

Appunti del corso
Probabilità e Finanza
A.A. 2006/2007

PAOLO BALDI
LUCIA CARAMELLINO

<http://www.mat.uniroma2.it/~caramell>
[versione: maggio 2007]

Ringrazio tutti gli studenti che hanno contribuito, con le loro segnalazioni, a diminuire drasticamente il numero di errori (di battitura e non) presenti in queste note. Tra tutti, vorrei citare Enzo Ferrazzano, la cui scrupolosa (e un po' avvilita...) errata corrige si è rivelata particolarmente utile.

L. C.
(Maggio 2007)

Indice

1	Introduzione alla finanza	1
1.1	Tassi d'interesse	1
1.2	Alcuni aspetti dei mercati finanziari	3
1.2.1	Vendita allo scoperto	3
1.2.2	Arbitraggio	4
1.2.3	Le proprietà del mercato finanziario	5
1.3	Prodotti finanziari, prodotti derivati	6
1.3.1	Contratti forward e futures	6
1.3.2	Opzioni	8
1.3.3	A cosa servono i derivati?	10
1.4	Soluzioni	13
2	Probabilità: richiami e complementi	15
2.1	Cenni di teoria della misura	15
2.1.1	σ -algebre	15
2.1.2	Spazi e funzioni misurabili	19
2.2	Spazi di probabilità	24
2.2.1	Prime proprietà	25
2.2.2	Indipendenza	27
2.3	Variabili aleatorie	28
2.3.1	Proprietà	29
2.3.2	Legge e distribuzione	30
2.3.3	Speranza matematica	34
2.4	Soluzioni	37
3	Condizionamento e martingale	41
3.1	Condizionamento	41
3.1.1	Media condizionale: il caso discreto	41
3.1.2	Il caso generale	44
3.1.3	Proprietà della speranza condizionale	46
3.1.4	Probabilità condizionale	49
3.2	Martingale	51
3.2.1	Definizioni e generalità	51

3.2.2	Prime proprietà	52
3.2.3	La decomposizione di Doob	53
3.2.4	Martingale trasformate	53
3.3	Tempi d'arresto	55
3.3.1	Il teorema d'arresto	57
3.4	Soluzioni	59
4	Modelli probabilistici per la finanza	62
4.1	Introduzione	62
4.2	Modelli discreti per la descrizione dei mercati finanziari . . .	65
4.3	Strategie ed arbitraggio	66
4.4	Arbitraggio e martingale	70
4.5	Mercati completi e prezzo delle opzioni europee	77
4.6	Il modello di Cox, Ross e Rubinstein	85
4.6.1	Assenza di arbitraggio e completezza	86
4.6.2	Prezzo e copertura delle opzioni nel modello CRR . . .	90
4.6.3	Passaggio al limite e la formula di Black e Scholes . . .	93
4.6.4	Studio empirico della velocità di convergenza	98
4.7	La volatilità	100
4.8	Le greche	101
4.9	Appendice al Capitolo 4	103
4.9.1	Teorema di separazione dei convessi	103
4.9.2	Un codice C per la funzione di ripartizione normale standard	104
5	Opzioni americane	105
5.1	Il modello	105
5.2	Il problema dell'arresto ottimo	106
5.3	Prezzo e copertura delle opzioni americane	110
5.4	La put americana nel modello CRR	113
6	Simulazione e metodi Monte Carlo per la finanza	118
6.1	Metodi Monte Carlo: generalità	118
6.2	Simulazione del modello CRR	120
6.3	Prezzo di opzioni europee con metodi Monte Carlo	122
6.3.1	Call e put standard	123
6.3.2	Opzioni call/put asiatiche	124
6.3.3	Opzioni call/put con barriere	125
6.4	Copertura dinamica di opzioni europee	125
6.5	Put americana	127
6.5.1	Prezzo	128
6.5.2	Copertura dinamica	129
	Bibliografia	131

Capitolo 1

Introduzione alla finanza

In questi primi paragrafi vedremo alcune nozioni elementari e introduttive. Da un punto di vista matematico si tratta di rivedere cose già note sulle proporzioni, più un po' di terminologia finanziaria.

1.1 Tassi d'interesse

Supponiamo di aprire un conto corrente versando un capitale pari a x . Indichiamo con $r \cdot 100\%$ il tasso d'interesse annuo che viene corrisposto. Quindi se, ad esempio $r = 0.05$, ciò corrisponde, nel linguaggio corrente ad un interesse annuo del $0.05 \cdot 100\% = 5\%$.

Alla fine dell'anno il valore del capitale sarà pari a

$$x + xr = x(1 + r)$$

Si dice che viene corrisposto un *interesse semplice* se negli anni successivi l'interesse viene calcolato sempre sul capitale iniziale x . Nella realtà è quello che succede se alla fine di ogni anno ritiriamo la porzione di capitale maturata, cioè xr . Dopo n anni, il capitale investito ad un interesse semplice pari a r avrà prodotto un capitale finale pari a

$$x + \underbrace{xr + \cdots + xr}_{n \text{ volte}} = x(1 + nr)$$

Si parla invece di *interesse composto* quando alla fine di ogni anno l'interesse viene calcolato sul capitale fino ad allora maturato. Dunque il capitale sarà pari a

$x(1 + r)$	alla fine del primo anno
$x(1 + r)^2$	alla fine del secondo anno
$\vdots$	$\vdots$
$x(1 + r)^n$	alla fine dell' n -esimo anno

È quindi evidente la differenza tra interesse semplice (il capitale cresce linearmente) e composto (il capitale cresce esponenzialmente). Nella realtà, vengono applicati interessi composti (e quasi mai semplici). Ovviamente si può procedere anche al contrario: se una quantità vale y all'anno n e se è applicato un tasso (composto) annuale r , il valore ad oggi, detto *valore attualizzato*, è

$$x = y(1 + r)^{-n}.$$

Esercizio 1.1. *Supponiamo di possedere un titolo “perpetuo”, che paga c Euro alla fine di ogni anno (cioè, riscuoteremo c Euro alla fine dell'anno n , per ogni $n = 1, 2, \dots$). Supponendo che il tasso di interesse composto sia r , qual è il valore attualizzato di questo titolo?*

Talvolta una banca versa l'interesse maturato alla fine di un periodo più piccolo di un anno. Se l'interesse annuo resta del $r \cdot 100\%$, allora se nell'anno sono previsti k periodi, il capitale sarà

$$\begin{array}{ll} x(1 + \frac{r}{k}) & \text{alla fine del primo periodo} \\ x(1 + \frac{r}{k})^2 & \text{alla fine del secondo periodo} \\ \vdots & \vdots \\ x(1 + \frac{r}{k})^k & \text{alla fine dell'anno (k -esimo periodo)} \end{array}$$

e su n anni si ottiene la seguente ricapitalizzazione:

$$x \left(1 + \frac{r}{k}\right)^{nk}.$$

Se facciamo tendere il numero di periodi nei quali l'anno è suddiviso all'infinito (cioè, $k \rightarrow \infty$), evidentemente otteniamo che il capitale alla fine dell' n -esimo anno sarà uguale a

$$\lim_{k \rightarrow \infty} x \left(1 + \frac{r}{k}\right)^{nk} = x e^{r n}$$

In questo caso, r prende il nome di *tasso istantaneo di interesse*. Ricordando che il limite $x(1 + \frac{r}{k})^k \uparrow e^r$ è crescente (qui infatti $r > 0$), si ha che $x e^{r n}$ è l'estremo superiore dei capitali che si possono ottenere su n anni con un interesse composto pari a r quando si suddivide l'anno in più periodi. Inoltre più sono i periodi, maggiore è il guadagno.

Esercizio 1.2. *La banca A offre, per il vostro conto corrente, un interesse del 3% semestrale; la banca B offre un interesse del 2% trimestrale. Qual è l'interesse annuale? A quale banca affidereste i vostri risparmi per sei mesi? E per un anno?*

Esercizio 1.3. Supponiamo di aver bisogno di 100mila Euro, ad esempio per acquistare una casa, e per questo chiediamo un mutuo ad una banca. La banca offre un mutuo di 15 anni, ad un interesse mensile dello 0.6%, da ripagare con rate mensili. Qual è il valore di ciascuna rata? In più, la banca chiede una tassa di 600 Euro per l'apertura del mutuo, un'altra tassa di 400 Euro per ispezionare la casa ed infine l'1% dell'ammontare di denaro richiesto. A conti fatti, qual è l'effettivo tasso di interesse applicato?

Esercizio 1.4. Un individuo ritiene che andrà in pensione tra 20 anni e decide di mettere in banca un certo ammontare di denaro alla fine di ogni mese, diciamo A , per i prossimi $12 \times 20 = 240$ mesi. Egli vuol essere sicuro che l'investimento gli garantisca una rendita di 1000 Euro mensili, per i successivi 30 anni. Assumendo che la banca gli dia un tasso di interesse del 6% annuale, composto mensilmente, qual è il valore di A ?

Esercizio 1.5. Supponiamo di poter depositare del denaro ad un interesse istantaneo del 10% annuo. Quanto occorre aspettare per raddoppiare il patrimonio?

Esercizio 1.6. Una banca propone due investimenti, che paga alla fine dell'anno i , per $i = 1, 2, 3$. Il primo, garantisce le quote 100, 140, 131 Euro; il secondo, 90, 160, 120 Euro. È possibile dire quale dei due investimenti sia migliore pur non conoscendo il tasso di interesse?

1.2 Alcuni aspetti dei mercati finanziari

Vediamo di descrivere alcuni aspetti dei mercati finanziari. Non andremo troppo in profondità, ma quando si modellizzano dei problemi reali, occorre averne un minimo di conoscenza. Vedremo poi che alcune nozioni tipiche dei mercati finanziari hanno una controparte interessante in matematica.

1.2.1 Vendita allo scoperto

Una cosa che può succedere in un mercato finanziario, è che qualcuno venda qualcosa che non possiede. Ciò è possibile in due modi diversi. Ad esempio giocando sul fatto che talvolta gli acquisti si regolano alla fine della giornata. Nel film “Un posto per due” i protagonisti sanno che il succo concentrato di frutta calerà di prezzo durante la giornata. Quindi, al mattino entrano nel mercato vendendo a prezzo elevato (non possiedono neanche una goccia...), mentre al pomeriggio lo acquistano a prezzo basso. Alla chiusura compensano quello venduto con quello acquistato e si intascano la differenza. Un'altra possibilità, per un investitore, consiste nel rivolgersi ad un broker (cioè un agente di cambio) che prende “in prestito” dal portafoglio di un altro cliente un pacchetto di titoli. Quando l'investitore avrà effettuato

l'operazione che gli interessa restituirà il pacchetto, che verrà “rimesso al suo posto”.

In gergo finanziario, colui che acquista il titolo prende una *posizione lunga*, mentre chi lo vende prende una *posizione corta*.

Come si può immaginare si tratta sempre di operazioni un po' rischiose e la vendita allo scoperto non è sempre una operazione consentita, oppure è comunque regolamentata. Noi considereremo un mercato finanziario in cui essa è permessa. Questo sarà anzi un elemento importante, che implicherà delle considerazioni teoriche importanti.

1.2.2 Arbitraggio

Per introdurre il concetto (fondamentale) di arbitraggio, cominciamo con un esempio.

Supponiamo che l'11 febbraio 2004 si osservino sul mercato dei cambi i tassi seguenti

Euro/USD	1.267
Euro/Yen	133.85
USD/Yen	106.75

Uno speculatore potrebbe fare l'operazione seguente: con 1 USD compra 106.75 Yen, che cambia in $106.75/133.85 = 0.798$ Euro e poi in $0.798 \cdot 1.267 = 1.011$ USD. Avrebbe quindi guadagnato 1.1 centesimi di USD. Anzi avrebbe potuto vendere allo scoperto il primo dollaro. Avrebbe quindi avuto un guadagno certo senza neanche avere bisogno di un capitale iniziale, cioè avrebbe fatto un arbitraggio.

In generale, si chiama *arbitraggio* un'operazione finanziaria che

- non necessita di un capitale iniziale;
- non può in nessun modo dare luogo ad una perdita e, con probabilità strettamente positiva, dà luogo ad un guadagno.

Nella realtà situazioni di arbitraggio si presentano, ma naturalmente tendono a riassorbirsi molto velocemente. Nella situazione immaginata precedentemente, inevitabilmente ci sarebbe stato un aumento di richieste di cambi vendite di USD contro Yen e questo avrebbe fatto rapidamente scendere il cambio USD/Yen, facendo sparire la possibilità di arbitraggio. Spesso questo ruolo di riequilibrio è svolto dalle banche centrali

Per questo motivo, in un modello di mercati finanziari si fa spesso l'ipotesi che l'arbitraggio non sia possibile. Questa è una ipotesi chiave per stabilire il prezzo di un bene. Ad esempio, dando per corretti i valori del cambio Euro/Yen e Euro/USD, quanto deve valere il cambio USD/Yen perché non ci sia arbitraggio?

Se indichiamo con x questo valore, occorrerà che facendo l'operazione

$$\text{USD} \rightarrow \text{Yen} \rightarrow \text{Euro} \rightarrow \text{USD}$$

non ci deve essere guadagno. Ora con questa operazione 1USD viene trasformato in

$$1 \text{ USD} = x \text{ Yen} = x \cdot \frac{1}{133.85} \text{ Euro} = x \cdot \frac{1}{133.85} \cdot 1.267 \text{ USD}.$$

Questa quantità è uguale a 1 se e solo se $x = 133.85/1.267 = 105.64$. Questo dunque deve essere il cambio USD/Yen perché non ci sia arbitraggio.

Giusto per curiosità, i valori dei cambi indicati sopra sono quelli reali alle 14h02 dell'11 febbraio 2004, per quanto riguarda i cambi USD/Euro e Euro/Yen. Invece il cambio USD/Yen era pari a 105.55. La lieve differenza con il valore ottenuto con il criterio dell'arbitraggio deriva dal fatto che nelle operazioni di cambio descritte sopra avevamo implicitamente supposto che non ci fossero spese di transazione. Si sa invece che passare da una valuta all'altra comporta delle spese.

1.2.3 Le proprietà del mercato finanziario

Nel resto di questo corso considereremo un mercato finanziario in cui valgono certe proprietà, che non si trovano in realtà nei mercati reali. Ciò si giustifica con la necessaria semplicità che si richiede ad un primo approccio ai problemi. Molte di queste ipotesi si possono oggi eliminare (si veda ad esempio l'assenza di spese di transazione), ma al prezzo di un trattamento matematico molto più elaborato.

Le proprietà che saranno richieste sono le seguenti.

1. Esiste nel mercato un tipo di investimento senza rischio, detto *titolo non rischioso* o *obbligazione* oppure *bond* (del tipo reddito da interesse versato per capitale depositato su un conto corrente), con un tasso, eventualmente istantaneo, costante, che indicheremo r . Questo investimento è calcolato mediante:
 - un interesse composto, nel qual caso al tempo t il valore di un investimento iniziale pari a x è dato da $x_t = x(1+r)^t$, oppure
 - un interesse composto a tempo continuo, dunque una somma pari a x , dopo un tempo t viene rivalutata a $x_t = xe^{rt}$.

Come vedremo, nel seguito useremo entrambi i modelli per l'evoluzione del titolo non rischioso. Per il momento, scegliamo il secondo tipo.

2. È possibile prendere in prestito somme di denaro allo stesso tasso e con le stesse regole.
3. I costi di transazione (spese per i cambi di valuta, acquisto o vendita di azioni o obbligazioni...) sono uguali a 0.
4. È permessa la vendita allo scoperto.

5. Sono permesse operazioni riguardanti frazioni di beni. Ad esempio è possibile acquistare 0.63 azioni di una società.

1.3 Prodotti finanziari, prodotti derivati

Nei mercati finanziari sono presenti molti prodotti (valute, azioni, obbligazioni...). Esistono però anche degli strumenti finanziari il cui valore è determinato da quello di altri prodotti. Questi ultimi si chiamano *prodotti derivati*. In questo paragrafo vedremo due degli esempi principali.

1.3.1 Contratti forward e futures

Esempio 1.3.1. (*Contratti forward*) Supponiamo che una industria americana sappia, oggi 16 gennaio 2004, che dovrà pagare una fattura in Euro all'inizio di marzo 2004. Ciò espone la società ad un rischio, legato alle fluttuazioni dei cambi. Per proteggere le società da questo genere di rischi esistono i *contratti forward*. Con questo contratto una seconda società s'impegna a fornire la merce al tempo indicato, *ad un prezzo fissato oggi*. Il prezzo di un contratto forward con scadenza 4 marzo 2004 è di 1.2462. Comperando questo contratto forward, la società si impegna a pagare alla data del 4 marzo 2004 una prefissata somma di Euro al tasso di cambio di 1.2462. Se il cambio Euro/Dollaro alla data indicata sarà superiore a 1.2462, la società che ha acquistato un contratto forward ci avrà guadagnato, se invece il cambio risulterà più basso, ci avrà rimesso. In ogni caso però si sarà garantita da improvvise oscillazioni dei cambi.

Un *contratto forward* è quindi un impegno sottoscritto da due parti, A e B, in cui B si impegna oggi a consegnare ad A alla data futura T una merce (frumento, valuta, azioni...) ad un prezzo F fissato oggi. Da parte sua, A si impegna ad acquistare da B la merce pattuita e a pagarla F .

Un po' di nomenclatura.

- La data di scadenza T (4 marzo 2004 nell'Esempio 1.3.1) si chiama la *data di maturità*.
- Il prezzo stabilito F (1.2462 nell'esempio precedente) si chiama il *prezzo di consegna*.
- Il prezzo corrente di un contratto forward si chiama il *prezzo forward*. Alla data di emissione prezzo di consegna e prezzo forward coincidono. Supponiamo però, acquistato un contratto forward come nell'Esempio 1.3.1, alla data del 10 febbraio, a seconda che il cambio sia salito o sceso, il prezzo di un contratto forward sarà salito o sceso di conseguenza.

- Il prezzo del bene sottostante il contratto (cioè il valore del cambio, nell'Esempio 1.3.1), si chiama il *prezzo spot*.
- Nel gergo del milieu, si dice che la società A che acquista un contratto forward prende una *posizione lunga*, mentre chi lo vende, B, prende una *posizione corta*, con una terminologia simile a quella che abbiamo visto per le vendite allo scoperto.

Precisiamo un paio di cose.

1. Entrare in un contratto forward non richiede nessun esborso di denaro, ma vincola la società all'acquisto della "merce" pattuita, alla data di maturità ed al prezzo di esercizio pattuito.
2. I contratti forward non sono regolamentati in borsa, ma possono comunque essere oggetti di scambio.

Alla fine del contratto è naturale considerare la quantità

$$S_T - F \quad (1.1)$$

Qui T è la data di maturità, S_T è il prezzo del sottostante a maturità, e F è il prezzo di consegna. Questa quantità rappresenta il guadagno (la perdita, se negativa...) del detentore del contratto forward, rispetto alla strategia ingenua che consisteva nell'attendere il 4 marzo ed acquistare gli Euro sul mercato. Il guadagno/perdita della società che ha preso una posizione corta sarà invece

$$F - S_T \quad (1.2)$$

I contratti futures sono simili ai forward, con la differenza che sono regolamentati dalla borsa, che impone una serie di obblighi, tra cui l'apertura di speciali conti correnti su cui le due controparti devono versare delle somme a garanzia. Tralascieremo di descrivere questi dettagli, concentrandoci sui contratti forward che coinvolgono meno dettagli di tecnica finanziaria e, soprattutto, sulle opzioni.

Vediamo ora come l'ipotesi di assenza di arbitraggio consenta di stabilire il giusto prezzo di consegna di un contratto forward. Consideriamo quindi un bene, di cui indichiamo con S_t il prezzo al tempo t . Qual è il prezzo di consegna per un contratto forward, di maturità T emesso a $t = 0$? Il criterio di assenza di arbitraggio impone che il prezzo di consegna F debba essere uguale a

$$S_0 e^{rT}$$

dove r indica al solito il tasso d'interesse istantaneo di un conto corrente. Vediamo perché.

Infatti, supponiamo $F < S_0 e^{rT}$; allora un investitore potrebbe vendere allo scoperto una unità del bene in questione al prezzo odierno, S_0 , versarlo in

banca e simultaneamente sottoscrivere un contratto forward al prezzo F . A maturità ritira dalla banca il capitale maturato, che vale ora $S_0 e^{rT}$. Può ora onorare il contratto acquistando una unità di sottostante al prezzo F e restituirla per compensare la vendita allo scoperto. Il bilancio dell'operazione è $S_0 e^{rT} - F > 0$. Questa operazione realizza un guadagno strettamente positivo, senza rischio e senza bisogno di disporre di un capitale; si tratta quindi di un arbitraggio.

Viceversa se $F > S_0 e^{rT}$ si può costruire una operazione di arbitraggio in maniera simile. Si vende un contratto forward con prezzo di consegna F e simultaneamente si prende in prestito in banca un ammontare S_0 con il quale si acquista una unità di sottostante. A maturità si cede l'unità di sottostante, incassando la quantità F . Occorre ora restituire alla banca il capitale preso a prestito, che ora vale $S_0 e^{rT}$. Il bilancio dell'operazione è ancora positivo: $F - S_0 e^{rT} > 0$.

1.3.2 Opzioni

Le opzioni sono un tipo di prodotti derivati diverso dai contratti forward e futures. Una opzione è un contratto che dà il diritto (ma non l'obbligo) a chi lo detiene di acquistare (o vendere, a seconda del tipo d'opzione) un bene (il *sottostante*) ad una data prefissata (la data di *maturità*) e ad un prezzo prefissato (il prezzo di *esercizio* o *strike*).

Le opzioni che danno diritto ad acquistare si chiamano *call*, quelle che danno diritto a vendere si chiamano *put*.

In realtà ci sono molti tipi di opzione. Noi ci occuperemo dei due tipi principali.

- Le opzioni *europee*, per le quali si ha il diritto di esercitare l'opzione unicamente al tempo di maturità.
- Le opzioni *americane*, per le quali il detentore ha il diritto di esercitare l'opzione in qualunque istante compreso tra il tempo di emissione ed il tempo di maturità.

Esempio 1.3.2. Il prezzo dell'opzione call sull'Euro al Chicago Board of Trade al 16 gennaio 2004, con maturità $T = 4$ Marzo e prezzo di esercizio $K = 1.26$ era di 0.017USD. Questo significa che l'acquirente, pagando il prezzo di 0.017USD acquisisce il diritto ad acquistare un euro al prezzo di 1.26USD il 4 Marzo. Quindi se il cambio Euro/USD sarà inferiore a $K = 1.26$, il detentore dell'opzione rinuncerà ad esercitare l'opzione ed acquisterà gli Euro che gli servono sul mercato. Avrà quindi pagato inutilmente il prezzo dell'opzione. Se invece il cambio Euro/USD il 4 Marzo sarà superiore al prezzo di esercizio 1.26, eserciterà l'opzione e si procurerà la valuta di cui ha bisogno al prezzo di esercizio. In ogni caso, mediante l'opzione si sarà garantito da un aumento dei cambi.

Indichiamo con S_t il prezzo del sottostante al tempo t . Immaginando di avere a che fare con un'opzione europea, al tempo $t = T$ la società che ha venduto l'opzione deve pagare, nel caso di una call europea, un prezzo pari a

$$C_T = (S_T - K)_+ = \max(S_T - K, 0) \quad (1.3)$$

Infatti, se $S_T < K$, l'opzione non viene esercitata e la società non deve pagare nulla. Se invece $S_T > K$, allora per onorare l'opzione occorre sborsare un ammontare pari a $S_T - K$ per ogni unità di sottostante. La quantità (1.3) si chiama il *payoff* dell'opzione, e rappresenta quindi la quantità di denaro che deve sborsare colui che cede il diritto di opzione.

Per una put europea, si vede subito che il payoff è

$$P_T = (K - S_T)_+ \quad (1.4)$$

Vedremo più tardi come si esprime il payoff delle opzioni americane. Sia per le opzioni put che per le call, il payoff è della forma $\Phi(S_T)$, dove

$$\begin{aligned} \Phi(x) &= (x - K)_+ && \text{per le opzioni call} \\ \Phi(x) &= (K - x)_+ && \text{per le opzioni put} \end{aligned}$$

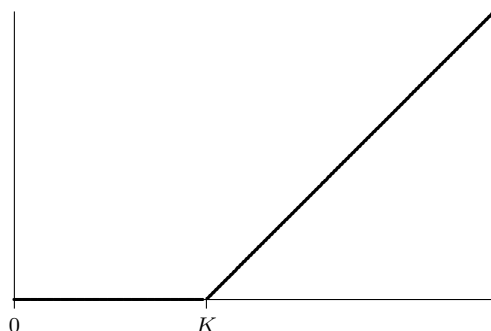


Figura 1.1 Grafico per il payoff di una opzione call.

