 'A discrete introduction: Mathematical finance and probability' Birkhäuser-Verlag, Basel, Suiza, 2003.
- Lamberton, D., Lapeyre, B. 'Introduction to Stochastic Calculus applied to finance', Chapman & Hall, Londres, 1996.
- Price, J. 'Optional Mathematics is not optional' *Notices of the AMS* Vol. 43 No. 9 pp.964-971, 1996.
- Roman, S. 'Introduction to the mathematics of finance: from risk management to options pricing', Springer-Verlag, 2004.