Nel gergo finanziario:

- chi acquista un'opzione call o una unità di sottostante si dice che prende una posizione *lunga*;
- chi vende una opzione call oppure una unità di sottostante si dice che prende una posizione *corta*.

Le opzioni, come i contratti forward, sono state concepite per ridurre i rischi degli operatori finanziari. Nell'Esempio 1.3.2, lo strumento dell'opzione veniva usato per garantirsi di acquisire la valuta di cui si ha bisogno ad un prezzo controllato (sicuramente inferiore o uguale al prezzo di esercizio). Le opzioni possono però anche essere usate a scopo speculativo. Sempre nel quadro dell'Esempio 1.3.2, un operatore avrebbe potuto considerare una speculazione

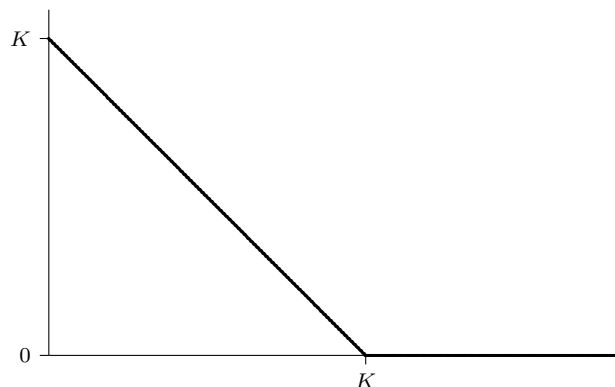


Figura 1.2 Grafico per il payoff di una opzione put. Rispetto all'opzione call c'è una notevole differenza: il payoff di una opzione put è limitato.

in cui acquista delle opzioni, sperando che il prezzo del sottostante a maturità sia superiore a quello d'esercizio, in maniera da lucrare la differenza. In genere si deve considerare che le speculazioni sulle opzioni sono ad alto rischio. Possono cioè produrre grossi guadagni, ma è anche molto elevato il rischio di perdere la totalità della somma investita. Torneremo più tardi su questo punto.

1.3.3 A cosa servono i derivati?

I derivati si possono usare essenzialmente per due scopi diversi:

- per coprire il rischio, oppure
- a scopo speculativo.

Qui “coprire il rischio” non vuole dire nascondere, ma ridurlo. Abbiamo già visto nell'Esempio 1.3.1 questo aspetto: con un contratto forward una compagnia si può garantire di disporre ad un tempo futuro fissato della valuta, o delle materie prime di cui ha bisogno. Oppure, pensando ad una opzione call oppure alla vendita di un contratto forward (oppure entrare corti su un contratto forward, come si dice) una società può garantirsi di vendere un suo prodotto ad un prezzo fissato. Le opzioni forniscono comunque anche uno strumento per gli speculatori, cioè per quegli operatori che cercano investimenti a rischio elevato, che possono cioè produrre grossi guadagni ma anche grosse perdite.

Esempio 1.3.3. Supponiamo che il prezzo di un'azione al 4 gennaio 2003 sia pari a 39 Euro e che il costo di una opzione call con uno strike di 40 Euro e maturità 4 aprile sia di 2 Euro. Uno speculatore che pensa che il prezzo dell'azione aumenterà nel futuro, ha a disposizione due strategie, supponendo un capitale di 3900 Euro:

- può acquistare 100 azioni, oppure
- può acquistare 1950 opzioni call.

Indichiamo con S_T il valore dell'azione alla maturità $T = 4$ aprile e vediamo quanto vale il portafoglio V_T dell'investitore in ciascuna delle due strategie considerate.

- Nel primo caso il valore è semplicemente $V_T = 100 \cdot S_T$.
- Nel secondo, ci sono due possibilità. Se $S_T > 40$, allora egli eserciterà l'opzione, acquisendo le azioni al prezzo di 40 Euro, che poi venderà al prezzo di mercato S_T . Il suo capitale sarà quindi pari a $1950 \cdot (S_T - 40)$. Se invece il prezzo dell'azione S_T sarà sotto i 40 Euro, l'opzione non verrà esercitata e l'investitore avrà perso il capitale investito. In conclusione il valore finale dell'investimento in questo secondo caso è di $V_T = 1950 \cdot (S_T - 40)_+$.

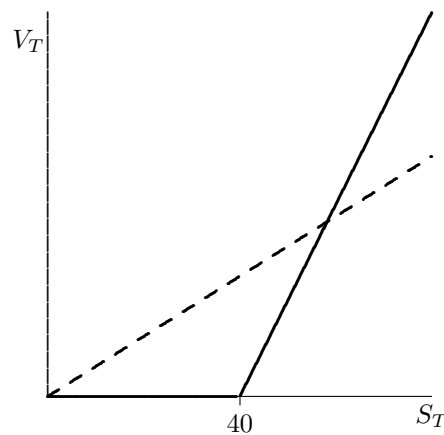


Figura 1.3 Grafico del valore dell'investimento V_T in funzione del prezzo S_T ; in tratteggio, il valore di V_T nella prima strategia.

È evidente che il secondo tipo d'investimento è molto più sensibile alle variazioni del prezzo S_T . Ad esempio, se $S_T = 44$, il valore del capitale per ciascuno dei due investimenti sarà rispettivamente

$$V_T = 4400 \quad \text{e} \quad V_T = 4 \cdot 1950 = 7800$$

Invece se fosse $S_T = 41$, allora avremmo 4100 e 1950. Soprattutto, se fosse $x \leq 40$, il valore del secondo investimento sarebbe uguale a 0. Dunque i derivati, che sono stati introdotti allo scopo di ridurre i rischi, possono anche essere usati a scopo speculativo. Un investitore con un obiettivo di questo

tipo potrà apprezzare le grosse possibilità di guadagno che essi offrono, ma dovrà tenere conto che il rischio di subire grosse perdite è molto elevato. Per quanto riguarda le opzioni, il problema del calcolo del prezzo non è così semplice da affrontare. Scopo principale di queste note e del corso, è appunto il calcolo del prezzo delle opzioni.

1.4 Soluzioni

Soluzione dell'esercizio 1.1. Se alla fine dell' n -esimo anno avremo c , allora il valore ad oggi di questa quantità è $c(1+r)^{-n}$. Dunque, il valore attualizzato di questo titolo è

$$\sum_{n=1}^{\infty} c(1+r)^{-n} = \frac{c}{1+r} \frac{1}{1-1/(1+r)} = \frac{c}{r}.$$

Soluzione dell'esercizio 1.2. Se immaginiamo che in un anno vi siano k rivalutazioni, ciascuna con tasso r_{per} , allora l'interesse annuale è $r_{\text{ann.}} = (1+r_{\text{per}})^k - 1 \simeq k \cdot r_{\text{per}}$ (quest'ultima approssimazione vale perché in genere i tassi sono percentuali non troppo elevate). Dunque, l'interesse annuale è $r_{\text{ann.}} \simeq 3 \times 2\% = 6\%$ nel primo caso (anno diviso in 2 periodi) e $r_{\text{ann.}} \simeq 2 \times 4\% = 8\%$ nel secondo (anno diviso in 4 periodi). Su un solo semestre, il tasso applicato è 3% nel primo caso e $[(1+0.02)^2 - 1] \times 100\% \simeq 2 \times 2\% = 4\%$ nel secondo. Evidentemente, è più vantaggioso il secondo caso.

Soluzione dell'esercizio 1.3. Osserviamo che in totale occorre far fronte a $12 \times 15 = 180$ rate. Sia A il valore di ciascuna rata. Per calcolare A , basta imporre che il valore attualizzato dei soldi che nei 15 anni saranno versati alla banca sia uguale alla somma presa in prestito. Ora, il valore attualizzato dell' n -esima rata è $A \cdot \rho^n$, con $\rho = (1+0.006)^{-1}$, dunque dev'essere

$$\sum_{n=1}^{180} A \rho^n = 100000,$$

da cui segue che

$$A = 100000 \frac{1-\rho}{\rho(1-\rho^{180})} = 910.05$$

Quindi, un prestito di 100mila Euro in 15 anni ad un tasso mensile dello 0.06% costa 910.05 Euro di rata mensile. Dovendo però pagare varie tasse, in realtà il prestito è pari a $100.000 - 600 - 400 - 1000 = 98.000$ Euro. Dunque, detto r_m l'effettivo tasso di interesse mensile, dev'essere

$$\sum_{n=1}^{180} 910.05 (1+r_m)^{-n} = 98000,$$

cioè

$$\frac{1-1/(1+r_m)^{180}}{r_m} = 107.69.$$

Risolvendo numericamente (tenendo conto del fatto che sappiamo che $r > 0.006$), si ottiene

$$r_m = 0.00627.$$

In termini di annualità, la banca dice di praticare un interesse annuale pari a $r = (1.006)^{12} - 1 \simeq 7.4\%$ mentre in realtà il tasso effettivo è $r = (1.00627)^{12} - 1 \simeq 7.8\%$.

Soluzione dell'esercizio 1.4. L'interesse mensile è $0.06/12 \times 100\% = 0.005 \times 100\%$. Poniamo $\alpha = 1 + 0.005 = 1.005$. Calcoliamo il patrimonio tra vent'anni: poiché il valore tra 20 anni del k -esimo versamento è $A \alpha^{240-(k-1)}$, alla fine l'individuo accumulerà

$$\sum_{k=1}^{240} A \alpha^{240-(k-1)} = A \alpha \frac{\alpha^{240} - 1}{\alpha - 1}.$$

Questa cifra deve rappresentare il valore attualizzato di 1000 Euro al mese per 30 anni. Dunque, posto $\beta = 1/1.005 = 1/\alpha$, dev'essere

$$A \alpha \frac{\alpha^{240} - 1}{\alpha - 1} = \sum_{k=0}^{359} 1000 \beta^k = 1000 \frac{1 - \beta^{360}}{1 - \beta}.$$

Sostituendo ad esempio $\alpha = 1/\beta$, si ottiene

$$A = 1000 \beta^{240} \frac{1 - \beta^{360}}{1 - \beta^{240}} = 360.99$$

Ciò significa che, ad un tasso del 6% annuo, versando 361 Euro per vent'anni sarà possibile prelevare 1000 Euro al mese per i successivi trent'anni.

Soluzione dell'esercizio 1.5. Se x denota il denaro investito, dopo un tempo di T anni il valore è $x e^{rT}$. Dunque, dev'essere $2x = x e^{rT}$, cioè

$$T = \frac{1}{r} \log 2 = 6.9 \text{ anni.}$$

Soluzione dell'esercizio 1.6. Sia r il tasso di interesse annuale. Alla fine dei tre anni, i due investimenti varranno:

$$\begin{aligned} I_1 &= 100(1+r)^2 + 140(1+r) + 131 \text{ e} \\ I_2 &= 90(1+r)^2 + 160(1+r) + 120. \end{aligned}$$

Il primo investimento è preferibile se e solo se $I_1 > I_2$, cioè, posto $x = 1+r$,

$$100x^2 + 140x + 131 > 90x^2 + 160x + 120.$$

Risolvendo, si ottiene che questa disuguaglianza è vera per ogni valore di x , dunque il primo investimento è senz'altro preferibile.

Capitolo 2

Probabilità: richiami e complementi

2.1 Cenni di teoria della misura

2.1.1 σ -algebre

Sia S un insieme e indichiamo con $\mathcal{P}(S)$ l'insieme delle sue parti (cioè la collezione di tutti i suoi sottoinsiemi). Siano $\mathcal{A}$, $\mathcal{F}$ due collezioni di sottoinsiemi di S : $\mathcal{A}, \mathcal{F} \subset \mathcal{P}(S)$.

Definizione 2.1.1. Si dice che la classe $\mathcal{A}$ è un'algebra su S se valgono le seguenti proprietà:

1. $\emptyset, S \in \mathcal{A}$;
2. se $A \in \mathcal{A}$ allora $A^c \in \mathcal{A}$ (A^c = il complementare dell'insieme A);
3. se $\{A_n\}_{n=1, \dots, N} \subset \mathcal{A}$ allora $\cup_{n=1}^N A_n \in \mathcal{A}$.

Definizione 2.1.2. $\mathcal{F}$ è detta σ -algebra di S se valgono le seguenti proprietà:

1. $\emptyset, S \in \mathcal{F}$;
2. se $A \in \mathcal{F}$ allora $A^c \in \mathcal{F}$;
3. se $\{A_n\}_n \subset \mathcal{F}$ allora $\cup_{n=1}^{\infty} A_n \in \mathcal{F}$.

In breve, una σ -algebra, così come un'algebra, contiene sempre gli insiemi $\emptyset$ e S ed inoltre, è chiusa sotto operazione di complementare. L'unica differenza sta nella terza proprietà: una σ -algebra è chiusa sotto unioni numerabili di insiemi, mentre in un'algebra tale proprietà è vera solo per unioni finite. Ovviamente, una σ -algebra è anche un'algebra ma non vale il viceversa.

Esercizio 2.1. Dimostrare che una σ -algebra $\mathcal{F}$ [risp. un'algebra $\mathcal{A}$] è chiusa sotto intersezioni numerabili [risp. finite] di insiemi di $\mathcal{F}$ [risp. di $\mathcal{A}$].

Naturalmente sono σ -algebrae $\mathcal{F}_1 = \{\emptyset, S\}$ (la più piccola di tutte) e $\mathcal{F}_2 = \mathcal{P}(S)$ (la più grande). Vediamo ora alcuni esempi meno banali.

Esercizio 2.2. Sia $B \subset S$ e poniamo $\mathcal{F} = \{\emptyset, B, B^c, S\}$. Provare che $\mathcal{F}$ è una σ -algebra.

Esercizio 2.3. Siano $\mathcal{F}_1$ e $\mathcal{F}_2$ due σ -algebrae. Provare che:

- $\mathcal{F}_1 \cup \mathcal{F}_2$ non è in generale una σ -algebra (né un'algebra);
- $\mathcal{F}_1 \cap \mathcal{F}_2$ è una σ -algebra.

Per quanto riguarda la seconda proprietà, più in generale si può dire che:

- se $\{\mathcal{F}_\gamma\}_{\gamma \in \Gamma}$ è una famiglia (qualsiasi) di σ -algebrae allora $\mathcal{F} = \bigcap_{\gamma \in \Gamma} \mathcal{F}_\gamma$ è ancora una σ -algebra.

Dunque, mentre l'unione di σ -algebrae non è in generale una σ -algebra, l'intersezione di un numero qualsiasi di σ -algebrae rimane una σ -algebra.

Esempio 2.1.3. (Partizioni finite o numerabili) Sia $\mathcal{C} = \{C_i\}_{i \in I}$ una partizione al più numerabile di S : l'insieme di indici I è finito o numerabile, i sottoinsiemi C_i sono a due a due disgiunti e la loro riunione è S . Sia $\mathcal{F}$ la classe di tutti gli insiemi che sono riunioni (finite o numerabili) degli insiemi C_i , cioè

$$\mathcal{F} = \{A \subset S; A = \bigcup_{i \in I_A} C_i, \text{ per qualche } I_A \subset I\},$$

(con la convenzione di porre $A = \emptyset$ se $I_A = \emptyset$). Allora, $\mathcal{F}$ è una σ -algebra. Infatti, ovviamente $\emptyset, S \in \mathcal{F}$ (si prenda $I_\emptyset = \emptyset$ e $I_S = I$). Prendiamo ora un elemento $A \in \mathcal{F}$. Allora

$$A^c = \left(\bigcup_{i \in I_A} C_i \right)^c = S \setminus \left(\bigcup_{i \in I_A} C_i \right) = \left(\bigcup_{i \in I} C_i \right) \setminus \left(\bigcup_{i \in I_A} C_i \right) = \bigcup_{i \in I_A^c} C_i,$$

dunque $A^c \in \mathcal{F}$. Infine, se $A_1, A_2, \dots \in \mathcal{F}$ allora, posto $A_n = \bigcup_{i \in I_{A_n}} C_i$,

$$\bigcup_n A_n = \bigcup_n \left(\bigcup_{i \in I_{A_n}} C_i \right) = \bigcup_{i \in \bigcup_n I_{A_n}} C_i,$$

e ancora $\bigcup_n A_n \in \mathcal{F}$. Osserviamo infine che i passaggi sopra scritti sono giustificati dal fatto che gli insiemi C_i sono a due a due disgiunti.

Definizione 2.1.4. Sia $\mathcal{C}$ una classe di sottoinsiemi di S . La σ -algebra generata da $\mathcal{C}$, in simboli $\sigma(\mathcal{C})$, è la più piccola σ -algebra di S che contiene la classe $\mathcal{C}$. In altre parole, $\sigma(\mathcal{C})$ è la σ -algebra generata da $\mathcal{C}$ se e solo se $\sigma(\mathcal{C})$ è una σ -algebra di S tale che $\sigma(\mathcal{C}) \supset \mathcal{C}$ e per ogni σ -algebra $\mathcal{F}_{\mathcal{C}}$ di S tale che $\mathcal{F}_{\mathcal{C}} \supset \mathcal{C}$ allora $\mathcal{F}_{\mathcal{C}} \supset \sigma(\mathcal{C})$.

Osserviamo che la Definizione 2.1.4 è ben posta, cioè la “più piccola σ -algebra contenente $\mathcal{C}$ ” esiste sempre. Infatti, consideriamo la classe $\Gamma_{\mathcal{C}}$ di tutte le σ -algebra contenenti $\mathcal{C}$. $\Gamma_{\mathcal{C}}$ è non vuota, dal momento che contiene la σ -algebra $\mathcal{P}(S)$. Inoltre, dall'Esercizio 2.3 segue facilmente che

$$\mathcal{G} = \bigcap_{\mathcal{F} \in \Gamma_{\mathcal{C}}} \mathcal{F}$$

è una σ -algebra. Evidentemente $\mathcal{G}$ è contenuta in ogni σ -algebra contenente $\mathcal{C}$. Dunque $\mathcal{G} = \sigma(\mathcal{C})$ è la più piccola σ -algebra contenente $\mathcal{C}$ e la Definizione 2.1.4 è ben posta.

Esercizio 2.4. Siano $\mathcal{C}, \mathcal{C}_1, \mathcal{C}_2$ collezioni di sottoinsiemi di S .

1. Scrivere esplicitamente $\sigma(\mathcal{C})$ quando $\mathcal{C} = \{A\}$.
2. Dimostrare che se in particolare $\mathcal{C}$ è una σ -algebra allora $\sigma(\mathcal{C}) = \mathcal{C}$.
3. Dimostrare che se $\mathcal{C}_1 \subset \mathcal{C}_2$ allora $\sigma(\mathcal{C}_1) \subset \sigma(\mathcal{C}_2)$.

Esempio 2.1.5. (σ -algebra generata da una partizione numerabile) Sia $\mathcal{C} = \{C_i\}_{i \in I}$ una partizione al più numerabile di S . Allora, la σ -algebra $\mathcal{F}$ costruita nell'Esempio 2.1.3 è la σ -algebra $\sigma(\mathcal{C})$. Verifichiamolo. Intanto $\mathcal{F}$ è una σ -algebra che ovviamente contiene $\mathcal{C}$ (basta prendere $I_A = \{i\}$, per $i \in I$), dunque $\sigma(\mathcal{C}) \subset \mathcal{F}$. Mostriamo ora che $\mathcal{F} \subset \sigma(\mathcal{C})$. Preso $A \in \mathcal{F}$, allora

$$A = \bigcup_{i \in I_A} C_i \in \sigma(\mathcal{C})$$

perché $C_i \in \mathcal{C} \subset \sigma(\mathcal{C})$ per ogni $i \in I_A$ e $\#(I_A) \leq \#(\mathbb{N})$.

Esercizio 2.5. Scrivere esplicitamente $\sigma(\{A, B\})$, con $A, B \subset S$.

Esempio 2.1.6. (σ -algebra di Borel ed insiemi boreliani) Se S è uno spazio topologico, la σ -algebra di Borel su S è la σ -algebra $\mathcal{B}(S)$ generata dagli aperti di S , cioè, se poniamo $\mathcal{O}$ =classe degli insiemi aperti di S , $\mathcal{B}(S) = \sigma(\mathcal{O})$. Un generico elemento di $\mathcal{B}(S)$ prende il nome di insieme di Borel o, più semplicemente, di boreliano.

La σ -algebra di Borel interviene molto spesso in probabilità ed in teoria della misura. In questo esempio sviluppiamo un po' i metodi per lavorarci. Intanto osserviamo che ci sono altre classi di insiemi che generano la σ -algebra di Borel. Ad esempio, se poniamo $\mathcal{C}$ =classe degli insiemi chiusi di S , allora si ha anche $\mathcal{B}(S) = \sigma(\mathcal{C})$. Infatti se chiamiamo $\mathcal{F}$ la σ -algebra generata da $\mathcal{C}$, allora certamente $\mathcal{F}$ contiene $\mathcal{O}$, dato che gli aperti sono i complementari dei chiusi, e dunque anche $\mathcal{B}(S)$, dato che quest'ultima è la più piccola σ -algebra contenente $\mathcal{O}$. Con lo stesso ragionamento si mostra che $\mathcal{B}(S) \supset \mathcal{F}$.

Il caso $S = \mathbb{R}$ è chiaramente di particolare interesse. Non è facile immaginare come siano fatti i boreliani di $\mathbb{R}$. È però di solito facile mostrare che un dato

insieme è un boreliano. Ad esempio sono boreliani gli intervalli semiaperti, cioè della forma

$$(a, b]$$

Infatti è facile verificare che

$$(a, b] = \bigcap_{n=1}^{\infty} (a, b + \frac{1}{n})$$

e dunque l'intervallo semiaperto si può scrivere come intersezione numerabile di aperti.

Si può dimostrare, ma non è proprio semplice, che non tutti i sottoinsiemi di $\mathbb{R}$ sono boreliani. In genere, per gli insiemi che si incontrano di solito, è facile dimostrare che sono boreliani con ragionamenti del tipo appena visto. Altre classi di sottoinsiemi di $\mathbb{R}$ che generano i boreliani sono:

- *Gli intervalli aperti.* Infatti ogni insieme aperto di $\mathbb{R}$ si può scrivere come unione numerabile di intervalli aperti ($\mathbb{R}$ è uno spazio metrico separabile...).
- *Gli intervalli semi aperti.* Infatti ogni intervallo aperto si può scrivere come unione numerabile di intervalli semi aperti:

$$(a, b) = \bigcup_{n=1}^{\infty} (a, b - \frac{1}{n}]$$

- *La classe delle semirette chiuse a destra* Indichiamola con $\pi(\mathbb{R})$, cioè

$$\pi(\mathbb{R}) = \{(-\infty, x]; x \in \mathbb{R}\}.$$

Infatti, mostriamo dapprima che $\mathcal{B}(\mathbb{R}) \supset \sigma(\pi(\mathbb{R}))$. Poniamo $\mathcal{O} = \{\text{insiemi aperti di } \mathbb{R}\}$. Preso $x \in \mathbb{R}$, allora

$$(-\infty, x] = \bigcap_{n=1}^{\infty} \underbrace{(-\infty, x + \frac{1}{n})}_{\in \mathcal{O} \subset \mathcal{B}(\mathbb{R}), \forall n} \in \mathcal{B}(\mathbb{R}).$$

Ciò prova che $\pi(\mathbb{R}) \subset \mathcal{B}(\mathbb{R})$ e quindi $\sigma(\pi(\mathbb{R})) \subset \mathcal{B}(\mathbb{R})$.

Viceversa, proviamo che $\mathcal{B}(\mathbb{R}) \subset \sigma(\pi(\mathbb{R}))$. Per la definizione di σ -algebra generata, basta mostrare che tutti gli intervalli aperti sono contenuti in $\sigma(\pi(\mathbb{R}))$. Ma questo è facile. Infatti $\sigma(\pi(\mathbb{R}))$ contiene gli intervalli semi aperti: presi $a < b$, allora

$$(a, b] = (-\infty, b] \cap (-\infty, a]^c$$

da cui segue che $(a, b] \in \sigma(\pi(\mathbb{R}))$ per ogni $a < b$. Poiché sappiamo che $\mathcal{B}(\mathbb{R})$ è generata dagli intervalli semiaperti, segue che $\mathcal{B}(\mathbb{R}) \subset \sigma(\pi(\mathbb{R}))$.

Esercizio 2.6. Fissato S , sia $\mathcal{A}$ la seguente classe di sottoinsiemi di S : $A \in \mathcal{A}$ se e solo se A è finito oppure A^c è finito. Mostrare che:

1. $\mathcal{A}$ è un'algebra;
2. se S è finito allora $\mathcal{A}$ è una σ -algebra;
3. se S non è finito allora $\mathcal{A}$ non è una σ -algebra.

Esercizio 2.7. Sia S un insieme non finito e $\mathcal{F}$ la seguente classe di sottoinsiemi di S : $A \in \mathcal{F}$ se e solo se A è finito o numerabile oppure A^c è finito o numerabile.

1. Mostrare che $\mathcal{F}$ è una σ -algebra.
2. Si prenda, ad esempio, $S = \mathbb{R}$: mostrare che $\mathcal{F}$ è strettamente contenuta in $\mathcal{B}(\mathbb{R})$.
3. Usare i punti precedenti per trovare un (contro)esempio al fatto che l'unione più che numerabile di insiemi di una σ -algebra non è in generale un elemento della σ -algebra.

2.1.2 Spazi e funzioni misurabili

Definizione 2.1.7. Siano S un insieme e $\mathcal{F}$ una σ -algebra di S . La coppia $(S, \mathcal{F})$ è detta uno *spazio misurabile*.

Definizione 2.1.8. Siano $(S, \mathcal{S})$ e $(E, \mathcal{E})$ due spazi misurabili e sia $f : S \rightarrow E$ una funzione su S a valori in E . f è detta una *funzione misurabile* se la controimmagine tramite f di un qualsiasi insieme $A \in \mathcal{E}$ è un sottoinsieme che appartiene alla σ -algebra $\mathcal{S}$. Ovvero se

$$\text{per ogni } A \in \mathcal{E} \text{ allora } f^{-1}(A) = \{s; f(s) \in A\} \in \mathcal{S},$$

Nel seguito useremo, per indicare la controimmagine di un insieme A una qualunque delle notazioni equivalenti

$$f^{-1}(A) \quad \text{oppure} \quad \{s; f(s) \in A\} \quad \text{oppure} \quad \{f \in A\}.$$

Le funzioni misurabili rivestono un'importanza fondamentale in particolare nel calcolo delle probabilità (più in generale, nella teoria della misura). È chiaro dalla definizione che la proprietà di misurabilità di una funzione dipende dalle σ -algebre $\mathcal{S}$ e $\mathcal{E}$. E infatti, prima di andare avanti, vediamo alcune osservazioni.

- Talvolta ci troveremo in situazioni in cui sull'insieme S ci sono più di una σ -algebra. In questo caso, per precisare rispetto a quale di esse la funzione f è misurabile, diremo che f è $\mathcal{S}$ -misurabile. Parleremo quindi di funzioni $\mathcal{S}$ -misurabili se sull'insieme S si considera la σ -algebra $\mathcal{S}$.

- Altra specifica notazionale. Anche la σ -algebra di riferimento su E può variare. Normalmente, qualora $E = \mathbb{R}$ o anche $E = \mathbb{R}^d$, se non verrà fatto riferimento alla σ -algebra $\mathcal{E}$ significherà che $\mathcal{E} = \mathcal{B}(E)$.
- Spesso per non incorrere ad equivoci si usa anche la notazione $f : (S, \mathcal{S}) \rightarrow (E, \mathcal{E})$, che sottolinea che $f : S \rightarrow E$ è misurabile quando le σ -algebre di riferimento sono $\mathcal{S}$ su S e $\mathcal{E}$ su E .
- Presa $f : (S, \mathcal{S}) \rightarrow (E, \mathcal{E})$ allora per ogni $\mathcal{S}' \supset \mathcal{S}$ e $\mathcal{E}' \subset \mathcal{E}$ si ha anche $f : (S, \mathcal{S}') \rightarrow (E, \mathcal{E}')$. Si noti però che la misurabilità può essere persa se invece si prende $\mathcal{S}' \subset \mathcal{S}$ oppure $\mathcal{E}' \supset \mathcal{E}$.

Definizione 2.1.9. (*σ -algebra generata da una funzione misurabile*) Siano $(S, \mathcal{S})$ e $(E, \mathcal{E})$ due spazi misurabili e sia $f : S \rightarrow E$ una funzione $\mathcal{S}$ -misurabile. La *σ -algebra generata da f* , in simboli $\sigma(f)$, è la più piccola σ -algebra di S rispetto alla quale f risulti misurabile.

In altre parole, $\sigma(f)$ è la σ -algebra tale che f è $\sigma(f)$ misurabile e per ogni σ -algebra $\mathcal{F}_f$ di S tale che f è $\mathcal{F}_f$ misurabile, allora $\mathcal{F}_f \supset \sigma(f)$.

Osservazione 2.1.10. È facile vedere che $\sigma(f)$ è esplicitamente data dagli insiemi che sono controimmagine tramite f di insiemi di $\mathcal{E}$, cioè

$$\sigma(f) = \{f^{-1}(A); A \in \mathcal{E}\}.$$

Vediamo perché. Intanto, poniamo per il momento $\mathcal{F}_f = \{f^{-1}(A); A \in \mathcal{E}\}$ (si noti che, poiché $f^{-1}(A) \subset S$, $\mathcal{F}_f$ è una classe di sottoinsiemi di S) e mostriamo che $\mathcal{F}_f$ è una σ -algebra di S . Innanzitutto essa contiene $\emptyset$ e S , poiché

$$\emptyset = f^{-1}\left(\underbrace{\emptyset}_{\in \mathcal{E}}\right) \in \mathcal{F}_f \quad \text{e} \quad S = f^{-1}\left(\underbrace{E}_{\in \mathcal{E}}\right) \in \mathcal{F}_f.$$

Inoltre, preso $A' \in \mathcal{F}_f$ allora $A' = f^{-1}(A)$ per qualche $A \in \mathcal{E}$, quindi

$$A'^c = (f^{-1}(A))^c = f^{-1}\left(\underbrace{A^c}_{\in \mathcal{E}}\right) \in \mathcal{F}_f.$$

Infine, presi $A'_n \in \mathcal{F}_f$ per $n = 1, 2, \dots$, allora $A'_n = f^{-1}(A_n)$ per qualche $A_n \in \mathcal{E}$ e si vede che

$$\bigcup_n A'_n = \bigcup_n f^{-1}(A_n) = f^{-1}\left(\underbrace{\bigcup_n A_n}_{\in \mathcal{E}}\right) \in \mathcal{F}_f.$$

Abbiamo usato qui due formule che sono sempre valide:

$$(f^{-1}(A))^c = f^{-1}(A^c) \quad \text{e} \quad \bigcup_n f^{-1}(A_n) = f^{-1}\left(\bigcup_n A_n\right) \quad (2.1)$$

La seconda di queste vale anche se alle riunioni si sostituiscono le intersezioni. Attenzione, le stesse formule scritte con f al posto di f^{-1} (e con A , A_n sottoinsiemi di S) *non* sono vere.

Mostriamo ora che $\sigma(f) = \mathcal{F}_f$. Ovviamente, f è $\mathcal{F}_f$ misurabile per costruzione, dunque $\mathcal{F}_f \supset \sigma(f)$. Per verificare l'inclusione opposta, osserviamo che, per ipotesi, f è $\sigma(f)$ -misurabile, quindi $f^{-1}(A)$ deve appartenere a $\sigma(f)$ per ogni $A \in \mathcal{E}$, dunque, poiché tutti gli insiemi di $\mathcal{F}_f$ sono di questa forma, $\mathcal{F}_f \subset \sigma(f)$.

Può succedere che una σ -algebra $\mathcal{F}$ sia definita in un modo poco esplicito (si pensi ad esempio al caso $\mathcal{S} = \mathcal{B}(\mathbb{R})$). Per questo verificare la misurabilità di una funzione usando direttamente la Definizione 2.1.8 può risultare complicato. La proposizione che segue fornisce un criterio pratico molto utile per la verifica della misurabilità:

Proposizione 2.1.11. *Siano $(S, \mathcal{S})$ e $(E, \mathcal{E})$ due spazi misurabili e sia $f : S \rightarrow E$ una funzione.*

1. *Sia $\mathcal{H} \subset \mathcal{E}$ tale che $\sigma(\mathcal{H}) = \mathcal{E}$. Se $f^{-1}(H) \in \mathcal{S}$ per ogni $H \in \mathcal{H}$, allora $f : (S, \mathcal{S}) \rightarrow (E, \mathcal{E})$ è misurabile.*
2. *Supponiamo $E = \mathbb{R}$, $\mathcal{E} = \mathcal{B}(\mathbb{R})$. Se $f^{-1}((-\infty, c]) = \{s : f(s) \leq c\} = \{f \leq c\} \in \mathcal{S}$ per ogni $c \in \mathbb{R}$, allora f è $\mathcal{S}$ -misurabile.*

Dimostrazione. 1. Poniamo

$$\mathcal{G} = \{A \in \mathcal{E}; f^{-1}(A) \in \mathcal{S}\}$$

$\mathcal{G}$ è una classe di sottoinsiemi di E che è contenuta in $\mathcal{E}$, evidentemente. È facile verificare che $\mathcal{G}$ è una σ -algebra. Infatti, $\emptyset \in \mathcal{G}$, $(f^{-1}(\emptyset) = \emptyset \in \mathcal{S})$ e $E \in \mathcal{G}$ ($f^{-1}(E) = S \in \mathcal{S}$). Inoltre, se $A \in \mathcal{G}$ allora $f^{-1}(A) \in \mathcal{S}$, dunque $f^{-1}(A^c) = (f^{-1}(A))^c \in \mathcal{S}$ per la prima delle (2.1) e dunque $A^c \in \mathcal{G}$. Infine, se $A_1, A_2, \dots \in \mathcal{G}$, allora $f^{-1}(A_n) \in \mathcal{S}$ per ogni n , dunque $f^{-1}(\bigcup_n A_n) = \bigcup_n f^{-1}(A_n) \in \mathcal{S}$, per la seconda delle (2.1) e dunque $\bigcup_n A_n \in \mathcal{G}$.

L'ipotesi assicura che $\mathcal{H} \subset \mathcal{G} \subset \mathcal{E}$, da cui segue $\sigma(\mathcal{H}) \subset \mathcal{G} \subset \mathcal{E}$. Poiché per ipotesi $\sigma(\mathcal{H}) = \mathcal{E}$, deve essere $\mathcal{G} = \mathcal{E}$, ovvero: per ogni $A \in \mathcal{E}$ allora $f^{-1}(A) \in \mathcal{S}$, che non è altro che la definizione di $\mathcal{S}$ -misurabilità.

2. Supponiamo $E = \mathbb{R}$, $\mathcal{E} = \mathcal{B}(\mathbb{R})$. La tesi segue dal punto precedente, prendendo $\mathcal{H} = \pi(\mathbb{R})$ e ricordando che $\sigma(\pi(\mathbb{R})) = \mathcal{B}(\mathbb{R})$ (si veda l'Esempio 2.1.6).

□

Esercizio 2.8. *Sia $(S, \mathcal{S})$ uno spazio misurabile e sia $f : S \rightarrow \mathbb{R}$ una funzione. Dimostrare che $f : (S, \mathcal{S}) \rightarrow (\mathbb{R}, \mathcal{B}(\mathbb{R}))$ è misurabile se e solo se $\{f > c\} \in \mathcal{S}$ per ogni $c \in \mathbb{R}$.*

Esercizio 2.9. Sia $f : \mathbb{R}^n \rightarrow \mathbb{R}$ una funzione continua. Dimostrare che f è boreliana, cioè $f : (\mathbb{R}^n, \mathcal{B}(\mathbb{R}^n)) \rightarrow (\mathbb{R}, \mathcal{B}(\mathbb{R}))$ è misurabile.

Esercizio 2.10. Sia f una funzione a scalino: $f(x) = \sum_{i \in I} \alpha_i 1_{A_i}(x)$, dove I è un insieme di indici al più numerabile e per ogni $i \in I$, $\alpha_i \in \mathbb{R}$ e $A_i \in \mathcal{B}(\mathbb{R}^n)$. Dimostrare che f è boreliana, cioè $f : (\mathbb{R}^n, \mathcal{B}(\mathbb{R}^n)) \rightarrow (\mathbb{R}, \mathcal{B}(\mathbb{R}))$ è misurabile.

Esempio 2.1.12. (*Funzioni discrete: misurabilità*) Siano $(S, \mathcal{S})$ e $(E, \mathcal{E})$ due spazi misurabili e sia $f : S \rightarrow E$ una funzione. Supponiamo che l'insieme $E = \{e_1, e_2, \dots\}$ sia discreto e che $\mathcal{E} = \mathcal{P}(E)$. Allora, f è misurabile se e solo se

$$\{f = e_i\} \in \mathcal{S} \quad \text{per ogni } e_i \in E.$$

Infatti, fissato $e_i \in E$, ovviamente $\{f = e_i\} = f^{-1}(\{e_i\}) \in \mathcal{S}$ perché tutti gli insiemi della forma $\{e_i\}$ appartengono a $\mathcal{E}$ per ogni $e_i \in E$. Viceversa, supponiamo che $\{f = e_i\} \in \mathcal{S}$ per ogni $e_i \in E$. Prendiamo allora $\mathcal{H} = \{\{e_1\}, \{e_2\}, \dots\}$: basta osservare che $\sigma(\mathcal{H}) = \mathcal{E}$ ed applicare la Proposizione 2.1.11.

Esempio 2.1.13. (*Funzioni discrete: σ -algebra generata*) Consideriamo la stessa situazione dell'esempio precedente. Vogliamo vedere come è fatta la σ -algebra $\sigma(f)$ generata da f . Più precisamente, poniamo $C_i = \{f = e_i\}$. Evidentemente $C_i \in \mathcal{S}$. Mostriamo che $\mathcal{C} = (C_i)_i$ è una partizione di S e che $\sigma(f)$ coincide con la σ -algebra generata dalla partizione $\mathcal{C}$, che abbiamo visto nell'Esempio 2.1.5

Mostriamo che $\mathcal{C}$ è una partizione. Se $i \neq j$, allora se esistesse un elemento $s \in S$ appartenente a $C_i \cap C_j$, dovrebbe essere simultaneamente $f(s) = e_i$ e $f(s) = e_j$, che è impossibile. Dunque C_i e C_j devono essere disgiunti. Inoltre se $s \in S$, allora esiste $e_i \in E$ tale che $f(s) = e_i$ e dunque $s \in C_i$. Ne segue che la riunione degli insiemi C_i è uguale a S .

Mostriamo ora che $\sigma(f) = \sigma(\mathcal{C})$. Al solito useremo il metodo della doppia inclusione.

1) $\sigma(f) \subset \sigma(\mathcal{C})$. Se $A' \in \sigma(f)$ allora (cfr. Esempio 2.1.9) si ha, per qualche $A \in \mathcal{E}$, $A' = f^{-1}(A)$. Possiamo allora scrivere $A = \bigcup_{e_i \in A} \{e_i\}$. Dunque, sempre usando la seconda delle (2.1)

$$A' = f^{-1}(A) = \bigcup_{e_i \in A} f^{-1}(\{e_i\}) = \bigcup_{e_i \in A} C_i \in \sigma(\mathcal{C}),$$

2) $\sigma(\mathcal{C}) \subset \sigma(f)$. Se $A \in \sigma(\mathcal{C})$ allora per qualche $I_A \subset I$ si ha (si veda l'Esempio 2.1.9)

$$A = \bigcup_{i \in I_A} \{f = e_i\} = f^{-1}\left(\bigcup_{i \in I_A} \{e_i\}\right) \in \sigma(f)$$

perché $\bigcup_{i \in I_A} \{e_i\} \in \mathcal{E}$.

Riassumiamo nella proposizione che segue alcune proprietà basilari delle funzioni misurabili.

Proposizione 2.1.14. 1. Se f, f_1, f_2 sono funzioni $\mathcal{S}$ -misurabili e se $\alpha \in \mathbb{R}$ allora¹ $f_1 + f_2, f_1 \cdot f_2, \alpha f$ sono $\mathcal{S}$ -misurabili.

2. Siano $(S, \mathcal{S}), (E, \mathcal{E})$ e $(F, \mathcal{F})$ spazi misurabili. Se $f : (S, \mathcal{S}) \rightarrow (E, \mathcal{E})$ è misurabile e $h : (E, \mathcal{E}) \rightarrow (F, \mathcal{F})$ è misurabile allora la funzione composta $h \circ f : (S, \mathcal{S}) \rightarrow (F, \mathcal{F})$ è misurabile.

3. Se $\{f_n\}_n$ è una successione di funzioni $\mathcal{S}$ -misurabili a valori in $\mathbb{R}$, allora $\inf_n f_n, \sup_n f_n, \liminf_n f_n, \limsup_n f_n$ sono² $\mathcal{S}$ -misurabili. Inoltre,

$$\{s : \text{esiste } \lim_n f_n(s)\} \in \mathcal{S}.$$

Dimostrazione. 1. Per verificare la prima proprietà, basta osservare che

$$\{f_1 + f_2 > c\} = \bigcup_{q \in \mathbb{Q}} (\{f_1 > q\} \cap \{f_2 > c - q\}).$$

Usando (opportunamente) la 2. della Proposizione 2.1.11, si ottiene la tesi. Per le altre, si ragiona in modo analogo.

2. Preso $\Lambda \in \mathcal{F}$, si ha

$$\begin{aligned} (h \circ f)^{-1}(\Lambda) &= f^{-1}(\underbrace{h^{-1}(\Lambda)}_{\substack{\in \mathcal{E} \text{ [} h \text{ è} \\ \mathcal{E}\text{-mis.}]}}) \in \mathcal{S}, \\ &\underbrace{\hspace{10em}}_{\in \mathcal{S} \text{ [} f \text{ è } \mathcal{S}\text{-mis.]}]} \end{aligned}$$

quindi $h \circ f$ è $\mathcal{S}$ -misurabile.

3. Basta osservare che, per ogni $c \in \mathbb{R}$,

$$\{\inf_n f_n > c\} = \bigcap_n \{f_n > c\}$$

ed usare 2. della proposizione 2.1.11. Analogamente si mostra la $\mathcal{S}$ -misurabilità di $\sup_n f_n, \liminf_n f_n$ e $\limsup_n f_n$. Infine

$$\begin{aligned} &\{s : \text{esiste } \lim_n f_n(s)\} \\ &= \{\liminf_n f_n > -\infty\} \cap \{\limsup_n f_n < +\infty\} \cap \{\liminf_n f_n = \limsup_n f_n\} \\ &= \{\liminf_n f_n > -\infty\} \cap \{\limsup_n f_n < +\infty\} \cap \{\liminf_n f_n - \limsup_n f_n = 0\} \end{aligned}$$

¹Qui, la notazione “.” ha senso per $d = 1$.

²Tali funzioni sono, in generale, a valori in $[-\infty, +\infty]$: ad essere precisi, occorrerebbe definire la σ -algebra $\mathcal{B}([-\infty, +\infty])$. Essa si definisce come la σ -algebra generata dagli aperti di $[-\infty, +\infty]$, o anche come la σ -algebra generata dalla classe $\mathcal{C}$ formata da $\mathcal{B}(\mathbb{R})$ e dai singleton $\{-\infty\}$ e $\{+\infty\}$. Ma non vogliamo entrare più in dettaglio...

che è un elemento di $\mathcal{S}$ perché $\liminf_n f_n$, $\limsup_n f_n$ sono $\mathcal{S}$ -misurabili e dunque anche $\liminf_n f_n - \limsup_n f_n$ lo è.

□

Esercizio 2.11. Siano $(S, \mathcal{S})$ e $(E, \mathcal{E})$ spazi misurabili e $f, g : (S, \mathcal{S}) \rightarrow (E, \mathcal{E})$ due funzioni misurabili. Supponiamo che esistano due funzioni misurabili $h_1, h_2 : (E, \mathcal{E}) \rightarrow (E, \mathcal{E})$ tali che $f = h_1 \circ g$ e $g = h_2 \circ f$. Dimostrare allora che $\sigma(f) = \sigma(g)$.

Nella proposizione che segue, vediamo che quando $\mathcal{S}$ è una σ -algebra generata da una partizione allora la $\mathcal{S}$ -misurabilità dà informazioni molto precise su com'è fatta f . Infatti,

Proposizione 2.1.15. Sia $f : S \rightarrow \mathbb{R}^d$ una funzione misurabile. Se f è $\sigma(\mathcal{C})$ -misurabile, con $\mathcal{C}$ partizione al più numerabile di S , allora f è costante su ciascun elemento C di $\mathcal{C}$.

Dimostrazione. Supponiamo, per assurdo, che esista $C \in \mathcal{C}$ tale che $s_1, s_2 \in C$, con $s_1 \neq s_2$, e $a_1 = f(s_1) \neq f(s_2)$. Poniamo $A_1 = f^{-1}(\{a_1\})$. Poiché $\{a_1\} \in \mathcal{B}(\mathbb{R})$, si ha $A_1 \in \mathcal{S} = \sigma(\mathcal{C})$. Ricordando che $s_1 \in A_1$ e che $\sigma(\mathcal{C})$ è costituita da tutte le possibili unioni al più numerabili di elementi di $\mathcal{C}$ (cfr. Esempio 2.1.5), necessariamente dev'essere $A_1 \supseteq C$, il che è assurdo perché $s_2 \in C$ ma $s_2 \notin A_1$.

□

Esercizio 2.12. Sia $f : S \rightarrow \mathbb{R}^d$ una funzione misurabile. Dimostrare che f è costante se e solo se $\sigma(f) = \{\emptyset, \Omega\}$ (la σ -algebra banale).

Presentiamo un ulteriore risultato, di cui omettiamo la dimostrazione (per la quale si rimanda, ad esempio, al Paragrafo A.3.2 di Williams [10]).

Proposizione 2.1.16. Siano $f : S \rightarrow \mathbb{R}^d$, e $g : S \rightarrow \mathbb{R}^m$ due funzioni $\mathcal{S}$ -misurabili. Supponiamo che g sia in particolare $\sigma(f)$ -misurabile. Allora, esiste una funzione $h : \mathbb{R}^d \rightarrow \mathbb{R}^m$ che sia $\mathcal{B}(\mathbb{R}^d)$ -misurabile tale che $g = h \circ f$.

2.2 Spazi di probabilità

Definizione 2.2.1. Uno spazio di probabilità è una tripla $(\Omega, \mathcal{F}, \mathbb{P})$ costituita da:

- uno spazio (astratto) Ω , detto *spazio campionario*;
- una σ -algebra $\mathcal{F}$ su Ω , che rappresenta l'insieme di tutti i possibili eventi; un elemento $A \in \mathcal{F}$ prende quindi il nome di *evento*;

- una *misura di probabilità*, o semplicemente una *probabilità*, $\mathbb{P}$ sullo spazio misurabile $(\Omega, \mathcal{F})$, finita e di massa 1:

$\mathbb{P} : \mathcal{F} \rightarrow [0, 1]$ è tale che

$$* \mathbb{P}(\Omega) = 1,$$

$$* \mathbb{P}(\cup_{n=1}^{\infty} A_n) = \sum_{n=1}^{\infty} \mathbb{P}(A_n), \text{ se } A_n \cap A_m = \emptyset \text{ per ogni } n \neq m.$$

La seconda proprietà è nota col nome di σ -*additività*.

Osserviamo che poiché $A_1, A_2, \dots \in \mathcal{F}$ e $\mathcal{F}$ è una σ -algebra, allora $A = \bigcup_n A_n \in \mathcal{F}$, dunque $\mathbb{P}$ è effettivamente definita su $\mathcal{F}$.

Osserviamo inoltre che, alla luce di quanto visto in precedenza, la coppia $(\Omega, \mathcal{F})$ è uno spazio di misura.

Esempio 2.2.2. Supponiamo $\mathcal{F} = \sigma(\mathcal{C})$, dove $\mathcal{C}$ è una partizione al più numerabile di Ω . Allora, una qualsiasi probabilità $\mathbb{P}$ su $(\Omega, \mathcal{F}) = (\Omega, \sigma(\mathcal{C}))$ è univocamente individuata dal valore che assume sugli elementi della partizione $\mathcal{C}$. Cioè, posto $\mathcal{C} = \{C_i; i \in I\}$, allora $\mathbb{P}$ è perfettamente individuata dai numeri

$$0 \leq p_i = \mathbb{P}(C_i), \quad i \in I, \quad \text{tali che} \quad \sum_{i \in I} p_i = 1.$$

Questo perché ogni $A \in \sigma(\mathcal{C})$ è della forma $A = \bigcup_{i \in I_A} C_i$, con $I_A \subseteq I$ finito o numerabile, dunque la σ -additività di $\mathbb{P}$ garantisce che

$$\mathbb{P}(A) = \mathbb{P}\left(\bigcup_{i \in I_A} C_i\right) = \sum_{i \in I_A} \mathbb{P}(C_i) = \sum_{i \in I_A} p_i.$$

2.2.1 Prime proprietà

Vediamo alcune proprietà elementari della probabilità.

Proposizione 2.2.3. *Sia $(\Omega, \mathcal{F}, \mathbb{P})$ uno spazio di probabilità.*

1. *Se $A, B \in \mathcal{F}$ sono tali che $A \subset B$ allora $\mathbb{P}(A) \leq \mathbb{P}(B)$.*
2. *(Sub-additività) Per ogni $A_1, A_2 \in \mathcal{F}$ allora $\mathbb{P}(A_1 \cup A_2) \leq \mathbb{P}(A_1) + \mathbb{P}(A_2)$. Più in generale, per ogni $A_1, A_2, \dots \in \mathcal{F}$,*

$$\mathbb{P}\left(\bigcup_n A_n\right) \leq \sum_n \mathbb{P}(A_n).$$

3. *Per ogni $A_1, A_2, \dots \in \mathcal{F}$ tali che $\mathbb{P}(A_n) = 0$ per ogni n , allora*

$$\mathbb{P}\left(\bigcup_n A_n\right) = 0.$$

4. Per ogni $A \in \mathcal{F}$, $\mathbb{P}(A^c) = 1 - \mathbb{P}(A)$.
5. (Formula di inclusione-esclusione) Per ogni $A_1, A_2 \in \mathcal{F}$ allora $\mathbb{P}(A_1 \cup A_2) = \mathbb{P}(A_1) + \mathbb{P}(A_2) - \mathbb{P}(A_1 \cap A_2)$. Più in generale, per ogni $A_1, \dots, A_N \in \mathcal{F}$,

$$\begin{aligned} \mathbb{P}\left(\bigcup_{n=1}^N A_n\right) &= \sum_{n=1}^N \mathbb{P}(A_n) - \sum_{n < m \leq N} \mathbb{P}(A_n \cap A_m) \\ &+ \sum_{n < m < \ell \leq N} \mathbb{P}(A_n \cap A_m \cap A_\ell) - \dots + (-1)^{N-1} \mathbb{P}(A_1 \cap \dots \cap A_N). \end{aligned}$$

Dimostrazione. 1. Poiché $B = A \cup (A^c \cap B)$ e quest'ultima unione è disgiunta,

$$\mathbb{P}(B) = \mathbb{P}(A) + \mathbb{P}(A^c \cap B) \geq \mathbb{P}(A).$$

2. Poniamo $B_1 = A_1$ e per $n > 1$, $B_n = A_n \cap (\bigcup_{k=1}^{n-1} A_k)^c$. B_n sono a due a due disgiunti, $\bigcup_n B_n = \bigcup_n A_n$ ed inoltre $B_n \subseteq A_n$ per ogni n . Quindi,

$$\mathbb{P}(\bigcup_n A_n) = \mathbb{P}(\bigcup_n B_n) = \sum_n \mathbb{P}(B_n) \leq \sum_n \mathbb{P}(A_n).$$

3. Segue immediatamente da 2.

4. Poiché $\Omega = A \cup A^c$, con A, A^c ovviamente disgiunti, si ha $1 = \mathbb{P}(\Omega) = \mathbb{P}(A) + \mathbb{P}(A^c)$, quindi $\mathbb{P}(A) = 1 - \mathbb{P}(A^c)$.

5. $A_1 \cup A_2 = (A_1 \cap A_2^c) \cup (A_1^c \cap A_2) \cup (A_1 \cap A_2)$. Poiché si tratta di elementi di $\mathcal{F}$ a due a due disgiunti, otteniamo

$$\mathbb{P}(A_1 \cup A_2) = \mathbb{P}(A_1 \cap A_2^c) + \mathbb{P}(A_1^c \cap A_2) + \mathbb{P}(A_1 \cap A_2).$$

Ora, $A_1 = (A_1 \cap A_2^c) \cup (A_1 \cap A_2)$ ed essendo disgiunti, $\mathbb{P}(A_1) = \mathbb{P}(A_1 \cap A_2^c) + \mathbb{P}(A_1 \cap A_2)$. Analogamente, $\mathbb{P}(A_2) = \mathbb{P}(A_1^c \cap A_2) + \mathbb{P}(A_1 \cap A_2)$. Sostituendo, si ottiene la tesi.

La formula generale di inclusione-esclusione è lasciata per esercizio.

Osserviamo che, poiché $\mathcal{F}$ è una σ -algebra, tutti gli insiemi presi in questa dimostrazione di cui viene valutata la probabilità $\mathbb{P}$ appartengono alla σ -algebra $\mathcal{F}$, cioè appartengono effettivamente al dominio di definizione dell'applicazione $\mathbb{P}$.

□

Nella teoria della probabilità (e più in generale in teoria della misura), particolare rilevanza hanno gli insiemi di probabilità nulla, per i quali sottolineiamo vale la proprietà 3. della Proposizione 2.2.3: l'unione al più numerabile di eventi di probabilità nulla ha ancora probabilità nulla.

Presentiamo inoltre le seguenti proprietà di convergenza monotona della probabilità. A tale scopo, ricordiamo che, prese due successioni $\{a_n\}_n$ e $\{b_n\}_n$ di numeri reali, si scrive:

- $a_n \uparrow a$ se $a_n \leq a_{n+1}$ per ogni n e $a = \sup_n a_n = \lim_n a_n$;
- $b_n \downarrow b$ se $b_n \geq b_{n+1}$ per ogni n e $b = \inf_n b_n = \lim_n b_n$.

In qualche modo si può generalizzare a successioni $\{A_n\}_n$ e $\{B_n\}_n$ di insiemi, nel modo seguente:

- $A_n \uparrow A$ se $A_n \subset A_{n+1}$ per ogni n e $A = \cup_n A_n$;
- $B_n \downarrow B$ se $B_n \supset B_{n+1}$ per ogni n e $B = \cap_n B_n$.

Allora,

Proposizione 2.2.4. *Sia $(\Omega, \mathcal{F}, \mathbb{P})$ uno spazio di probabilità e siano $\{A_n\}_n, \{B_n\}_n \subset \mathcal{F}$.*

1. Se $A_n \uparrow A$ allora $\mathbb{P}(A_n) \uparrow \mathbb{P}(A)$.
2. Se $B_n \downarrow B$ allora $\mathbb{P}(B_n) \downarrow \mathbb{P}(B)$.

Dimostrazione. 1. Ovviamente $\mathbb{P}(A_n)$ è una successione (numerica) crescente, perché $A_n \subset A_{n+1}$. Poniamo

$$A'_1 = A_1 \quad \text{e per } n > 1, \quad A'_n = A_n \setminus A_{n-1}.$$

Allora: gli A'_n sono a due a due disgiunti, tali che $A = \cup_n A_n = \cup_n A'_n$ ed inoltre $A_n = \cup_{k=1}^n A'_k$. Ma allora,

$$\begin{aligned} \mathbb{P}(A) &= \mathbb{P}(\cup_n A'_n) = \sum_n \mathbb{P}(A'_n) = \lim_{n \rightarrow \infty} \sum_{k=1}^n \mathbb{P}(A'_k) \\ &= \lim_{n \rightarrow \infty} \mathbb{P}(\cup_{k=1}^n A'_k) = \lim_{n \rightarrow \infty} \mathbb{P}(A_n) = \sup_n \mathbb{P}(A_n). \end{aligned}$$

2. Poniamo $A_n = B_n^c$: $A_n \uparrow A = \cup_n A_n = \cup_n B_n^c = (\cap_n B_n)^c = B^c$, quindi da 1. segue che $\mathbb{P}(A_n) \uparrow \mathbb{P}(B^c)$. Allora, $\mathbb{P}(B_n) = 1 - \mathbb{P}(A_n) \downarrow 1 - \mathbb{P}(B^c) = \mathbb{P}(B)$. $\square$

2.2.2 Indipendenza

Sia $(\Omega, \mathcal{F}, \mathbb{P})$ uno spazio di probabilità. Ricordiamo che n eventi $A_1, \dots, A_n \in \mathcal{F}$ si dicono indipendenti se per ogni sottoinsieme $\{i_1, \dots, i_k\}$ di $\{1, \dots, n\}$ si ha

$$\mathbb{P}(A_{i_1} \cap \dots \cap A_{i_k}) = \mathbb{P}(A_{i_1}) \dots \mathbb{P}(A_{i_k}).$$

Ad esempio, limitiamoci al caso $n = 2$ e prendiamo due eventi indipendenti A e B : $\mathbb{P}(A \cap B) = \mathbb{P}(A)\mathbb{P}(B)$. Allora,

$$\begin{aligned} \mathbb{P}(A \cap B^c) &= \mathbb{P}(A) - \mathbb{P}(A \cap B) = \mathbb{P}(A) - \mathbb{P}(A)\mathbb{P}(B) \\ &= \mathbb{P}(A)(1 - \mathbb{P}(B)) = \mathbb{P}(A)\mathbb{P}(B^c). \end{aligned}$$

Analogamente, si fa vedere che $\mathbb{P}(A^c \cap B) = \mathbb{P}(A^c)\mathbb{P}(B)$ e $\mathbb{P}(A^c \cap B^c) = \mathbb{P}(A^c)\mathbb{P}(B^c)$. Ciò significa che, preso comunque un elemento $G_1 \in \mathcal{G}_1 = \sigma(\{A\}) = \{\emptyset, \Omega, A, A^c\}$ e un elemento $G_2 \in \mathcal{G}_2 = \sigma(\{B\}) = \{\emptyset, \Omega, B, B^c\}$, allora $\mathbb{P}(G_1 \cap G_2) = \mathbb{P}(G_1)\mathbb{P}(G_2)$, cioè G_1 e G_2 sono indipendenti.

Ciò giustifica le seguente definizione, più generale, di indipendenza.

Definizione 2.2.5. Siano $\mathcal{G}_1, \dots, \mathcal{G}_n$ sotto σ -algebre i $\mathcal{F}$, cioè $\mathcal{G}_i \subset \mathcal{F}$ è una σ -algebra, $i = 1, \dots, n$. $\mathcal{G}_1, \dots, \mathcal{G}_n$ si dicono *indipendenti* se

$$\mathbb{P}(G_1 \cap \dots \cap G_n) = \mathbb{P}(G_1) \dots \mathbb{P}(G_n), \quad \text{per ogni } G_1 \in \mathcal{G}_1, \dots, G_n \in \mathcal{G}_n.$$

Pensando a $\mathcal{G}_1, \dots, \mathcal{G}_n$ come a n possibili esperimenti, la Definizione 2.2.5 dice che gli n esperimenti sono indipendenti se presi comunque n possibili risultati $G_1 \in \mathcal{G}_1, \dots, G_n \in \mathcal{G}_n$ allora gli eventi $G_1, \dots, G_n$ sono indipendenti.

Esempio 2.2.6. Supponiamo $\mathcal{G}_i = \sigma(\mathcal{C}_i)$, con $\mathcal{C}_i$ partizione al più numerabile di Ω , per $i = 1, \dots, n$. Allora, $\mathcal{G}_1, \dots, \mathcal{G}_n$ sono indipendenti se e solo se

$$\mathbb{P}(C_1 \cap \dots \cap C_n) = \mathbb{P}(C_1) \dots \mathbb{P}(C_n), \quad \text{per ogni } C_1 \in \mathcal{C}_1, \dots, C_n \in \mathcal{C}_n$$

In altre parole, quando le σ -algebre sono generate da una partizione al più numerabile, è sufficiente provare l'indipendenza tra gli atomi che compongono le partizioni.

Per semplicità di notazione, dimostriamolo solo nel caso $n = 2$. Siano quindi $G_1 \in \mathcal{G}_1$ e $G_2 \in \mathcal{G}_2$. Allora

$$\begin{aligned} G_1 &= \cup_{j \in I_1} A_j, & A_j &\in \mathcal{C}_1, j \in I_1 \\ G_2 &= \cup_{k \in I_2} B_k, & B_k &\in \mathcal{C}_2, k \in I_2 \end{aligned}$$

e si ha:

$$\begin{aligned} \mathbb{P}(G_1 \cap G_2) &= \mathbb{P}\left(\cup_{j \in I_1, k \in I_2} (A_j \cap B_k)\right) = \sum_{j \in I_1, k \in I_2} \mathbb{P}(A_j \cap B_k) \\ &= \sum_{j \in I_1} \sum_{k \in I_2} \mathbb{P}(A_j)\mathbb{P}(B_k) = \mathbb{P}\left(\cup_{j \in I_1} A_j\right)\mathbb{P}\left(\cup_{k \in I_2} B_k\right) = \mathbb{P}(G_1)\mathbb{P}(G_2). \end{aligned}$$

2.3 Variabili aleatorie

Abbiamo visto che se $(\Omega, \mathcal{F}, \mathbb{P})$ è uno spazio di probabilità allora in particolare $(\Omega, \mathcal{F})$ è uno spazio misurabile. Le variabili aleatorie sono delle funzioni definite su Ω , e in particolare misurabili. Infatti,

Definizione 2.3.1. Sia $(E, \mathcal{E})$ uno spazio misurabile. Una *variabile aleatoria* (v.a.) X su $(\Omega, \mathcal{F}, \mathbb{P})$ è una funzione

$$X : (\Omega, \mathcal{F}) \longrightarrow (E, \mathcal{E}) \quad \text{misurabile.}$$

Una v.a. X è detta *discreta* se l'insieme dei valori che assume $E_X = X(\Omega) \equiv \{x \in \mathbb{R} : x = X(\omega), \omega \in \Omega\}$ è un insieme discreto, cioè finito o numerabile. In particolare, diremo che una v.a. discreta X è *finita* se E_X è un insieme finito. Senza perdere in generalità, in tal caso sceglieremo $\mathcal{E} = \mathcal{P}(E_X)$.

Quando non tratteremo v.a. discrete, saranno per lo più continue e a valori in $E = \mathbb{R}$ oppure $E = \mathbb{R}^d$, con la scelta di $\mathcal{E} = \mathcal{B}(E)$. Dunque, una v.a. X è semplicemente una funzione su Ω a valori in $\mathbb{R}$ o in $\mathbb{R}^d$ che sia $\mathcal{F}$ -misurabile.

Anche nel caso di v.a. si definisce l'indipendenza:

Definizione 2.3.2. n v.a. $X_1, \dots, X_n$ si dicono indipendenti se

$$\mathbb{P}(X_1 \in A_1, \dots, X_n \in A_n) = \mathbb{P}(X_1 \in A_1) \cdots \mathbb{P}(X_n \in A_n)$$

per ogni scelta di $A_1, \dots, A_n$ tali che le probabilità sopra scritte abbiano senso, e cioè $A_1 \in \mathcal{E}_1, \dots, A_n \in \mathcal{E}_n$ (si sta infatti supponendo che X_i sia a valori nello spazio misurabile $(E_i, \mathcal{E}_i)$).

Esercizio 2.13. Ricordando che $\sigma(X_i) = \{X_i^{-1}(A); A \in \mathcal{E}_i\}$, verificare che la Definizione 2.3.2 equivale alla seguente: $X_1, \dots, X_n$ sono indipendenti se e solo se $\sigma(X_1), \dots, \sigma(X_n)$ sono indipendenti.

A volte, ci troveremo a parlare di una v.a. indipendente da una σ -algebra. Sulla base delle definizioni viste, il modo più naturale di descrivere matematicamente tale situazione è il seguente:

Definizione 2.3.3. Sia X una v.a. e $\mathcal{G}$ una sotto σ -algebra di $\mathcal{F}$. Diremo che X e $\mathcal{G}$ sono indipendenti se $\sigma(X)$ e $\mathcal{G}$ sono indipendenti, o in modo equivalente se

$$\mathbb{P}((X \in A) \cap G) = \mathbb{P}(X \in A)\mathbb{P}(G) \quad \text{per ogni } A \in \mathcal{E}, G \in \mathcal{G}.$$

Nei paragrafi che seguono, ricordiamo brevemente le proprietà principali e gli strumenti comunemente utilizzati nella trattazione delle variabili aleatorie.

2.3.1 Proprietà

Poiché una v.a. X è semplicemente una funzione misurabile, valgono tutte le proprietà viste nel Paragrafo 2.1.2, che riassumiamo qui di seguito.

- La σ -algebra generata da X , in simboli $\sigma(X)$, ovvero la più piccola σ -algebra contenuta in $\mathcal{F}$ rispetto alla quale X rimane misurabile, è data da

$$\sigma(X) = \{X^{-1}(A); A \in \mathcal{E}\}.$$

In particolare, se X è discreta allora

$$\sigma(X) = \sigma(\mathcal{C}_X), \quad \text{dove } \mathcal{C}_X = \{X^{-1}(\{x\}); x \in E_X = X(\Omega)\},$$

e ricordiamo che $\mathcal{C}_X$ è una partizione al più numerabile di Ω (cfr. Osservazione 2.1.10 ed Esempio 2.1.13).

- Sia $X : \Omega \rightarrow \mathbb{R}$. X è una v.a. se e solo se $X^{-1}((-\infty, c]) \in \mathcal{F}$ per ogni $c \in \mathbb{R}$ (cfr. Proposizione 2.1.11).
- In particolare, se X è discreta allora X è una v.a. se e solo se $X^{-1}(\{x\}) \in \mathcal{F}$ per ogni $x \in E_X = X(\Omega)$ (cfr. Esempio 2.1.12).
- Se $X, Y, \{X_n\}_n$ sono v.a., $\alpha, \beta \in \mathbb{R}$ e $h : E_X \rightarrow \mathbb{R}^m$ è misurabile allora (cfr. Proposizione 2.1.14): $\alpha X + \beta Y, X \cdot Y$ sono v.a.; $h \circ X$ è una v.a.; $\inf_n X_n, \sup_n X_n, \liminf_n X_n, \limsup_n X_n$ sono v.a. ed inoltre $\{\omega : \text{esiste } \lim_n X_n(\omega)\} \in \mathcal{F}$.
- Se $\sigma(X) = \sigma(\mathcal{C})$ con $\mathcal{C}$ partizione al più numerabile di Ω allora X è discreta ed assume valore costante su ogni elemento C di $\mathcal{C}$ (cfr. Proposizione 2.1.15).
- Se X è una v.a. a valori in $\mathbb{R}^d$ e Y è una v.a. a valori in $\mathbb{R}^m$ e $\sigma(X)$ -misurabile allora esiste $h : \mathbb{R}^d \rightarrow \mathbb{R}^m$ misurabile e tale che $Y = h(X)$ (cfr. Proposizione 2.1.16).

2.3.2 Legge e distribuzione

Com'è noto, nella pratica il formalismo introdotto per le v.a. non viene usato. Infatti, quando si lavora con variabili aleatorie in realtà spesso si perde di vista qual è l'originario spazio di probabilità $(\Omega, \mathcal{F}, \mathbb{P})$ e X vista come *applicazione* (misurabile) da Ω su E , anzi tali ingredienti neanche si conoscono. Piuttosto, si lavora con la distribuzione di X e con la sua densità (discreta o continua). Oppure, con la legge di X , che dà il legame tra variabili aleatorie e probabilità. Ricordiamo infatti che nei precedenti corsi di probabilità si è lavorato con v.a. bernoulliane, binomiali, geometriche oppure gaussiane, esponenziali etc. Ovvero, si sta specificando qual è la legge della v.a. presa in considerazione. In generale, la *legge* della v.a. $X = (X_1, \dots, X_d)$ a valori in $E \subset \mathbb{R}^d$, o equivalentemente la *legge congiunta* di $X_1, \dots, X_d$, è la misura di probabilità $\mathbb{P}_X : \mathcal{E} \rightarrow [0, 1]$ definita da

$$A \mapsto \mathbb{P}_X(A) = \mathbb{P}(X \in A), \quad A \in \mathcal{E}.$$

Esercizio 2.14. Verificare che effettivamente $\mathbb{P}_X$ è una probabilità, non più su $(\Omega, \mathcal{F})$ bensì su $(E, \mathcal{E})$.

Nei due paragrafi che seguono, studiamo o meglio rivediamo un po' più in dettaglio due esempi particolarmente rilevanti: il caso in cui X è una v.a. discreta ed il caso in cui X è assolutamente continua.

Variabili aleatorie discrete

In questo caso, X è una v.a. a valori in $E = E_X = X(\Omega) = \{x^1, x^2, \dots\}$, con $x^i = (x_1^i, \dots, x_d^i)$, per ogni i , e la σ -algebra di riferimento su E si prende, senza perdere in generalità, come $\mathcal{E} = \mathcal{P}(E)$.

Quando si lavora con v.a. discrete, particolarmente utile è la *densità discreta* di X , o equivalentemente la *densità discreta congiunta* di $X_1, \dots, X_d$. Essa è data da $p_X : E_X \rightarrow [0, 1]$, definita come

$$p_X(x) = \mathbb{P}(X = x) = \mathbb{P}(X = x_1, \dots, X_d = x_d), \quad x = (x_1, \dots, x_d) \in E_X.$$

In generale, per comodità è bene definire p_X anche al di fuori di E_X , ponendo ovviamente $p_X(x) = 0$ se $x \notin E_X$.

Ovviamente,

$$p_X(x) > 0 \text{ per ogni } x \in E_X \text{ e } \sum_{x \in E_X} p_X(x) = 1$$

e la legge di X è

$$\mathbb{P}_X(A) = \sum_{x \in A} p_X(x), \quad A \in \mathcal{P}(E).$$

Nota p_X , è completamente noto il comportamento (probabilistico) di X . Inoltre, se p è una funzione su un insieme discreto E e tale che

$$p(x) > 0 \text{ per ogni } x \in E \text{ e } \sum_{x \in E} p(x) = 1$$

allora (Teorema di Skorohod) si può dimostrare che esiste uno spazio di probabilità $(\Omega, \mathcal{F}, \mathbb{P})$ su cui è definita una v.a. X tale che $E_X = E$ e $p_X = p$.

Riassumiamo brevemente alcune delle proprietà principali della densità discreta p_X .

1. Nota p_X , è nota la densità discreta di un qualsiasi sotto-vettore di $X = (X_1, \dots, X_d)$, e in particolare quindi tutte le densità marginali. Ad esempio, supponiamo per semplicità $d = 3$: $X = (X_1, X_2, X_3)$. La v.a. (X_1, X_2) assume valori in $\{z = (z_1, z_2) : \text{esiste } i \text{ tale che } z_1 = x_1^i, z_2 = x_2^i\}$ e la densità discreta (congiunta) di X_1 e X_2 è data da

$$p_{X_1, X_2}(z_1, z_2) = \sum_{i : x_1^i = z_1, x_2^i = z_2} p_X(x^i) = \sum_i p_X(z_1, z_2, x_3^i).$$

Procedendo analogamente, la densità marginale, ad esempio, di X_1 è invece data da

$$p_{X_1}(\xi) = \sum_{i : x_1^i = \xi} p_X(x^i) = \sum_i p_X(\xi, x_2^i, x_3^i).$$

2. Le v.a. (discrete) $X_1, \dots, X_d$ sono indipendenti se e solo se per ogni $x = (x_1, \dots, x_d) \in E_X$,

$$p_X(x_1, \dots, x_d) = p_{X_1}(x_1) \cdots p_{X_d}(x_d)$$

dove p_{X_i} denota l' i -esima densità marginale.

3. Supponiamo, per semplicità, $d = 2$. La densità di $Y = X_1 + X_2$ è data da

$$p_Y(y) = \sum_{i: x_1^i + x_2^i = y} p_X(x_1^i, x_2^i) = \sum_t p_X(t, y - t) = \sum_t p_X(y - t, t).$$

Se in particolare X_1 e X_2 sono indipendenti, allora

$$p_Y(y) = \sum_t p_{X_1}(t) p_{X_2}(y - t) = \sum_t p_{X_1}(y - t) p_{X_2}(t).$$

4. Supponiamo X e Y v.a. discrete, l'una a valori in $E_X \subset \mathbb{R}^d$ e l'altra in $E_Y \subset \mathbb{R}^m$. La *densità condizionale* di X dato Y è data da

$$p_{X|Y}(x|y) = \frac{p_{X,Y}(x, y)}{p_Y(y)}$$

purché ovviamente $p_Y(y) > 0$, dove $p_{X,Y}$ e p_Y denotano, rispettivamente, la densità di $(X, Y) = (X_1, \dots, X_d, Y_1, \dots, Y_m)$ e quella (marginale) di Y . Ricordiamo che, fissato $y \in E_Y$ tale che $p_Y(y) > 0$, $p_{X|Y}(x|y)$ è una densità discreta su E_X , cioè

$$p_{X|Y}(x|y) \geq 0 \text{ per ogni } x \in E_X \text{ e } \sum_{x \in E_X} p_{X|Y}(x|y) = 1$$

Variabili aleatorie assolutamente continue

In questo caso, X assume valori in $E = \mathbb{R}^d$ e la σ -algebra di riferimento su E si prende come $\mathcal{E} = \mathcal{B}(E)$. Inoltre, dire che X è assolutamente continua significa supporre l'esistenza della *densità di probabilità* p_X di $X = (X_1, \dots, X_d)$, o equivalente la *densità di probabilità congiunta* di $X_1, \dots, X_d$:

$$p_X : \mathbb{R}^d \rightarrow [0, +\infty) \text{ è integrabile su } \mathbb{R}^d \text{ e tale che}$$

$$\mathbb{P}(X \in A) = \int_A p_X(x) dx, \quad A \in \mathcal{E}.$$

Dunque, la legge $\mathbb{P}_X$ di una v.a. X assolutamente continua si rappresenta in forma integrale:

$$\mathbb{P}_X(A) = \mathbb{P}(X \in A) = \int_A p_X(x) dx, \quad A \in \mathcal{E}.$$

Ovviamente, la densità p_X è una funzione integrabile tale che

$$p_X(x) \geq 0 \quad \text{e} \quad \int_{\mathbb{R}^d} p_X(x) dx = 1.$$

Nota p_X , è completamente noto il comportamento (probabilistico) di X . Inoltre, se p è una funzione integrabile su $\mathbb{R}^d$ tale che

$$p(x) \geq 0 \text{ per ogni } x \text{ e } \int_{\mathbb{R}^d} p(x) dx = 1$$

allora (Teorema di Skorohod) si può dimostrare che esiste uno spazio di probabilità $(\Omega, \mathcal{F}, \mathbb{P})$ su cui è definita una v.a. assolutamente continua X con densità di probabilità $p_X = p$.

Riassumiamo brevemente alcune delle proprietà principali della densità (continua) p_X . Osserviamo che si tratta più o meno delle stesse proprietà già ricordate per una densità discreta: l'unica accortezza è quella di sostituire le somme con degli integrali!

1. Nota p_X , allora qualsiasi sotto-vettore di $X = (X_1, \dots, X_d)$ è assolutamente continuo, con densità ricavabile da p_X . In particolare, ogni componente X_i ha densità. Ad esempio, supponiamo per semplicità $X = (X_1, X_2, X_3)$: la densità (congiunta) di X_1 e X_2 è data da

$$p_{X_1, X_2}(x_1, x_2) = \int_{\mathbb{R}} p_X(x_1, x_2, x_3) dx_3$$

e la densità marginale, ad esempio, di X_1 è invece data da

$$p_{X_1}(x_1) = \int_{\mathbb{R}^2} p_X(x_1, x_2, x_3) dx_2 dx_3$$

Ricordiamo che non vale il viceversa: se esiste la densità (marginale) di X_1 e di X_2 , non è detto che esista la densità (congiunta) di (X_1, X_2) .

2. Le v.a. (congiuntamente) assolutamente continue $X_1, \dots, X_d$ sono indipendenti se e solo se per ogni $x = (x_1, \dots, x_d) \in \mathbb{R}^d$,

$$p_X(x_1, \dots, x_d) = p_{X_1}(x_1) \cdots p_{X_d}(x_d)$$

dove p_{X_i} denota l' i -esima densità marginale.

3. Supponiamo, per semplicità, $d = 2$. Se $X = (X_1, X_2)$ è assolutamente continua, allora $Y = X_1 + X_2$ è assolutamente continua ed ha densità

$$p_Y(y) = \int_{\mathbb{R}} p_X(t, y-t) dt = \int_{\mathbb{R}} p_X(y-t, t) dt$$

e se in particolare X_1 e X_2 sono indipendenti, allora

$$p_Y(y) = \int_{\mathbb{R}} p_{X_1}(t) p_{X_2}(y-t) dt = \int_{\mathbb{R}} p_{X_1}(y-t) p_{X_2}(t) dt.$$

4. Supponiamo X e Y v.a. congiuntamente assolutamente continue, l'una a valori in $\mathbb{R}^d$ e l'altra in $\mathbb{R}^m$. La densità condizionale di X dato Y è data da

$$p_{X|Y}(x|y) = \frac{p_{X,Y}(x,y)}{p_Y(y)}$$

purché ovviamente $p_Y(y) > 0$, dove $p_{X,Y}$ e p_Y denotano, rispettivamente, la densità di $(X,Y) = (X_1, \dots, X_d, Y_1, \dots, Y_m)$ e quella (marginale) di Y . Ricordiamo che, fissato y tale che $p_Y(y) > 0$, $p_{X|Y}(x|y)$ è una densità (continua) su $\mathbb{R}^d$, cioè

$$x \mapsto p_{X|Y}(x|y) \geq 0 \text{ è integrabile su } \mathbb{R}^d \text{ e } \int_{\mathbb{R}^d} p_{X|Y}(x|y) dx = 1$$

5. Nel caso di v.a. non discrete, particolare rilievo assume la funzione di distribuzione o di ripartizione (congiunta) $F_X : \mathbb{R}^d \rightarrow [0, 1]$: per $x = (x_1, \dots, x_d) \in \mathbb{R}^d$,

$$F_X(x) = \mathbb{P}(X_1 \leq x_1, \dots, X_d \leq x_d) = \mathbb{P}_X\left((-\infty, x_1] \times \dots \times (-\infty, x_d]\right).$$

Ricordiamo in particolare che X è assolutamente continua se e solo se esiste

$$g(x) = \frac{\partial^d}{\partial x_1 \dots \partial x_d} F_X(x),$$

con g integrabile su $\mathbb{R}^d$ tale che $\int_{\mathbb{R}^d} g(x) dx = 1$, e in tal caso

$$p_X(x) = g(x) \text{ per quasi ogni } x,$$

dove la dicitura “per quasi ogni x ” significa per tutti gli $x \in \mathbb{R}^d \setminus \mathcal{N}$, dove $\mathcal{N}$ è un insieme di $\mathbb{R}^d$ al più numerabile.

2.3.3 Speranza matematica

In questo paragrafo riassumiamo la definizione e le proprietà principali della speranza matematica. Prendiamo quindi una v.a. X , che per il momento supponiamo essere 1-dimensionale (X a valori in $\mathbb{R}$).

Ricordiamo quanto sviluppato nei precedenti corsi di probabilità.

- Sia X discreta, a valori in E_X , e sia p_X la densità discreta. Si dice che X ha *media* oppure *speranza matematica* oppure *aspettazione finita* se

$$\sum_{x \in E_X} |x| p_X(x) < \infty$$

e in tal caso, la *media* oppure *speranza matematica* oppure *aspettazione* di X è

$$\mathbb{E}(X) = \sum_{x \in E_X} x p_X(x).$$

- Sia X assolutamente continua e sia p_X la densità di probabilità (continua). Si dice che X ha *media* oppure *speranza matematica* oppure *aspettazione finita* se

$$\int_{\mathbb{R}} |x| p_X(x) dx < \infty$$

e in tal caso, la *media* oppure *speranza matematica* oppure *aspettazione* di X è

$$\mathbb{E}(X) = \int_{\mathbb{R}} x p_X(x) dx.$$

Osservazione 2.3.4. Supponiamo che X sia una v.a. discreta definita su uno spazio di probabilità discreto $(\Omega, \mathcal{F}, \mathbb{P})$, cioè Ω è finito o numerabile e $\mathcal{F} = \mathcal{P}(\Omega)$. In tal caso, per ogni $x \in E_X$, possiamo scrivere

$$p_X(x) = \mathbb{P}(X^{-1}(\{x\})) = \sum_{\omega \in X^{-1}(\{x\})} \mathbb{P}(\{\omega\}).$$

Allora,

$$\begin{aligned} \sum_{x \in E_X} |x| p_X(x) &= \sum_{x \in E_X} |x| \sum_{\omega \in X^{-1}(\{x\})} \mathbb{P}(\{\omega\}) \\ &= \sum_{x \in E_X} \sum_{\omega \in X^{-1}(\{x\})} |X(\omega)| \mathbb{P}(\{\omega\}) \\ &= \sum_{\omega \in \cup_{x \in E_X} X^{-1}(\{x\})} |X(\omega)| \mathbb{P}(\{\omega\}) \\ &= \sum_{\omega \in \Omega} |X(\omega)| \mathbb{P}(\{\omega\}) \end{aligned}$$

e, analogamente,

$$\sum_{x \in E_X} x p_X(x) = \sum_{\omega \in \Omega} X(\omega) \mathbb{P}(\{\omega\})$$

Allora, X ha speranza finita se e solo se $\sum_{\omega \in \Omega} |X(\omega)| \mathbb{P}(\{\omega\}) < \infty$ e in tal caso la speranza matematica è data da

$$\mathbb{E}(X) = \sum_{\omega \in \Omega} X(\omega) \mathbb{P}(\{\omega\}).$$

La quantità $\sum_{\omega \in \Omega} X(\omega) \mathbb{P}(\{\omega\})$ viene anche denotata con $\int_{\Omega} X(\omega) d\mathbb{P}$, così che

$$\mathbb{E}(X) = \int_{\Omega} X(\omega) d\mathbb{P}.$$

In realtà, la quantità sopra scritta è la vera definizione di *speranza matematica della v.a. X* , a prescindere che X sia discreta o assolutamente continua. È evidente che, perché abbia senso, occorre definire l'integrale rispetto a $\mathbb{P}$. Ciò può essere fatto: poiché $\mathbb{P}$ è una misura in generale astratta, la teoria della misura consente di sviluppare l'integrale rispetto ad una misura

qualsiasi, e la speranza matematica è semplicemente l'integrale rispetto a $\mathbb{P}$ della v.a. X . Questi argomenti saranno sviluppati nei corsi superiori di probabilità. Ciò nonostante, per le medie di v.a. discrete o assolutamente continue si ritroveranno le formule qui proposte, che rappresentano delle formule davvero operative: ricordiamo infatti che, nella pratica, spesso non si conoscono Ω , $\mathcal{F}$ e $\mathbb{P}$ ma si conosce la densità (discreta o continua) di X .

Riassumiamo qui di seguito alcune delle proprietà principali (di cui omettiamo le dimostrazioni rimandando ad un qualsiasi testo base in probabilità).

1. Sia X una v.a. discreta (risp. assolutamente continua), con densità discreta (risp. densità continua) p_X . Sia Y una v.a. su $\mathbb{R}$ della forma $Y = f(X)$. Allora Y ha speranza matematica finita se e solo se

$$\sum_{x \in E_X} |f(x)| p_X(x) dx < \infty \quad (\text{risp. } \int_{\mathbb{R}} |f(x)| p_X(x) dx < \infty)$$

e in tal caso, la speranza matematica di Y è

$$\sum_{x \in E_X} f(x) p_X(x) dx < \infty \quad (\text{risp. } \int_{\mathbb{R}} f(x) p_X(x) dx < \infty).$$

2. Nel seguito, X e Y denotano due v.a. con speranza finita. Si ha:

- (linearità) se $\alpha, \beta \in \mathbb{R}$ allora $\alpha X + \beta Y$ ha speranza finita e $\mathbb{E}(\alpha X + \beta Y) = \alpha \mathbb{E}(X) + \beta \mathbb{E}(Y)$;
- (positività) se $\mathbb{P}(X \geq 0) = 1$ allora $\mathbb{E}(X) \geq 0$; quindi, se $\mathbb{P}(X \geq Y) = 1$ allora $\mathbb{E}(X) \geq \mathbb{E}(Y)$;
- si ha sempre: $|\mathbb{E}(X)| \leq \mathbb{E}(|X|)$;
- se X e Y sono indipendenti allora $\mathbb{E}(XY) = \mathbb{E}(X)\mathbb{E}(Y)$; più in generale, X e Y sono indipendenti se e solo se prese due funzioni f e g tali che esistono le speranze matematiche di $f(X)$ e $g(Y)$ allora esiste la speranza di $f(X)g(Y)$ ed inoltre $\mathbb{E}(f(X)g(Y)) = \mathbb{E}(f(X))\mathbb{E}(g(Y))$.

In particolare, si dice che X ha *momento k -esimo finito* se esiste $\mathbb{E}(|X|^k)$ ed in tal caso $\mathbb{E}(X^k)$ prende il nome di momento k -esimo. È facile vedere che se esiste il momento k -esimo allora esiste il momento j -esimo, per ogni $j \leq k$.

Particolarmente rilevante è il caso $k = 2$, in cui si può definire la *varianza*:

$$\text{Var}(X) = \mathbb{E}\left((X - \mathbb{E}(X))^2\right) = \mathbb{E}(X^2) - \mathbb{E}(X)^2.$$

Ricordiamo che la media dà la migliore approssimazione costante della v.a. nel senso del momento secondo, e la varianza dà l'errore che si commette sostituendo la v.a. con questa costante. In termini matematici:

$$\inf_{c \in \mathbb{R}} \mathbb{E}\left((X - c)^2\right) = \min_{c \in \mathbb{R}} \mathbb{E}\left((X - c)^2\right) = \mathbb{E}\left((X - \mathbb{E}(X))^2\right) = \text{Var}(X).$$

Ricordiamo che valgono anche le seguenti disuguaglianze:

- (Markov) Se $p > 0$ e $|X|^p$ ha speranza finita allora per ogni $c > 0$ si ha

$$\mathbb{P}(X > c) \leq \frac{\mathbb{E}(|X|^p)}{c^p}.$$

- (Chebycev) Se X ha momento secondo finito allora per ogni $c > 0$ si ha

$$\mathbb{P}(|X - \mathbb{E}(X)| > c) \leq \frac{\text{Var}(X)}{c^2}.$$

- (Jensen) Se φ è una funzione convessa allora $\mathbb{E}(\varphi(X)) \geq \varphi(\mathbb{E}(X))$.
- (Schwarz) Se esistono le speranze di X^2 e Y^2 allora XY ha speranza finita e $|\mathbb{E}(XY)| \leq \mathbb{E}(|XY|) \leq \mathbb{E}(X^2) \mathbb{E}(Y^2)$.

Infine, se supponiamo che X sia una v.a. d -dimensionale, cioè $X = (X_1, \dots, X_d)$, discreta oppure assolutamente continua, definiamo la speranza matematica di X come il vettore

$$\mathbb{E}(X) = (\mathbb{E}(X_1), \dots, \mathbb{E}(X_d)).$$

Dunque, X ha speranza matematica finita se e solo se ciascuna componente X_i ha speranza matematica finita e la speranza di X è il vettore composto dalle speranze matematiche delle componenti. Ovviamente, tutte le proprietà viste continuano a valere (dando un senso opportuno alle quantità prima scalari e che ora diventano vettoriali). Ricordiamo, perché sarà spesso utilizzata in seguito, solo la seguente proprietà: presa $X = (X_1, \dots, X_d)$ di densità discreta (risp. continua) p_X , allora $Y = f(X_1, \dots, X_d)$ ha speranza finita se e solo se

$$\sum_{(x_1, \dots, x_d) \in E_X} |f(x_1, \dots, x_d)| p_X(x_1, \dots, x_d) < \infty$$

$$\left(\text{risp. } \int_{\mathbb{R}^d} |f(x_1, \dots, x_d)| p_X(x_1, \dots, x_d) dx_1, \dots, dx_d < \infty \right)$$

e in tal caso

$$\mathbb{E}(f(X_1, \dots, X_d)) = \sum_{(x_1, \dots, x_d) \in E_X} f(x_1, \dots, x_d) p_X(x_1, \dots, x_d)$$

$$\left(\text{risp. } \mathbb{E}(f(X_1, \dots, X_d)) = \int_{\mathbb{R}^d} f(x_1, \dots, x_d) p_X(x_1, \dots, x_d) dx_1, \dots, dx_d \right).$$

2.4 Soluzioni

Soluzione dell'esercizio 2.1. Se $\{A_n\}_n \in \mathcal{F}$ allora

$$\left(\bigcap_n A_n \right)^c = \bigcup_n \underbrace{A_n^c}_{\in \mathcal{F}, \forall n} \in \mathcal{F},$$

dunque $\cap_n A_n \in \mathcal{F}$.

Soluzione dell'esercizio 2.2. La dimostrazione è immediata: se $A \in \mathcal{F}$ allora anche $A^c \in \mathcal{F}$ e una qualsiasi unione di elementi di $\mathcal{F}$ appartiene a $\mathcal{F}$.

Soluzione dell'esercizio 2.3. Proviamo, con un controesempio, che l'unione di σ -algebre non è in generale una σ -algebra. Si prendano, ad esempio, $A, B \subset S$ e siano $\mathcal{F}_1 = \{\emptyset, A, A^c, S\}$ e $\mathcal{F}_2 = \{\emptyset, B, B^c, S\}$. $\mathcal{F}_1$ e $\mathcal{F}_2$ sono due σ -algebre (Esercizio 2.2) e $\mathcal{F} = \mathcal{F}_1 \cup \mathcal{F}_2 = \{\emptyset, A, A^c, B, B^c, S\}$: se $A \neq B, B^c$, con $A, B \neq \emptyset, S$, allora ad esempio $A \cup B \notin \mathcal{F}$, dunque $\mathcal{F}$ non è una σ -algebra né un'algebra.

L'intersezione di una collezione qualsiasi di σ -algebre $\mathcal{F} = \cap_{\gamma \in \Gamma} \mathcal{F}_\gamma$ è invece una σ -algebra. Infatti, $S \in \mathcal{F}$ perché $S \in \mathcal{F}_\gamma$ per ogni $\gamma \in \Gamma$ ($\mathcal{F}_\gamma$ è una σ -algebra), dunque $S \in \cap_{\gamma \in \Gamma} \mathcal{F}_\gamma = \mathcal{F}$. Sia ora $\{A_n\}_n \subset \mathcal{F}$: $A_n \in \mathcal{F}_\gamma$ per ogni n e γ , dunque $\cup_n A_n \in \mathcal{F}_\gamma$ per ogni γ e allora $\cup_n A_n \in \cap_{\gamma \in \Gamma} \mathcal{F}_\gamma = \mathcal{F}$. Infine, se $A \in \mathcal{F}$ allora $A \in \mathcal{F}_\gamma$ per ogni γ , dunque $A^c \in \mathcal{F}_\gamma$ per ogni γ e allora $A^c \in \cap_{\gamma \in \Gamma} \mathcal{F}_\gamma = \mathcal{F}$.

Soluzione dell'esercizio 2.4. 1. Si ha: $\sigma(\mathcal{C}_1) = \{\emptyset, S, A, A^c\}$.

Per gli altri punti, ricordiamo dapprima che $\sigma(\mathcal{C}) = \cap_{\mathcal{F} \in \mathcal{I}_{\mathcal{C}}} \mathcal{F}$, dove $\mathcal{I}_{\mathcal{C}}$ denota l'insieme delle σ -algebre contenenti $\mathcal{C}$. Dunque, 2. è immediata conseguenza del fatto che, essendo $\mathcal{C}$ una σ -algebra, allora $\mathcal{C} \in \mathcal{I}_{\mathcal{C}}$. Per 3., basta osservare che poiché $\mathcal{C}_1 \subset \mathcal{C}_2 \subset \sigma(\mathcal{C}_2)$ allora $\sigma(\mathcal{C}_2) \in \mathcal{I}_{\mathcal{C}_1}$, dunque $\sigma(\mathcal{C}_2) \supset \sigma(\mathcal{C}_1)$.

Soluzione dell'esercizio 2.5. Posto $C_1 = A \cap B^c$, $C_2 = A^c \cap B$, $C_3 = A \cap B$, $C_4 = A^c \cap B^c$ ed infine $\mathcal{C} = \{C_1, C_2, C_3, C_4\}$, allora si ha $\sigma(\{A, B\}) = \sigma(\mathcal{C})$. Infatti, $\{A, B\} \subset \sigma(\mathcal{C})$, dunque $\sigma(\{A, B\}) \subset \sigma(\mathcal{C})$; poiché $\mathcal{C} \subset \sigma(\{A, B\})$ si ha anche $\sigma(\mathcal{C}) \subset \sigma(\{A, B\})$. Ora, poiché $\mathcal{C}$ è una partizione di S , dall'Esempio 2.1.5 segue immediatamente che $\sigma(\mathcal{C}) = \{\emptyset, S, C_1, C_2, C_3, C_4, C_1 \cup C_2, C_1 \cup C_3, C_1 \cup C_4, C_2 \cup C_3, C_2 \cup C_4, C_3 \cup C_4, C_1 \cup C_2 \cup C_3, C_1 \cup C_2 \cup C_4, C_1 \cup C_3 \cup C_4, C_2 \cup C_3 \cup C_4\}$.

Soluzione dell'esercizio 2.6. 1. $S \in \mathcal{A}$: se S è finito, niente da dire; se invece S non è finito, $S^c = \emptyset$ è finito. $\mathcal{A}$ è chiusa sotto operazione di complementare: banale. Infine, siano $A, B \in \mathcal{A}$ e mostriamo che $A \cup B \in \mathcal{A}$. Occorre suddividere in casi:

- se A e B sono entrambi finiti allora $A \cup B$ è finito, dunque $A \cup B \in \mathcal{A}$;
- se A^c e B^c sono entrambi finiti allora $(A \cup B)^c = A^c \cap B^c$ è finito, dunque $A \cup B \in \mathcal{A}$;
- se A e B^c sono entrambi finiti allora: se B è finito anche $A \cup B$ è finito, dunque $A \cup B \in \mathcal{A}$; se invece B non è finito allora $(A \cup B)^c = A^c \cap B^c \subset B^c$ è finito e ancora $A \cup B \in \mathcal{A}$;
- se A^c e B sono entrambi finiti allora, procedendo in modo analogo a quanto visto sopra, segue facilmente che $A \cup B \in \mathcal{A}$.

2. Se S è finito, $\#\mathcal{P}(S) = 2^{\#S} < \infty$ ($\mathcal{P}(S)$ = insieme delle parti di S) quindi presa una qualsiasi successione $\{A_n\}_n$ di sottoinsiemi di S allora esiste un indice n_0 tale che per ogni $n > n_0$, A_n coincide con $\emptyset$ oppure S oppure con un qualche A_k , con $k \leq n_0$. Ciò significa che se $\{A_n\}_n \subset \mathcal{A}$ allora $\cup_n A_n = \cup_{k=1}^N B_k$ con $B_k \in \mathcal{A}$ opportuni e da 1. segue che $\cup_n A_n \in \mathcal{A}$.

3. Se S non è finito, sia $\{s_1, s_2, \dots\}$ un insieme numerabile di S . Posto $A = \{s_1, s_3, \dots, s_{2k+1}, \dots\}$ allora $A^c \supset \{s_2, s_4, \dots, s_{2k}, \dots\}$, quindi $A \notin \mathcal{A}$. Ma $A = \cup_{k \geq 0} \{s_{2k+1}\}$ e $\{s_{2k+1}\} \in \mathcal{A}$ per ogni $k \geq 1$, quindi $\mathcal{A}$ non è una σ -algebra.

Soluzione dell'esercizio 2.7. 1. Occorre solo dimostrare che $\mathcal{F}$ è chiusa sotto unioni numerabili. Sia $\{A_n\}_n \subset \mathcal{F}$. Se A_n finito o numerabile per ogni n allora $\cup_n A_n$ è finito o numerabile, quindi $\cup_n A_n \in \mathcal{F}$. Altrimenti, cioè se esiste un ℓ tale che A_ℓ non è finito né numerabile, allora A_ℓ^c lo è e $(\cup_n A_n)^c = \cap_n A_n^c \subset A_\ell^c$ è finito o numerabile, e ancora $\cup_n A_n \in \mathcal{F}$.

2. Sia $A \in \mathcal{F}$. Allora o A o A^c sono numerabili, cioè $A = \{a_k\}_k = \cup_k \{a_k\}$ oppure $A^c = \{a_k\}_k = \cup_k \{a_k\}$, con $a_k \in S = \mathbb{R}$ per ogni k . Ora, $\{a_k\} \in \mathcal{B}(\mathbb{R})$:

$$\{a_k\} = \cap_n \left(a_k - \frac{1}{n}, a_k + \frac{1}{n} \right)$$

da cui segue che o A oppure A^c appartiene a $\mathcal{B}(\mathbb{R})$, quindi $A \in \mathcal{B}(\mathbb{R})$. L'inclusione è stretta: basta considerare l'intervallo $(0, 1)$, che è in $\mathcal{B}(\mathbb{R})$ ma non in $\mathcal{F}$. Si noti che $(0, 1) = \cup_{x: x \in (0, 1)} \{x\} \notin \mathcal{F}$ pur essendo $\{x\} \in \mathcal{F}$ per ogni $x \in (0, 1)$: l'unione più che numerabile di elementi di una σ -algebra non appartiene in generale alla σ -algebra. Ciò risponde al punto 3.

Soluzione dell'esercizio 2.8. Essendo $\{f > c\} = f^{-1}((c, +\infty))$, se f è misurabile allora ovviamente $\{f > c\} \in \mathcal{S}$, perché $(c, +\infty) \in \mathcal{B}(\mathbb{R})$. Viceversa, basta usare il punto 1. della Proposizione 2.1.11 prendendo $\mathcal{H} = \{(c, +\infty); c \in \mathbb{R}\}$. Infatti, si ha che $\sigma(\mathcal{H}) = \mathcal{B}(\mathbb{R})$, perché $\mathcal{H} \subset \sigma(\pi(\mathbb{R})) = \mathcal{B}(\mathbb{R})$, dunque $\sigma(\mathcal{H}) \subset \mathcal{B}(\mathbb{R})$, e $\pi(\mathbb{R}) \subset \sigma(\mathcal{H})$, dunque $\mathcal{B}(\mathbb{R}) = \sigma(\pi(\mathbb{R})) \subset \sigma(\mathcal{H})$.

Oppure, più semplicemente, basta usare il punto 2. della Proposizione 2.1.11 ed osservare che $\{f > c\} = \{f \leq c\}^c$.

Soluzione dell'esercizio 2.9. Se f è continua, per ogni $c \in \mathbb{R}$ l'insieme $\{f \leq c\}$ è un chiuso. Infatti, se $\{x_n\}_n \subset \{f \leq c\}$ e $x_n \rightarrow x$ per $n \rightarrow \infty$, allora $f(x) = \lim_{n \rightarrow \infty} f(x_n) \leq c$, dunque $x \in \{f \leq c\}$. Ora, la tesi segue direttamente applicando la 2. della Proposizione 2.1.11.

Soluzione dell'esercizio 2.10. Useremo ancora la 2. della Proposizione 2.1.11. Sia $c \in \mathbb{R}$. Definiamo $I_c = \{i \in I : \alpha_i \leq c\}$. Allora, $\{f \leq c\} = \cup_{i \in I_c} A_i \in \mathcal{B}(\mathbb{R}^n)$ in quanto unione al più numerabile di elementi di $\mathcal{B}(\mathbb{R}^n)$.

Soluzione dell'esercizio 2.11. In particolare $g : (S, \sigma(g)) \rightarrow (E, \mathcal{E})$ è misurabile, dunque per il punto 2. della Proposizione 2.1.14 $f = h_1 \circ g$ è

$\sigma(g)$ -misurabile. Dunque, $\sigma(f) \subset \sigma(g)$. Scambiando i ruoli a f e g e a h_1 e h_2 , si ottiene $\sigma(g) \subset \sigma(f)$, da cui la tesi.

Soluzione dell'esercizio 2.12. Supponiamo che f sia costante: $f \equiv c$, $c \in \mathbb{R}^d$. Allora, per ogni boreliano A di $\mathbb{R}^d$, si ha:

$$f^{-1}(A) = \begin{cases} \emptyset & \text{se } c \notin A \\ S & \text{se } c \in A \end{cases}$$

dunque $\sigma(f) = \{f^{-1}(A) : A \in \mathcal{B}(\mathbb{R}^d)\} = \{\emptyset, S\}$.

Viceversa, supponiamo che $\sigma(f) = \{\emptyset, S\}$ e dimostriamo che f è costante. Supponiamo quindi per assurdo che f assuma due valori diversi, diciamo c_1 e c_2 . Allora, posto $A_i = f^{-1}(\{c_i\})$ per $i = 1, 2$, allora $A_1 \cap A_2 = \emptyset$ e $A_i \in \sigma(f)$ per $i = 1, 2$, il che è assurdo a meno che uno dei due sia il vuoto (dunque il c_i corrispondente in realtà non è assunto).

Soluzione dell'esercizio 2.13. Supponiamo che $X_1, \dots, X_n$ siano indipendenti. Presi $G_1 \in \sigma(X_1), \dots, G_n \in \sigma(X_n)$, allora esistono $A_1 \in \mathcal{E}_1, \dots, A_n \in \mathcal{E}_n$ tali che $G_i = X_i^{-1}(A_i)$, per $i = 1, \dots, n$. Allora

$$\begin{aligned} \mathbb{P}(G_1 \cap \dots \cap G_n) &= \mathbb{P}(X_1 \in A_1, \dots, X_n \in A_n) \\ &= \mathbb{P}(X_1 \in A_1) \cdots \mathbb{P}(X_n \in A_n) = \mathbb{P}(G_1) \cdots \mathbb{P}(G_n), \end{aligned}$$

dunque $\sigma(X_1), \dots, \sigma(X_n)$ sono indipendenti.

Viceversa, supponiamo $\sigma(X_1), \dots, \sigma(X_n)$ indipendenti. Allora, per ogni $A_1 \in \mathcal{E}_1, \dots, A_n \in \mathcal{E}_n$, si ha che $G_i = \{X_i \in A_i\} \in \sigma(X_i)$, $i = 1, \dots, n$. Allora,

$$\begin{aligned} \mathbb{P}(X_1 \in A_1, \dots, X_n \in A_n) &= \mathbb{P}(G_1 \cap \dots \cap G_n) \\ &= \mathbb{P}(G_1) \cdots \mathbb{P}(G_n) = \mathbb{P}(X_1 \in A_1) \cdots \mathbb{P}(X_n \in A_n), \end{aligned}$$

dunque $X_1, \dots, X_n$ sono indipendenti.

Soluzione dell'esercizio 2.14. Sicuramente $\mathbb{P}_X(A) \in [0, 1]$, ed inoltre $\mathbb{P}_X(E) = \mathbb{P}(X \in E) = 1$. Poi, se $A_1, A_2, \dots \in \mathcal{E}$ sono a due a due disgiunti, allora gli eventi $\{X \in A_1\}, \{X \in A_2\}, \dots$ rimangono a due a due disgiunti e

$$\begin{aligned} \mathbb{P}_X(\cup_n A_n) &= \mathbb{P}(X \in \cup_n A_n) \\ &= \mathbb{P}(\cup_n \{X \in A_n\}) = \sum_n \mathbb{P}(X \in A_n) = \sum_n \mathbb{P}_X(A_n). \end{aligned}$$

Dunque, $\mathbb{P}_X : \mathcal{E} \rightarrow [0, 1]$ ha massa 1 ed è σ -additiva, cioè una probabilità.

Capitolo 3

Condizionamento e martingale

3.1 Condizionamento

Sia $(\Omega, \mathcal{F}, \mathbb{P})$ uno spazio di probabilità. Dunque, ricordiamo,

Ω è un insieme (lo *spazio campionario*);

$\mathcal{F}$ è una σ -algebra di sottoinsiemi di Ω (gli *eventi*);

$\mathbb{P}$ è una probabilità su $(\Omega, \mathcal{F})$, cioè $\mathbb{P}$ è un'applicazione $\mathcal{F} \rightarrow [0, 1]$ t.c.

$$\mathbb{P}(\Omega) = 1,$$

$$\text{se } A_n \cap A_m = \emptyset, n \neq m, \mathbb{P}\left(\bigcup_{n=1}^{\infty} A_n\right) = \sum_{n=1}^{\infty} \mathbb{P}(A_n) \text{ } (\sigma\text{-additività}).$$

Ricordiamo che, dati due eventi $A, B \in \mathcal{F}$ con $\mathbb{P}(B) > 0$, la probabilità condizionata di A dato B è definita da

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

Intuitivamente, $\mathbb{P}(A|B)$ rappresenta la probabilità del verificarsi dell'evento A , una volta noto il fatto che l'evento B si è verificato.

Vogliamo ora definire la *media condizionale* di una v.a. X data una σ -algebra $\mathcal{G} \subset \mathcal{F}$ e, in seguito, la *probabilità condizionale* di un evento qualsiasi $A \in \mathcal{F}$ data una σ -algebra $\mathcal{G}$. Poiché si tratta di oggetti alquanto delicati e spesso inizialmente di difficile comprensione, procederemo per gradi.

3.1.1 Media condizionale: il caso discreto

Cominciamo con un caso semplice, anche se tipico di condizionamento.

Esempio 3.1.1. Supponiamo che (X, Y) sia una coppia di v.a. discrete, con densità (discreta) congiunta p_{XY} e densità marginali p_X e p_Y . Siano E_X ed E_Y i possibili valori per X e Y rispettivamente, e supponiamo che $p_Y(y) > 0$ per ogni $y \in E_Y$. Per X , supponiamo che abbia speranza matematica finita. Abbiamo visto che la densità discreta di X condizionata a $\{Y = y\}$, è data da

$$p_{X|Y}(x|y) = \frac{p_{XY}(x, y)}{p_Y(y)}, \quad \text{purché } y \in E_Y.$$

La media condizionale di X dato $\{Y = y\}$ è semplicemente la media di X quando per X si consideri la distribuzione condizionale:

$$\begin{aligned} \mathbb{E}(X | Y = y) &= \sum_{x \in E_X} x p_{X|Y}(x|y) = \frac{1}{p_Y(y)} \sum_{x \in E_X} x p_{XY}(x, y) \\ &= \frac{1}{p_Y(y)} \sum_{(x, y) \in E_X \times \{y\}} x p_{XY}(x, y) \\ &= \frac{1}{p_Y(y)} \mathbb{E}(X 1_{\{Y=y\}}). \end{aligned}$$

Se poniamo, per $y \in E_Y$, $h_X(y) = \mathbb{E}(X 1_{\{Y=y\}})/p_Y(y)$, la v.a. $Z = h_X(Y)$ dà la media di X dato il risultato (aleatorio) Y . Ad essa può quindi essere dato il significato di *media condizionale di X dato Y* .

Osserviamo che $Z = h_X(Y)$ verifica le seguenti proprietà.

1. $Z = h_X(Y)$ è $\sigma(Y)$ -misurabile. Infatti è una funzione (misurabile) di Y .
2. Per ogni $B \subset E_Y$, si ha

$$\mathbb{E}(h_X(Y) 1_{\{Y \in B\}}) = \mathbb{E}(X 1_{\{Y \in B\}})$$

Infatti,

$$\mathbb{E}(h_X(Y) 1_{\{Y \in B\}}) = \sum_{y \in B} h_X(y) p_Y(y) = \sum_{y \in B} \mathbb{E}(X 1_{\{Y=y\}}) = \mathbb{E}(X 1_{\{Y \in B\}}).$$

Infine, osserviamo che la v.a. $Z = h_X(Y)$ è l'unica v.a. che soddisfa alle condizioni 1. e 2. Infatti, supponiamo che $\tilde{Z}$ sia un'altra v.a. che soddisfa a 1. e 2. Poiché da 1. $\tilde{Z}$ è $\sigma(Y)$ -misurabile dev'essere (cfr. Proposizione 2.1.16) $\tilde{Z} = \tilde{h}(Y)$, per qualche funzione (misurabile) $\tilde{h}$. Ora, usando 2. con la scelta $B = \{Y = y\} = Y^{-1}(y) = \{Y = y\} \in \sigma(Y)$, si ha

$$\mathbb{E}(X 1_{\{Y=y\}}) = \mathbb{E}(\tilde{Z} 1_{\{Y=y\}}) = \mathbb{E}(\tilde{h}(Y) 1_{\{Y=y\}}) = \tilde{h}(y) \mathbb{E}(1_{\{Y=y\}})$$

e cioè $\tilde{h}(y) = \mathbb{E}(X 1_{\{Y=y\}})/\mathbb{P}(Y = y) = h_X(y)$.

Analizziamo ora un caso un po' più generale.

Esempio 3.1.2. Sia X una v.a. con speranza matematica finita e sia $\mathcal{G}$ una σ -algebra, con $\mathcal{G} \subset \mathcal{F}$. Supponiamo che $\mathcal{G}$ sia generata da una partizione al più numerabile: $\mathcal{G} = \sigma(\mathcal{C})$, con $\mathcal{C} = \{C_i\}_{i \in I}$. Supponiamo anche che $\mathbb{P}(C_i) > 0$ per ogni $i \in I$.

Osserviamo che nell'Esempio 3.1.1 abbiamo implicitamente lavorato con questo tipo di ingredienti. Infatti,

$$\sigma(Y) = \mathcal{G} = \sigma(\mathcal{C}_Y), \quad C \in \mathcal{C}_Y \text{ se e solo se } C = Y^{-1}(\{y\}), \text{ con } y \in E_Y.$$

Ricordiamo inoltre che la media condizionale di X dato Y è

$$Z(\omega) = h_X(Y(\omega)) = \frac{\mathbb{E}(X 1_{\{Y=y\}})}{\mathbb{P}(Y=y)}, \quad \text{quando } \omega \in \{Y=y\}.$$

Ora, gli eventi $\{Y=y\}$, con $y \in E_Y$, sono gli atomi della partizione $\mathcal{C}_Y$ che genera $\mathcal{G} = \sigma(Y)$. Volendo quindi generalizzare ad una qualsiasi σ -algebra $\mathcal{G}$ generata da una partizione $\mathcal{C}$, viene naturale considerare la v.a.

$$Z(\omega) = \sum_{i \in I} \frac{\mathbb{E}(X 1_{C_i})}{\mathbb{P}(C_i)} 1_{\{\omega \in C_i\}} = \frac{\mathbb{E}(X 1_{C_k})}{\mathbb{P}(C_k)} \quad \text{quando } \omega \in C_k$$

Come vedremo in seguito, la v.a. Z è la media condizionale di X data la σ -algebra $\mathcal{G}$. Osserviamo che abbiamo fatto un paio di passi in avanti rispetto all'Esempio 3.1.1: abbiamo una σ -algebra $\mathcal{G}$ (almeno apparentemente) più generale e, soprattutto, non stiamo facendo ipotesi speciali sulla v.a. X : l'importante è che abbia senso la quantità $\mathbb{E}(X 1_{C_k})$, dunque l'importante è che X abbia media.

La v.a. Z verifica le seguenti proprietà, analoghe a quelle già viste nell'Esempio 3.1.1.

1. Z è $\mathcal{G}$ -misurabile. Infatti, 1_{C_i} è $\mathcal{G}$ -misurabile per ogni i , ed essendo Z una combinazione lineare delle funzioni 1_{C_i} , è $\mathcal{G}$ -misurabile.
2. Per ogni $G \in \mathcal{G}$, si ha

$$\mathbb{E}(Z 1_G) = \mathbb{E}(X 1_G).$$

Infatti, preso $G \in \mathcal{G}$, sia $I_G \subset I$ tale che $G = \cup_{i \in I_G} C_i$. Allora, $Z 1_G = \sum_{i \in I_G} \frac{\mathbb{E}(X 1_{C_i})}{\mathbb{P}(C_i)} 1_{C_i}$ e quindi

$$\begin{aligned} \mathbb{E}(Z 1_G) &= \sum_{i \in I_G} \frac{\mathbb{E}(X 1_{C_i})}{\mathbb{P}(C_i)} \mathbb{P}(C_i) = \sum_{i \in I_G} \mathbb{E}(X 1_{C_i}) \\ &= \mathbb{E}(X 1_{\cup_{i \in I_G} C_i}) = \mathbb{E}(X 1_G). \end{aligned}$$

Infine, osserviamo che anche in questo caso la v.a. Z è l'unica v.a. che soddisfa alle condizioni 1. e 2. Infatti, da 1. Z è $\mathcal{G}$ -misurabile, dunque per la Proposizione 2.1.18 è costante sugli elementi della partizione: $Z(\omega) = z_i$ per ogni $\omega \in C_i$. Allora, da 2. si ha

$$\mathbb{E}(X 1_{C_i}) = \mathbb{E}(Z 1_{C_i}) = z_i \mathbb{E}(1_{C_i}) = z_i \mathbb{P}(C_i),$$

cioè $z_i = \mathbb{E}(X 1_{C_i}) / \mathbb{P}(C_i)$, da cui la tesi.

3.1.2 Il caso generale

Negli Esempi 3.1.1 e 3.1.2, abbiamo visto che le proprietà 1. e 2. si rivelano caratteristiche della v.a. Z , che intuitivamente abbiamo chiamato *speranza condizionale* di X data la v.a. Y oppure data la σ -algebra $\mathcal{G}$. La definizione generale prende infatti in considerazione proprio queste due proprietà caratteristiche:

Definizione 3.1.3. Siano $(\Omega, \mathcal{F}, \mathbb{P})$ uno spazio di probabilità, X una v.a. avente speranza matematica finita e $\mathcal{G}$ una sotto σ -algebra di $\mathcal{F}$. Una versione $\mathbb{E}(X | \mathcal{G})$ della *media (o speranza) condizionale di X dato $\mathcal{G}$* è una v.a. tale che:

(i) $\mathbb{E}(X | \mathcal{G})$ è $\mathcal{G}$ -misurabile e integrabile;

(ii) per ogni $G \in \mathcal{G}$,

$$\mathbb{E}(\mathbb{E}(X | \mathcal{G}) 1_G) = \mathbb{E}(X 1_G)$$

Qualora $\mathcal{G} = \sigma(Y)$, con Y v.a., allora $\mathbb{E}(X | \sigma(Y))$ prende il nome di *speranza condizionale di X dato Y* ed è denotata con il simbolo $\mathbb{E}(X | Y)$.

Abbiamo visto negli esempi che la speranza condizionale esiste se la σ -algebra $\mathcal{G}$ è generata da una partizione al più numerabile. Lasciamo a corsi più avanzati la dimostrazione dell'esistenza quando la σ -algebra è più generale. Del resto, in questo corso il caso di σ -algebre generate da una partizione sarà quello che interessa.

Nella Definizione 3.1.3 si parla di *versione* perché è abbastanza chiaro che non c'è unicità. Infatti, se Z è una v.a. $\mathcal{G}$ -misurabile e avente speranza matematica finita e tale che $\mathbb{E}(1_G Z) = \mathbb{E}(1_G X)$ per ogni $G \in \mathcal{G}$ e Z' è un'altra v.a., sempre $\mathcal{G}$ -misurabile, ma che differisce da Z solo su un evento di probabilità 0, allora si ha anche $\mathbb{E}(1_G Z') = \mathbb{E}(1_G X)$ per ogni $G \in \mathcal{G}$.

Quando $\mathcal{G} = \sigma(Y)$, il calcolo della speranza condizionale prende un aspetto particolare. Abbiamo infatti visto (Proposizione 2.1.16) che ogni v.a. $\sigma(Y)$ -misurabile è della forma $h(Y)$, dove h è un'opportuna funzione misurabile. Quindi esiste una funzione misurabile Ψ_X tale che

$$\mathbb{E}(X | Y) = \Psi_X(Y).$$

Il calcolo della speranza condizionale, in questo caso, si riconduce quindi a quello della funzione Ψ_X . Con abuso di linguaggio, denoteremo

$$\Psi_X(y) = \mathbb{E}(X | Y = y),$$

ma attenzione a non confondere la notazione “ $\mathbb{E}(X | Y = y)$ ” con la media condizionale di X dato l’evento $\{Y = y\}$ (che, osserviamo, in molti casi di interesse è un evento di probabilità nulla...). Nell’Esempio 3.1.1, abbiamo visto che

$$\Psi_X(y) = h_X(y) = \frac{\mathbb{E}(X 1_{\{Y=y\}})}{\mathbb{P}(Y = y)}.$$

In generale, la funzione Ψ_X resta determinata dalla relazione

$$\mathbb{E}(X 1_{\{Y \in B\}}) = \mathbb{E}(\Psi_X(Y) 1_{\{Y \in B\}}), \quad \text{per ogni boreliano } B \quad (3.1)$$

(ricordiamo infatti che $\sigma(Y) = \{Y^{-1}(B); B \in \mathcal{B}(E)\}$). Osserviamo che, se X, Y sono v.a. a valori reali, allora siamo in grado di calcolare le due speranze matematiche che compaiono nella (3.1) non appena conosciamo le distribuzioni congiunte di X e Y . Vediamo, nel prossimo esempio, come questo fatto si può sfruttare per calcolare $\mathbb{E}(X | Y)$.

Esempio 3.1.4. Siano X e Y due v.a. reali congiuntamente assolutamente continue, con X avente speranza matematica finita. Indichiamo con f la densità di probabilità congiunta e con f_X e f_Y le due densità marginali. Supponiamo, ma solo per semplicità, che $f_Y(y) > 0$.

La speranza matematica a destra nella (3.1) vale dunque

$$\mathbb{E}(X 1_{\{Y \in B\}}) = \int_B dy \int_{\mathbb{R}} x f(x, y) dx$$

mentre

$$\mathbb{E}(\Psi_X(Y) 1_{\{Y \in B\}}) = \int_B \Psi_X(y) f_Y(y) dy$$

Queste due quantità risultano uguali per ogni scelta del boreliano B se scegliamo

$$\Psi_X(y) = \frac{1}{f_Y(y)} \int_{\mathbb{R}} x f(x, y) dx = \int_{\mathbb{R}} x f_{X|Y}(x | y) dx$$

dove

$$f_{X|Y}(x | y) = \frac{f(x, y)}{f_Y(y)}$$

è la (vecchia!) densità condizionale. Per concludere dobbiamo solo verificare che la v.a. $\Psi_X(Y)$ ha speranza matematica finita. Infatti

$$\int_{\mathbb{R}} |\Psi_X(y)| f_Y(y) dy \leq \int_{\mathbb{R}} dx \int_{\mathbb{R}} |x| f(x, y) dy = \int_{\mathbb{R}} |x| f_X(x) dx < \infty$$

poiché supponiamo che X abbia speranza matematica finita.

In conclusione, la media di X condizionata a Y è la media di X fatta sotto la densità condizionale $f_{X|Y}(x | y)$, sostituendo poi ad y il suo valore generico (aleatorio) Y : esattamente come vorrebbe l’intuizione!

3.1.3 Proprietà della speranza condizionale

Vediamo ora alcune proprietà delle speranze condizionali. Avvertenza: nel seguito l'abbreviazione “q.c.” sta per “quasi certamente”. Significa che le proprietà in questione (uguaglianze, disuguaglianze etc.) sono verificate per $\omega \in \Omega \setminus \mathcal{N}$, dove $\mathcal{N}$ denota un evento tale che $\mathbb{P}(\mathcal{N}) = 0$.

1. Se $X = c$, allora $\mathbb{E}(X | \mathcal{G}) = c$ q.c. Immediato!

2. $\mathbb{E}(\mathbb{E}(X | \mathcal{G})) = \mathbb{E}(X)$.

Infatti basta prendere $G = \Omega \in \mathcal{G}$ nel punto (ii) della Definizione 3.1.3.

3. (Linearità) Se X e Y hanno media allora $\mathbb{E}(\alpha X + \beta Y | \mathcal{G}) = \alpha \mathbb{E}(X | \mathcal{G}) + \beta \mathbb{E}(Y | \mathcal{G})$ q.c., per ogni $\alpha, \beta \in \mathbb{R}$.

Per ipotesi $\mathbb{E}(X | \mathcal{G})$ e $\mathbb{E}(Y | \mathcal{G})$ sono $\mathcal{G}$ -misurabili ed hanno speranza matematica finita. Quindi lo stesso vale per $Z := \alpha \mathbb{E}(X | \mathcal{G}) + \beta \mathbb{E}(Y | \mathcal{G})$. Inoltre, se $G \in \mathcal{G}$,

$$\begin{aligned} \mathbb{E}(Z 1_G) &= \alpha \mathbb{E}(\mathbb{E}(X | \mathcal{G}) 1_G) + \beta \mathbb{E}(\mathbb{E}(Y | \mathcal{G}) 1_G) \\ &= \alpha \mathbb{E}(X 1_G) + \beta \mathbb{E}(Y 1_G) = \mathbb{E}((\alpha X + \beta Y) 1_G), \end{aligned}$$

dunque $Z := \alpha \mathbb{E}(X | \mathcal{G}) + \beta \mathbb{E}(Y | \mathcal{G})$ è una versione della speranza condizionale di $\alpha X + \beta Y$ dato $\mathcal{G}$.

4. (Monotonia) Se $\mathbb{P}(X \geq Y) = 1$ allora $\mathbb{E}(X | \mathcal{G}) \geq \mathbb{E}(Y | \mathcal{G})$ q.c.

Poniamo $G = \{\omega; \mathbb{E}(X | \mathcal{G}) < \mathbb{E}(Y | \mathcal{G})\}$. Evidentemente $G \in \mathcal{G}$. Se fosse $\mathbb{P}(G) > 0$ allora si avrebbe $0 > \mathbb{E}((\mathbb{E}(X | \mathcal{G}) - \mathbb{E}(Y | \mathcal{G})) 1_G)$, mentre invece si ha

$$\mathbb{E}((\mathbb{E}(X | \mathcal{G}) - \mathbb{E}(Y | \mathcal{G})) 1_G) = \mathbb{E}(\mathbb{E}(X - Y | \mathcal{G}) 1_G) = \mathbb{E}((X - Y) 1_G) \geq 0,$$

il che è assurdo.

5. (Disuguaglianza di Jensen) Se Φ è una funzione continua convessa e tale che $\Phi(X)$ abbia speranza matematica finita, allora

$$\mathbb{E}(\Phi(X) | \mathcal{G}) \geq \Phi(\mathbb{E}(X | \mathcal{G})) \quad \text{q.c.}$$

In particolare, se $p \geq 1$ e $\mathbb{E}(|X|^p) < +\infty$, allora possiamo scegliere $\Phi(x) = |x|^p$ e dunque $\mathbb{E}(|X|^p | \mathcal{G}) \geq |\mathbb{E}(X | \mathcal{G})|^p$, q.c.

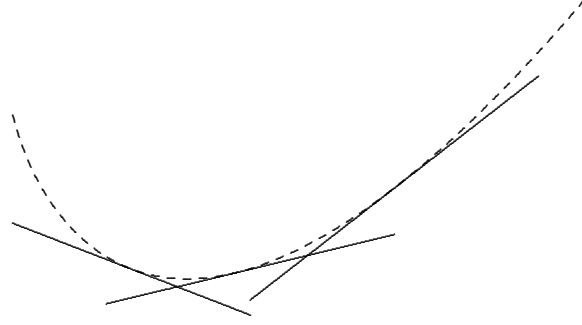


Figura 3.1 Una funzione convessa e continua è sempre involucro superiore delle funzioni lineari affini che la minorano.

Se Φ è convessa allora si può dimostrare che $\Phi(x) = \sup_{f \in A_\Phi} f(x)$, dove A_Φ è l'insieme delle funzioni lineari affini f che minorano Φ : $A_\Phi = \{f; f(x) = ax + b \text{ e } f(x) \leq \Phi(x)\}$ (si veda figura). Allora,

$$\mathbb{E}(\Phi(X) | \mathcal{G}) \geq \mathbb{E}(f(X) | \mathcal{G}) = f(\mathbb{E}(X | \mathcal{G})),$$

da cui segue che

$$\mathbb{E}(\Phi(X) | \mathcal{G}) \geq \sup_{f \in A_\Phi} f(\mathbb{E}(X | \mathcal{G})) = \Phi(\mathbb{E}(X | \mathcal{G})).$$

6. Se Y è $\mathcal{G}$ -misurabile allora $\mathbb{E}(XY | \mathcal{G}) = Y \mathbb{E}(X | \mathcal{G})$, q.c. In particolare, $\mathbb{E}(Y | \mathcal{G}) = Y$, q.c.

Faremo la dimostrazione solo nel caso in cui Y prende un numero finito di valori, $y_1, \dots, y_n$: $Y = \sum_{i \leq n} y_i 1_{\{Y=y_i\}}$. Se poniamo $G_i = \{Y = y_i\}$, allora $G_i \in \mathcal{G}$ perché Y è $\mathcal{G}$ -misurabile. Ora, preso $G \in \mathcal{G}$, si ha:

$$\begin{aligned} \mathbb{E}(XY 1_G) &= \sum_{i \leq n} y_i \mathbb{E}(X 1_{G_i \cap G}) = \sum_{i \leq n} y_i \mathbb{E}(\mathbb{E}(X | \mathcal{G}) 1_{G_i \cap G}) \\ &= \mathbb{E}\left(\sum_{i \leq n} y_i 1_{G_i} \mathbb{E}(X | \mathcal{G}) 1_G\right) = \mathbb{E}(Y \mathbb{E}(X | \mathcal{G}) 1_G). \end{aligned}$$

7. Siano $\mathcal{A}, \mathcal{G} \subset \mathcal{F}$ due σ -algebre tali che $\mathcal{A} \subset \mathcal{G}$. Allora: $\mathbb{E}(X | \mathcal{A}) = \mathbb{E}(\mathbb{E}(X | \mathcal{A}) | \mathcal{G}) = \mathbb{E}(\mathbb{E}(X | \mathcal{G}) | \mathcal{A})$, q.c.

$\mathbb{E}(X | \mathcal{A})$ è $\mathcal{A}$ -misurabile e quindi è $\mathcal{G}$ -misurabile perché $\mathcal{A} \subset \mathcal{G}$. Dunque, per il punto precedente, $\mathbb{E}(X | \mathcal{A}) = \mathbb{E}(\mathbb{E}(X | \mathcal{A}) | \mathcal{G})$. Dimostriamo ora che $\mathbb{E}(X | \mathcal{A}) = \mathbb{E}(\mathbb{E}(X | \mathcal{G}) | \mathcal{A})$: preso $A \in \mathcal{A}$, allora A appartiene anche a $\mathcal{G}$ e

$$\mathbb{E}(\mathbb{E}(X | \mathcal{G}) 1_A) = \mathbb{E}(X 1_A).$$

8. Sia X indipendente da $\mathcal{G}$. Allora, $\mathbb{E}(X | \mathcal{G}) = \mathbb{E}(X)$, q.c.

Se $G \in \mathcal{G}$ allora X e $Y = 1_G$ sono indipendenti, quindi

$$\mathbb{E}(X 1_G) = \mathbb{E}(X)\mathbb{E}(1_G) = \mathbb{E}(\mathbb{E}(X) 1_G).$$

Ora sappiamo che la speranza condizionale $\mathbb{E}(X | \mathcal{G})$ si calcola facilmente se X è indipendente da $\mathcal{G}$ (si veda 8.) oppure se X è già $\mathcal{G}$ -misurabile. C'è un'altra situazione, molto frequente e che combina queste due situazioni, in cui il calcolo è abbastanza facile.

9. Siano X una v.a. $\mathcal{G}$ -misurabile e Y una v.a. indipendente da $\mathcal{G}$. Per ogni f misurabile tale che $f(X, Y)$ ha speranza finita si ha $\mathbb{E}(f(X, Y) | \mathcal{G}) = \bar{f}(X)$ q.c., dove $\bar{f}(x) = \mathbb{E}(f(x, Y))$.

Lo dimostreremo solo nel caso in cui la funzione f è della forma $f(x, y) = f_1(x)f_2(y)$. Il caso generale si dimostra usando il fatto che ogni funzione $f(x, y)$ si può approssimare con delle combinazioni lineari di funzioni del tipo $f_1(x)f_2(y)$ e ne lasceremo la dimostrazione ad un corso più avanzato. Se f è della forma indicata, allora $\bar{f}(x) = f_1(x)\mathbb{E}(f_2(Y))$. Usando 6. e 8.,

$$\begin{aligned}\mathbb{E}(f(X, Y) | \mathcal{G}) &= \mathbb{E}(f_1(X) f_2(Y) | \mathcal{G}) = f_1(X) \mathbb{E}(f_2(Y) | \mathcal{G}) = \\ &= f_1(X) \mathbb{E}(f_2(Y)) = \bar{f}(X).\end{aligned}$$

10. Si ha, per ogni v.a. Z che sia $\mathcal{G}$ -misurabile e di quadrato integrabile,

$$\mathbb{E}[(X - \mathbb{E}(X | \mathcal{G}))^2] \leq \mathbb{E}[(X - Z)^2] \quad (3.2)$$

e, per di più la disuguaglianza precedente è stretta a meno che non sia $Z = \mathbb{E}(X | \mathcal{G})$ q.c. In altre parole $\mathbb{E}(X | \mathcal{G})$ è la v.a. $\mathcal{G}$ -misurabile che meglio approssima X nel senso dei minimi quadrati.

È questa una delle proprietà più importanti della speranza condizionale. Si noti l'analogia con $\mathbb{E}(X)$: sia $\mathbb{E}(X)$ che $\mathbb{E}(X | \mathcal{G})$ sono le migliori stime di X nel senso dei minimi quadrati, nel primo caso nella classe delle costanti mentre nel secondo caso nello spazio delle v.a. con momento secondo finito e $\mathcal{G}$ -misurabili.

Dimostriamo la (3.2). Sia X una v.a. con momento secondo finito. Intanto, osserviamo che grazie alla disuguaglianza di Jensen, sappiamo che anche $\mathbb{E}(X | \mathcal{G})$ ha momento secondo finito (si scelga $p = 2$ nel punto 5.). Ora, presa una qualunque v.a. Z che sia $\mathcal{G}$ -misurabile e di quadrato integrabile, si ha

$$\begin{aligned}\mathbb{E}[(X - Z)^2] &= \mathbb{E}[(X - \mathbb{E}(X | \mathcal{G}))^2] + \mathbb{E}[(Z - \mathbb{E}(X | \mathcal{G}))^2] \\ &\quad - 2\mathbb{E}[(X - \mathbb{E}(X | \mathcal{G}))(Z - \mathbb{E}(X | \mathcal{G}))].\end{aligned}$$

Ma $Z - \mathbb{E}(X | \mathcal{G})$ è una v.a. $\mathcal{G}$ -misurabile e dunque

$$\begin{aligned} & \mathbb{E}[(X - \mathbb{E}(X | \mathcal{G}))(Z - \mathbb{E}(X | \mathcal{G}))] \\ &= \mathbb{E}[\mathbb{E}[(X - \mathbb{E}(X | \mathcal{G}))(Z - \mathbb{E}(X | \mathcal{G})) | \mathcal{G}]] \\ &= (Z - \mathbb{E}(X | \mathcal{G}))\mathbb{E}[\mathbb{E}[X - \mathbb{E}(X | \mathcal{G}) | \mathcal{G}]] = 0. \end{aligned}$$

Dunque,

$$\mathbb{E}[(X - Z)^2] \geq \mathbb{E}[(X - \mathbb{E}(X | \mathcal{G}))^2]$$

e si ha uguaglianza se e solo se $Z = \mathbb{E}(X | \mathcal{G})$, da cui la tesi.

3.1.4 Probabilità condizionale

Definizione 3.1.5. Sia $A \in \mathcal{F}$ e sia $\mathcal{G}$ una sotto- σ -algebra di $\mathcal{F}$. Una versione della *probabilità condizionale* di A dato $\mathcal{G}$ è data da

$$\mathbb{P}(A | \mathcal{G}) = \mathbb{E}(1_A | \mathcal{G}).$$

Così come nel caso della media condizionale, quando $\mathcal{G} = \sigma(Y)$ useremo la notazione $\mathbb{P}(A | Y)$.

Ricordiamo che, anche in questo caso, occorre parlare di “versione”: analogamente alla media condizionale, la probabilità condizionale è definita a meno di insiemi di probabilità nulla, ovvero è definita q.c.

Esempio 3.1.6. Riprendendo gli esempi 3.1.1, 3.1.2 e 3.1.4, possiamo scrivere:

- se (X, Y) è una coppia di v.a. discrete, di densità discreta congiunta p_{XY} e densità (discrete) marginali p_X e p_Y allora q.c.

$$\begin{aligned} \mathbb{P}(X \in B | Y) &= h_B(Y), \quad \text{dove } h_B(y) = \mathbb{P}(X \in B | Y = y) \\ &= \sum_{x \in B} \frac{p_{XY}(x, y)}{p_Y(y)} \end{aligned}$$

per ogni $B \subset E_X$ (abbiamo scelto in particolare $A = \{X \in B\}$);

- se $\mathcal{G} = \sigma(\mathcal{C})$, con $\mathcal{C} = \{C_i\}_{i \in I}$ partizione al più numerabile di Ω , e se $A \in \mathcal{F}$ allora q.c.

$$\mathbb{P}(A | \mathcal{G}) = \sum_i \frac{\mathbb{P}(A \cap C_i)}{\mathbb{P}(C_i)} 1_{C_i};$$

- se (X, Y) è una coppia di v.a. assolutamente continue, di densità continua congiunta p_{XY} e densità (continue) marginali p_X e p_Y allora q.c.

$$\mathbb{P}(X \in B | Y) = \Psi_B(Y), \quad \text{dove } \Psi_B(y) = \int_B \frac{p_{XY}(x, y)}{p_Y(y)} dy,$$

per ogni $B \in \mathcal{E}$ (abbiamo scelto in particolare $A = \{X \in B\}$).

La probabilità condizionale verifica le due proprietà basilari della probabilità:

- (i) per ogni $A \in \mathcal{F}$, $\mathbb{P}(A^c | \mathcal{G}) = 1 - \mathbb{P}(A | \mathcal{G})$, q.c.: immediata conseguenza del fatto che $1 = 1_A + 1_{A^c}$.
- (ii) per ogni $A_1, A_2, \dots \in \mathcal{F}$ tali che $A_n \cap A_m = \emptyset$ per $n \neq m$ allora $\mathbb{P}(\cup_n A_n | \mathcal{G}) = \sum_n \mathbb{P}(A_n | \mathcal{G})$, q.c.

Lo dimostriamo solo nel caso in cui $\mathcal{G}$ è generata da una partizione $\mathcal{C} = \{C_i\}_{i \in I}$ al più numerabile. In tal caso, ponendo $I_{\mathcal{N}} = \{i \in I : \mathbb{P}(C_i) = 0\}$, allora ovviamente $\mathbb{P}(\cup_{i \in I_{\mathcal{N}}} C_i) = 0$ ed è quindi immediato vedere che

$$\mathbb{E}(X | \mathcal{G}) = \sum_{i \notin I_{\mathcal{N}}} \frac{\mathbb{E}(X 1_{C_i})}{\mathbb{P}(C_i)} 1_{C_i}, \quad \text{q.c.}$$

Ma allora, q.c. si hanno le seguenti uguaglianze:

$$\begin{aligned} \mathbb{P}(\cup_n A_n | \mathcal{G})(\omega) &= \sum_{i \notin I_{\mathcal{N}}} \frac{\mathbb{P}((\cup_n A_n) \cap C_i)}{\mathbb{P}(C_i)} 1_{C_i} = \sum_{i \notin I_{\mathcal{N}}} \sum_n \frac{\mathbb{P}(A_n \cap C_i)}{\mathbb{P}(C_i)} 1_{C_i} \\ &= \sum_n \sum_{i \notin I_{\mathcal{N}}} \frac{\mathbb{P}(A_n \cap C_i)}{\mathbb{P}(C_i)} 1_{C_i} = \sum_n \mathbb{P}(A_n | \mathcal{G}), \end{aligned}$$

da cui la tesi.

Si potrebbe allora pensare di poter costruire una v.a.

$$\mathbb{P}(\cdot, \cdot) : \mathcal{F} \times \Omega \longrightarrow \mathbb{R}$$

tale che

1. fissato $A \in \mathcal{F}$, $\omega \mapsto \mathbb{P}(A, \omega)$ è una versione di $\mathbb{P}(A | \mathcal{G})(\omega)$;
2. fissato $\omega \in \Omega$, $A \mapsto \mathbb{P}(A, \omega)$ è una probabilità su $(\Omega, \mathcal{F})$.

Qualora sia possibile, $\mathbb{P}(A, \omega)$ prende il nome di *versione regolare della probabilità condizionata*. In generale però, tale funzione $\mathbb{P}(A, \omega)$ non si può costruire. Infatti, abbiamo visto che, fissato $A \in \mathcal{F}$, $\mathbb{P}(A | \mathcal{G})$ è definita q.c.: esiste $\mathcal{N}_A \in \mathcal{F}$ tale che $\mathbb{P}(\mathcal{N}_A) = 0$ e $\mathbb{P}(A | \mathcal{F})(\omega)$ è ben definita per $\omega \notin \mathcal{N}_A$. Ne segue che, in generale,

$$\omega \mapsto \mathbb{P}(A | \mathcal{F})(\omega) \text{ è ben definita per } \omega \notin \mathcal{N} = \cup_{A \in \mathcal{F}} \mathcal{N}_A.$$

Ora, perché il tutto abbia senso occorre che $\mathcal{N}$ sia “trascurabile”, cioè $\mathbb{P}(\mathcal{N}) = 0$. Ma, poiché $\mathcal{F}$ non è in generale numerabile e dunque $\mathcal{N} = \cup_{A \in \mathcal{F}} \mathcal{N}_A$ non è in generale un’unione numerabile di eventi,

- non è detto che $\mathcal{N}$ sia un evento, ovvero $\mathcal{N} \in \mathcal{F}$;

- se anche $\mathcal{N} \in \mathcal{F}$, non è detto che $\mathbb{P}(\mathcal{N}) = 0$.

Osserviamo infine che, poiché la probabilità condizionale di A è data dalla media condizionale della v.a. $X = 1_A$, molte delle proprietà viste per le medie condizionali valgono anche in questo caso. Ad esempio, se $A \in \mathcal{G}$ allora ovviamente $\mathbb{P}(A | \mathcal{G}) = 1_A$ q.c., così come se A è indipendente da $\mathcal{G}$ allora $\mathbb{P}(A | \mathcal{G}) = \mathbb{P}(A)$, q.c.

3.2 Martingale

3.2.1 Definizioni e generalità

Sia $(\Omega, \mathcal{F}, \mathbb{P})$ uno spazio di probabilità. Una *filtrazione* è una successione $(\mathcal{F}_n)_n$ *crescente* di sotto- σ -algebre di $\mathcal{F}$, cioè tale che $\mathcal{F} \supset \mathcal{F}_{n+1} \supset \mathcal{F}_n$. Presa una successione $(X_n)_n$, diremo che è *adattata* alla filtrazione $(\mathcal{F}_n)_n$, o anche che è $\mathcal{F}_n$ -adattata, se X_n è $\mathcal{F}_n$ -misurabile, per ogni n . In questo paragrafo ci collocheremo sempre in questo contesto.

Definizione 3.2.1. Una successione $(X_n)_n$ di v.a. è una martingala (risp. una supermartingala, una submartingala) rispetto alla filtrazione $(\mathcal{F}_n)_n$ se $(X_n)_n$ è $\mathcal{F}_n$ -adattata, la v.a. X_n ha speranza matematica finita per ogni n e se, per ogni n ,

$$\mathbb{E}(X_{n+1} | \mathcal{F}_n) = X_n \quad \text{q.c. (risp. } \leq, \geq \text{)}.$$

Esercizio 3.1. Verificare che $(X_n)_n$ è una martingala (risp. una supermartingala, una submartingala) se e solo se, per ogni $A \in \mathcal{F}_n$,

$$\mathbb{E}(1_A X_{n+1}) = \mathbb{E}(1_A X_n) \quad (\text{risp. } \leq, \geq). \quad (3.3)$$

Si noti che $(X_n)_n$ è una supermartingala se e solo se $(-X_n)_n$ è una submartingala e $(X_n)_n$ è una martingala se e solo se $(X_n)_n$ è simultaneamente una supermartingala e una submartingala.

Esempio 3.2.2. Sia $(Y_n)_n$ una successione di v.a. indipendenti e sia $\mathcal{F}_n = \sigma(Y_1, \dots, Y_n)$. Supponiamo inoltre che le Y_n abbiano tutte speranza matematica finita ed $= 0$ (risp. $\leq 0, \geq 0$). Allora: $X_n = Y_0 + Y_1 + \dots + Y_n$ è una $\mathcal{F}_n$ -martingala (risp. una supermartingala, una submartingala).

Infatti

$$\begin{aligned} \mathbb{E}(X_{n+1} | \mathcal{F}_n) &= \mathbb{E}(X_n + Y_{n+1} | \mathcal{F}_n) = \mathbb{E}(X_n | \mathcal{F}_n) + \mathbb{E}(Y_{n+1} | \mathcal{F}_n) \\ &= X_n + \mathbb{E}(Y_{n+1}) = (\text{risp. } \leq 0, \geq 0) X_n \end{aligned}$$

poiché X_n è $\mathcal{F}_n$ -misurabile, mentre Y_n è indipendente da $\mathcal{F}_n$.

Esercizio 3.2. Sia X una v.a. avente speranza matematica finita e $\{\mathcal{F}_n\}_n$ una filtrazione. Mostrare che $X_n = \mathbb{E}(X | \mathcal{F}_n)$ è una martingala.

Esercizio 3.3. Siano $Y_0, Y_1, \dots$ v.a. i.i.d., tali che $\mathbb{P}(Y_n = 1) = p = 1 - q = \mathbb{P}(Y_n = -1)$. Siano $\mathcal{F}_n = \sigma(Y_0, \dots, Y_n)$ e $X_n = Y_0 + \dots + Y_n$. Poniamo

$$Z_n = \left(\frac{q}{p}\right)^{X_n} \quad e \quad W_n = \left(\frac{p}{q}\right)^{X_n}.$$

Dire se la proprietà di martingala o supermartingala o submartingala rispetto a $\mathcal{F}_n$ è verificata da $(Z_n)_n$ e/o $(W_n)_n$.

Esercizio 3.4. Siano $Y_0, Y_1, \dots$ v.a. i.i.d., con media μ . Poniamo $\mathcal{F}_n = \sigma(Y_0, \dots, Y_n)$ e $X_n = Y_0 \cdots Y_n$. Dimostrare che $(X_n)_n$ è una $\mathcal{F}_n$ -martingala se e solo se $\mu = 1$. Supponiamo quindi $\mu = 1$ e che le $(Y_n)_n$ abbiano varianza σ^2 : verificare che $Z_n = X_n^2 - \sigma^2 \sum_{k=0}^{n-1} X_k^2$ è una $\mathcal{F}_n$ -martingala.

3.2.2 Prime proprietà

- Se $(X_n)_n$ è una martingala (risp. una supermartingala, una submartingala), la successione $(\mathbb{E}(X_n))_n$ è costante (risp. decrescente, crescente).

Basta infatti applicare la (3.3) con $A = \Omega$.

- Se $(X_n)_n$ è una martingala (risp. una supermartingala, una submartingala), per $m < n$ si ha $\mathbb{E}(X_n | \mathcal{F}_m) = X_m$ q.c. (risp. $\leq, \geq$).

Infatti, per ricorrenza,

$$\begin{aligned} \mathbb{E}[X_n | \mathcal{F}_m] &= \mathbb{E}[\mathbb{E}(X_n | \mathcal{F}_{n-1}) | \mathcal{F}_m] \\ &= \mathbb{E}[X_{n-1} | \mathcal{F}_m] = \dots = \mathbb{E}[X_{m+1} | \mathcal{F}_m] = X_m. \end{aligned}$$

- Se $(X_n)_n$ e $(Y_n)_n$ sono martingale (o supermartingale o submartingale), lo stesso vale per $X_n + Y_n$.

La dimostrazione è immediata.

- Se $(X_n)_n$ e $(Y_n)_n$ sono supermartingale anche $X_n \wedge Y_n$ è una supermartingala.

Infatti

$$\begin{aligned} \mathbb{E}(X_{n+1} \wedge Y_{n+1} | \mathcal{F}_n) &\leq \mathbb{E}(X_{n+1} | \mathcal{F}_n) \leq X_n \\ \mathbb{E}(X_{n+1} \wedge Y_{n+1} | \mathcal{F}_n) &\leq \mathbb{E}(Y_{n+1} | \mathcal{F}_n) \leq Y_n \end{aligned}$$

e dunque $\mathbb{E}(X_{n+1} \wedge Y_{n+1} | \mathcal{F}_n) \leq X_n \wedge Y_n$. Analogamente, si prova la seguente proprietà:

- Se $(X_n)_n$ e $(Y_n)_n$ sono submartingale anche $X_n \vee Y_n$ è una submartingala.
- Se $(X_n)_n$ è una martingala (risp. una submartingala) e $f : \mathbb{R} \rightarrow \mathbb{R}$ è una funzione convessa (risp. convessa e crescente) tale che $f(X_n)$ è integrabile per ogni $n \geq 0$, allora $Y_n = f(X_n)$ è una submartingala.

La dimostrazione segue direttamente applicando la disuguaglianza di Jensen (per la speranza condizionale).

3.2.3 La decomposizione di Doob

Si dice che $(A_n)_n$ è un *processo prevedibile crescente* se $A_0 = 0$, $A_n \leq A_{n+1}$ per ogni $n \geq 0$ e A_{n+1} è $\mathcal{F}_n$ -misurabile.

Sia $(X_n)_n$ una $\mathcal{F}_n$ -submartingala. Definiamo

$$A_0 = 0, \quad A_{n+1} = A_n + \mathbb{E}(X_{n+1} - X_n | \mathcal{F}_n) = \sum_{k=0}^n \mathbb{E}(X_{k+1} - X_k | \mathcal{F}_k).$$

Per costruzione, $(A_n)_n$ è un processo prevedibile crescente e $M_n = X_n - A_n$ verifica $\mathbb{E}(M_{n+1} - M_n | \mathcal{F}_n) = 0$ (perché A_{n+1} è $\mathcal{F}_n$ -misurabile!). Quindi, $(M_n)_n$ è una $\mathcal{F}_n$ -martingala ed inoltre

$$X_n = M_n + A_n.$$

Questa decomposizione è, per di più, unica. Infatti, se $X_n = M'_n + A'_n$ è un'altra decomposizione di questo tipo, allora $A'_0 = 0$ e

$$A'_{n+1} - A'_n = X_{n+1} - X_n - (M'_{n+1} - M'_n)$$

da cui, condizionando rispetto a $\mathcal{F}_n$, $A'_{n+1} - A'_n = \mathbb{E}(X_{n+1} - X_n | \mathcal{F}_n)$. Dunque, $A'_n = A_n$ e $M'_n = M_n$. Abbiamo quindi dimostrato che

Teorema 3.2.3. (*Decomposizione di Doob*) Ogni submartingala $(X_n)_n$ si può scrivere, in maniera unica, nella forma $X_n = M_n + A_n$, dove $(M_n)_n$ è una martingala e $(A_n)_n$ è un processo prevedibile crescente integrabile e tale che $A_0 = 0$.

Il processo prevedibile crescente $(A_n)_n$ del Teorema (3.2.3) si chiama il *compensatore* della submartingala $(X_n)_n$.

Presa una supermartingala (risp. martingala) $(X_n)_n$, allora $(-X_n)_n$ è una submartingala (risp. martingala). Applicando a $(-X_n)_n$ il Teorema 3.2.3, otteniamo la seguente decomposizione di Doob per supermartingale:

Teorema 3.2.4. Ogni supermartingala $(X_n)_n$ si può scrivere, in maniera unica, nella forma $X_n = M_n - A_n$, dove $(M_n)_n$ è una martingala e $(A_n)_n$ è un processo prevedibile crescente integrabile e tale che $A_0 = 0$.

Esercizio 3.5. Sia $(X_n)_n$ la martingala dell'Esempio 3.2.2: $X_n = Y_0 + \dots + Y_n$, con $Y_1, \dots$ i.i.d., di media 0, e $\mathcal{F}_n = \sigma(Y_0, \dots, Y_n)$. Supponiamo che le Y_i abbiano varianza σ^2 . Verificare che X_n^2 è una submartingala e determinarne il compensatore.

3.2.4 Martingale trasformate

Proposizione 3.2.5. Sia $(M_n)_n$ una $\mathcal{F}_n$ -martingala e $(H_n)_n$ un processo prevedibile, cioè tale che H_{n+1} sia $\mathcal{F}_n$ misurabile. Sia

$$X_0 = 0 \quad \text{e per } n > 0, \quad X_n = \sum_{j=1}^n H_j (M_j - M_{j-1})$$

Allora $(X_n)_n$ è ancora una martingala rispetto alla stessa filtrazione $\mathcal{F}_n$.

La martingala $(X_n)_n$ definita nella Proposizione 3.2.5 è detta la *martingala trasformata di $(M_n)_n$ tramite $(H_n)_n$* .

Dimostrazione della Proposizione 3.2.5. X_n è chiaramente $\mathcal{F}_n$ -misurabile e poiché H_{n+1} è $\mathcal{F}_n$ -misurabile, si ha

$$\mathbb{E}(H_{n+1}(M_{n+1} - M_n) \mid \mathcal{F}_n) = H_{n+1} \mathbb{E}(M_{n+1} - M_n \mid \mathcal{F}_n) = 0.$$

Essendo $X_{n+1} = H_{n+1}(M_{n+1} - M_n) + X_n$, si ha

$$\begin{aligned} \mathbb{E}(X_{n+1} \mid \mathcal{F}_n) &= \mathbb{E}(H_{n+1}(M_{n+1} - M_n) \mid \mathcal{F}_n) + \mathbb{E}(X_n \mid \mathcal{F}_n) \\ &= H_{n+1} \mathbb{E}(M_{n+1} - M_n \mid \mathcal{F}_n) + X_n = X_n \end{aligned}$$

□

Presentiamo infine un'utile caratterizzazione delle martingale, che fa uso delle martingale trasformate.

Proposizione 3.2.6. Sia $(M_n)_n$ un processo $\mathcal{F}_n$ -adattato e tale che M_n è integrabile per ogni n . $(M_n)_{0 \leq n \leq N}$ è una martingala se e solo se per ogni processo prevedibile $(H_n)_{0 \leq n \leq N}$ si ha

$$\mathbb{E}\left(\sum_{n=1}^N H_n(M_n - M_{n-1})\right) = 0.$$

Dimostrazione. Se $(M_n)_{0 \leq n \leq N}$ è una martingala, dalla Proposizione 3.2.5 anche $(X_n)_{0 \leq n \leq N}$ è una martingala, con $X_n = \sum_{j=1}^n H_j(M_j - M_{j-1})$ per $n \geq 1$ e $X_0 = 0$. Quindi,

$$\mathbb{E}\left(\sum_{n=1}^N H_n(M_n - M_{n-1})\right) = \mathbb{E}(X_N) = \mathbb{E}(X_0) = 0.$$

Viceversa, per dimostrare che $(M_n)_{0 \leq n \leq N}$ è una martingala basta far vedere che $\mathbb{E}(M_{n+1} 1_A) = \mathbb{E}(M_n 1_A)$, per ogni fissato $n = 0, 1, \dots, N-1$ e $A \in \mathcal{F}_n$. Fissiamo quindi $n = 0, 1, \dots, N-1$ e $A \in \mathcal{F}_n$. Sia

$$H_j = \begin{cases} 1_A & \text{se } j = n+1 \\ 0 & \text{altrimenti.} \end{cases}$$

H è ovviamente prevedibile e

$$0 = \mathbb{E}\left(\sum_{j=1}^N H_j(M_j - M_{j-1})\right) = \mathbb{E}\left(1_A(M_{n+1} - M_n)\right)$$

cioè $\mathbb{E}(M_{n+1} 1_A) = \mathbb{E}(M_n 1_A)$.

□

Vedremo in seguito che il processo-martingala con cui lavoreremo è dato dai prezzi (scontati) dei titoli presenti nel mercato finanziario. In particolare, quindi, si tratterà di un processo multidimensionale. Le martingale a valori in $\mathbb{R}^n$ si definiscono nel modo più naturale: il processo $\mathcal{F}_n$ -adattato e integrabile $(M_n)_n$, con $M_n = (M_n^1, \dots, M_n^d)$ a valori in $\mathbb{R}^d$, è una martingala se e solo se lo sono i processi $(M_n^i)_n$ per ogni $i = 1, \dots, d$. Ricordando che la media di un vettore aleatorio è il vettore delle medie delle coordinate (aleatorie), si ha semplicemente che

$$\mathbb{E}(M_{n+1} | \mathcal{F}_n) = M_n, \quad \text{per ogni } n.$$

Ovviamente, tutte le proprietà di martingala che abbiamo visto continuano ad essere valide. In particolare, la Proposizione 3.2.6 si può riscrivere nel modo seguente:

Proposizione 3.2.7. *Sia $(M_n)_n$ un processo a valori in $\mathbb{R}^d$, $\mathcal{F}_n$ -adattato e tale che M_n è integrabile per ogni n . $(M_n)_{0 \leq n \leq N}$ è una martingala se e solo se comunque si scelgano d processi prevedibili $(H_n^1)_{0 \leq n \leq N}, \dots, (H_n^d)_{0 \leq n \leq N}$ si ha*

$$\mathbb{E}\left(\sum_{i=1}^d \sum_{n=1}^N H_n^i (M_n^i - M_{n-1}^i)\right) = 0.$$

La dimostrazione della Proposizione 3.2.7 è lasciata per esercizio.

3.3 Tempi d'arresto

Definizione 3.3.1. (i) Un tempo d'arresto rispetto alla filtrazione $(\mathcal{F}_n)_n$, anche detto un $\mathcal{F}_n$ -tempo d'arresto, è una v.a. $\tau : \Omega \rightarrow \bar{\mathbb{N}}$ (quindi può anche prendere il valore $+\infty$) tale che, per ogni $n \geq 0$,

$$\{\tau \leq n\} \in \mathcal{F}_n.$$

(ii) Se τ è un $\mathcal{F}_n$ -tempo d'arresto, si chiama σ -algebra degli eventi antecedenti al tempo τ , in simboli $\mathcal{F}_\tau$, la σ -algebra

$$\mathcal{F}_\tau = \{A \in \mathcal{F}_\infty; \text{ per ogni } n \geq 0, A \cap \{\tau \leq n\} \in \mathcal{F}_n\},$$

dove $\mathcal{F}_\infty$ denota la σ -algebra generata da $\bigcup_n \mathcal{F}_n$.

Esercizio 3.6. Verificare che $\mathcal{F}_\tau$ è effettivamente una σ -algebra.

È utile osservare che nelle (i) e (ii) della Definizione 3.3.1, si possono sostituire le condizioni $\{\tau \leq n\} \in \mathcal{F}_n$ e $A \cap \{\tau \leq n\} \in \mathcal{F}_n$ con $\{\tau = n\} \in \mathcal{F}_n$ e $A \cap \{\tau = n\} \in \mathcal{F}_n$, rispettivamente, perché

$$\{\tau \leq n\} = \bigcup_{k=0}^n \{\tau = k\}, \quad \text{quindi } \{\tau = n\} = \{\tau \leq n\} \cap \{\tau \leq n-1\}^c.$$

Esempio 3.3.2. (*Esempio fondamentale*) Sia $(X_n)_n$ un processo a valori in $(E, \mathcal{E})$ e adattato alla filtrazione $(\mathcal{F}_n)_n$. Poniamo, per $A \in \mathcal{E}$,

$$\tau_A(\omega) = \inf\{n \geq 0; X_n(\omega) \in A\}, \quad (3.4)$$

con l'intesa che $\inf \emptyset = +\infty$. Allora τ_A è un tempo d'arresto, poichè

$$\{\tau_A = n\} = \{X_0 \notin A, X_1 \notin A, \dots, X_{n-1} \notin A, X_n \in A\} \in \mathcal{F}_n;$$

τ_A si chiama *tempo di passaggio o d'ingresso* in A .

Siano τ_1, τ_2 due tempi d'arresto della filtrazione $(\mathcal{F}_n)_n$. Vediamo ora alcune proprietà interessanti, di cui ci serviremo nel seguito. La loro verifica può apparire complicata a prima vista (soprattutto le σ -algebre $\mathcal{F}_\tau$ fanno un po' di paura) ma si rivela poi abbastanza semplice e scolastica. Lasciamo per esercizio la dimostrazione.

Esercizio 3.7. *Siano τ_1 e τ_2 due $\mathcal{F}_n$ -tempi d'arresto. Dimostrare che $\tau_1 + \tau_2$, $\tau_1 \vee \tau_2$, $\tau_1 \wedge \tau_2$ sono tempi d'arresto, rispetto alla stessa filtrazione.*

Esercizio 3.8. *Siano τ_1 e τ_2 due $\mathcal{F}_n$ -tempi d'arresto. Dimostrare che se $\tau_1 \leq \tau_2$, allora $\mathcal{F}_{\tau_1} \subset \mathcal{F}_{\tau_2}$.*

Esercizio 3.9. *Siano τ_1 e τ_2 due $\mathcal{F}_n$ -tempi d'arresto. Mostrare che $\mathcal{F}_{\tau_1 \wedge \tau_2} = \mathcal{F}_{\tau_1} \cap \mathcal{F}_{\tau_2}$.*

Esercizio 3.10. *Siano τ_1 e τ_2 due $\mathcal{F}_n$ -tempi d'arresto. Dimostrare che $\{\tau_1 < \tau_2\}$ e $\{\tau_1 = \tau_2\}$ appartengono entrambi a $\mathcal{F}_{\tau_1} \cap \mathcal{F}_{\tau_2}$.*

Sia $(X_n)_n$ un processo adattato; spesso avremo bisogno di considerare la posizione del processo al tempo τ , cioè $X_\tau(\omega) = X_{\tau(\omega)}(\omega)$. Possiamo quindi dire che

$$X_\tau = X_n \text{ sull'evento } \{\tau = n\}, \quad n \in \bar{\mathbb{N}}.$$

Osserviamo che la v.a. X_τ è $\mathcal{F}_\tau$ -misurabile perché

$$\{X_\tau \in A\} \cap \{\tau = n\} = \{X_n \in A\} \cap \{\tau = n\} \in \mathcal{F}_n.$$

Il risultato che segue è molto utile per calcolare la speranza condizionale rispetto alla σ -algebra $\mathcal{F}_\tau$.

Proposizione 3.3.3. *Sia X una v.a. integrabile e sia τ un tempo d'arresto. Si ponga, per ogni $n \in \bar{\mathbb{N}}$, $X_n = \mathbb{E}(X \mid \mathcal{F}_n)$. Allora sull'evento $\{\tau = n\}$, si ha $\mathbb{E}(X \mid \mathcal{F}_\tau) = \mathbb{E}(X \mid \mathcal{F}_n)$, cioè*

$$\mathbb{E}(X \mid \mathcal{F}_\tau) = X_\tau \quad (3.5)$$

Dimostrazione. Se $A \in \mathcal{F}_\tau$, allora $1_A 1_{\{\tau=n\}} = 1_{A \cap \{\tau=n\}}$ è $\mathcal{F}_n$ -misurabile. Quindi,

$$\mathbb{E}(1_A X_\tau) = \sum_{n \in \mathbb{N}} \mathbb{E}[1_A X_n 1_{\{\tau=n\}}] = \sum_{n \in \mathbb{N}} \mathbb{E}[1_A X 1_{\{\tau=n\}}] = \mathbb{E}(1_A X),$$

e dunque $\mathbb{E}(X | \mathcal{F}_\tau) = X_\tau$. □

Concludiamo con il prossimo risultato, molto semplice da dimostrare ma, come vedremo in seguito, particolarmente utile.

Proposizione 3.3.4. *Se $(X_n)_n$ è una martingala (resp. una supermartingala, una submartingala), lo stesso è vero anche per il processo arrestato*

$$X_n^\tau = X_{n \wedge \tau},$$

dove τ è una tempo d'arresto della filtrazione $(\mathcal{F}_n)_n$.

Dimostrazione. Per la definizione di tempo d'arresto, $\{\tau \geq n+1\} = \{\tau \leq n\}^c \in \mathcal{F}_n$; dunque

$$\begin{aligned} \mathbb{E}(X_{n+1}^\tau - X_n^\tau | \mathcal{F}_n) &= \mathbb{E}((X_{n+1} - X_n) 1_{\{\tau \geq n+1\}} | \mathcal{F}_n) \\ &= 1_{\{\tau \geq n+1\}} \mathbb{E}(X_{n+1} - X_n | \mathcal{F}_n) = (\text{risp. } \leq, \geq) 0. \end{aligned} \tag{3.6}$$

□

3.3.1 Il teorema d'arresto

Sia $(X_n)_n$ una supermartingala. Quindi, per $m < n$, si ha $\mathbb{E}(X_n | \mathcal{F}_m) \leq X_m$ q.c. Questa proprietà resta vera se si sostituiscono m e n con dei tempi d'arresto? Cioè: se τ_1 e τ_2 sono due tempi d'arresto con $\tau_1 \leq \tau_2$, è vero che

$$\mathbb{E}(X_{\tau_2} | \mathcal{F}_{\tau_1}) \leq X_{\tau_1} \tag{3.7}$$

Supponiamo che τ_1 e τ_2 siano limitati: esiste $k \in \mathbb{N}$ tale che $\tau_1 \leq \tau_2 \leq k$. Poiché $A \cap \{\tau_1 = j\} \in \mathcal{F}_j$ e $(X_n^{\tau_2})_n = (X_{\tau_2 \wedge n})_n$ è una supermartingala,

$$\mathbb{E}(X_{\tau_2} 1_A) = \sum_{j=0}^k \mathbb{E}(X_{\tau_2 \wedge k} 1_{A \cap \{\tau_1=j\}}) \leq \sum_{j=0}^k \mathbb{E}(X_{\tau_2 \wedge j} 1_{A \cap \{\tau_1=j\}}).$$

Ma $X_{\tau_2 \wedge j} = j$ su $\{\tau_1 = j\}$, dunque

$$\mathbb{E}(X_{\tau_2} 1_A) \leq \sum_{j=0}^k \mathbb{E}(X_j 1_{A \cap \{\tau_1=j\}}) = \mathbb{E}(X_{\tau_1} 1_A).$$

Quindi abbiamo dimostrato che

Teorema 3.3.5. *(Teorema d'arresto) Siano $(X_n)_n$ una martingala (resp. supermartingala, submartingala) e τ_1, τ_2 due tempi d'arresto limitati tali che $\tau_1 \leq \tau_2$. Allora $\mathbb{E}(X_{\tau_2} | \mathcal{F}_{\tau_1}) = X_{\tau_1}$ (resp. $\leq, \geq$).*

Prendendo le speranze matematiche,

Corollario 3.3.6. *Sia $(X_n)_n$ una martingala (resp. supermartingala, submartingala) e sia τ un tempo d'arresto limitato. Allora $\mathbb{E}(X_\tau) = \mathbb{E}(X_0)$ (resp. $\leq, \geq$).*

3.4 Soluzioni

Soluzione dell'esercizio 3.1. La dimostrazione è immediata conseguenza della definizione di media condizionata: posto $Y_n = \mathbb{E}(X_{n+1} | \mathcal{F}_n)$ allora dev'essere

$$\mathbb{E}(1_A Y_n) = \mathbb{E}(1_A X_{n+1}) \quad \text{per ogni } A \in \mathcal{F}_n.$$

Dunque, $Y_n = X_n$ (risp. $\leq, \geq$) se e solo se per ogni $A \in \mathcal{F}_n$ si ha

$$\mathbb{E}(1_A X_{n+1}) = \mathbb{E}(1_A X_n) \quad (\text{risp. } \leq, \geq).$$

Soluzione dell'esercizio 3.2. Usando il fatto che $\mathcal{F}_n \subset \mathcal{F}_{n+1}$ e la proprietà 7 del Paragrafo 3.1.3, si ha

$$\mathbb{E}(X_{n+1} | \mathcal{F}_n) = \mathbb{E}\left(\mathbb{E}(X | \mathcal{F}_{n+1}) \mid \mathcal{F}_n\right) = \mathbb{E}(X | \mathcal{F}_n) = X_n,$$

dunque $(X_n)_n$ è una martingala.

Soluzione dell'esercizio 3.3. Intanto, è immediato verificare che sia Z_n che W_n sono integrabili. Studiamo dapprima $(Z_n)_n$:

$$\begin{aligned} \mathbb{E}(Z_{n+1} | \mathcal{F}_n) &= \mathbb{E}\left(\left(\frac{q}{p}\right)^{X_n+Y_{n+1}} \mid \mathcal{F}_n\right) = \left(\frac{q}{p}\right)^{X_n} \mathbb{E}\left(\left(\frac{q}{p}\right)^{Y_{n+1}} \mid \mathcal{F}_n\right) \\ &= Z_n \left(\frac{q}{p} \cdot \mathbb{P}(Y_{n+1} = 1) + \frac{p}{q} \cdot \mathbb{P}(Y_{n+1} = -1)\right) = Z_n \left(\frac{q}{p} \cdot p + \frac{p}{q} \cdot q\right) = Z_n, \end{aligned}$$

dunque $(Z_n)_n$ è una martingala. Ora, essendo $W_n = 1/Z_n$, con $f(z) = 1/z$ convessa, $(W_n)_n$ è una submartingala. Verifichiamolo anche direttamente:

$$\begin{aligned} \mathbb{E}(W_{n+1} | \mathcal{F}_n) &= \mathbb{E}\left(\left(\frac{p}{q}\right)^{X_n+Y_{n+1}} \mid \mathcal{F}_n\right) = \left(\frac{p}{q}\right)^{X_n} \mathbb{E}\left(\left(\frac{p}{q}\right)^{Y_{n+1}} \mid \mathcal{F}_n\right) \\ &= W_n \left(\frac{p}{q} \cdot \mathbb{P}(Y_{n+1} = 1) + \frac{q}{p} \cdot \mathbb{P}(Y_{n+1} = -1)\right) = W_n \left(\frac{p}{q} \cdot p + \frac{q}{p} \cdot q\right) = c W_n, \end{aligned}$$

dove $c = (p^3 + q^3)/(pq) = (p^2 - pq + q^2)/(pq)$. Osserviamo che $c \geq 1$ per ogni p, q , con $c = 1$ se e solo se $p = q = 1/2$. Dunque, $(W_n)_n$ è una submartingala (e in particolare una-banale-martingala se $p = q$).

Soluzione dell'esercizio 3.4. Poiché $X_{n+1} = X_n \cdot Y_{n+1}$, si ha:

$$\mathbb{E}(X_{n+1} | \mathcal{F}_n) = \mathbb{E}(X_n Y_{n+1} | \mathcal{F}_n) = X_n \mathbb{E}(Y_{n+1} | \mathcal{F}_n) = X_n \mathbb{E}(Y_{n+1}) = \mu X_n,$$

dunque $(X_n)_n$ è una martingala se e solo se $\mu = 1$. Supponiamo ora che $\mathbb{E}(Y_n) = \mu$ e $\text{Var}(Y_n) = \sigma^2$, per ogni n . Allora, usando che $X_{n+1}^2 = X_n^2 \cdot Y_{n+1}^2$:

$$\begin{aligned} \mathbb{E}(Z_{n+1} | \mathcal{F}_n) &= \mathbb{E}\left(X_{n+1}^2 - \sigma^2 \sum_{k=0}^n X_k^2 \mid \mathcal{F}_n\right) \\ &= X_n^2 \mathbb{E}(Y_{n+1}^2 | \mathcal{F}_n) - \sigma^2 \sum_{k=0}^n X_k^2 = X_n^2 \mathbb{E}(Y_{n+1}^2) - \sigma^2 \sum_{k=0}^n X_k^2 \\ &= X_n^2(\sigma^2 + 1) - \sigma^2 \sum_{k=0}^n X_k^2 = X_n^2 - \sigma^2 \sum_{k=0}^{n-1} X_k^2 = Z_n \end{aligned}$$

Soluzione dell'esercizio 3.5. X_n^2 è una submartingala perché $X_n^2 = f(X_n)$, con $f(x) = x^2$ convessa e $f(X_n)$ è integrabile. Il compensatore è $A_n = \sum_{k=0}^{n-1} \mathbb{E}(X_{k+1}^2 - X_k^2 | \mathcal{F}_k)$. Ora,

$$\begin{aligned} \mathbb{E}(X_{k+1}^2 - X_k^2 | \mathcal{F}_k) &= \mathbb{E}((X_k + Y_{k+1})^2 - X_k^2 | \mathcal{F}_k) \\ &= \mathbb{E}(Y_{k+1}^2 + 2X_k Y_{k+1} | \mathcal{F}_k) = \mathbb{E}(Y_{k+1}^2) + 2X_k \mathbb{E}(Y_{k+1}) = \sigma^2, \end{aligned}$$

dunque $A_n = n\sigma^2$.

Soluzione dell'esercizio 3.6. 1. $\Omega \in \mathcal{F}_\tau$: infatti, $\Omega \in \mathcal{F}_\infty$ ($\mathcal{F}_\infty$ è una σ -algebra) e, per ogni n , $\Omega \cap \{\tau \leq n\} = \{\tau \leq n\} \in \mathcal{F}_n$ perché τ è $\mathcal{F}_n$ -misurabile.

2. Se $A \in \mathcal{F}_\tau$ allora $A^c \in \mathcal{F}_\tau$: se $A \in \mathcal{F}_\tau$ allora $A^c \in \mathcal{F}_\infty$ e $A^c \cap \{\tau \leq n\} = \{\tau \leq n\} \cap (A \cap \{\tau \leq n\})^c \in \mathcal{F}_n$ per ogni n (perché sia $\{\tau \leq n\}$ che $A \cap \{\tau \leq n\}$ stanno in $\mathcal{F}_n$).

3. Se $\{A_k\} \subset \mathcal{F}_\tau$ allora $\cup_k A_k \in \mathcal{F}_\tau$: in tal caso, ovviamente $\cup_k A_k \in \mathcal{F}_\infty$ ed inoltre $(\cup_k A_k) \cap \{\tau \leq n\} = \cup_k (A_k \cap \{\tau \leq n\}) \in \mathcal{F}_n$.

Soluzione dell'esercizio 3.7. Per la prima, abbiamo visto che basta verificare che, per ogni n , $\{\tau_1 + \tau_2 = n\} \in \mathcal{F}_n$. Ma si può scrivere

$$\{\tau_1 + \tau_2 = n\} = \bigcup_{k=0}^n \{\tau_1 = k, \tau_2 \leq n - k\}.$$

Ma, se $0 \leq k \leq n$, gli eventi $\{\tau_1 = k\}$ e $\{\tau_2 \leq n - k\}$ appartengono a $\mathcal{F}_n$ (ricordiamo che le σ -algebre $\mathcal{F}_n$ sono crescenti). Dunque, nell'unione a destra tutti gli eventi appartengono a $\mathcal{F}_n$.

Per la seconda, osserviamo che

$$\{\tau_1 \vee \tau_2 \leq n\} = \underbrace{\{\tau_1 \leq n\}}_{\in \mathcal{F}_n} \cap \underbrace{\{\tau_2 \leq n\}}_{\in \mathcal{F}_n} \in \mathcal{F}_n.$$

Infine,

$$\{\tau_1 \wedge \tau_2 > n\} = \underbrace{\{\tau_1 > n\}}_{\in \mathcal{F}_n} \cap \underbrace{\{\tau_2 > n\}}_{\in \mathcal{F}_n} \in \mathcal{F}_n,$$

dunque $\{\tau_1 \wedge \tau_2 \leq n\} = \{\tau_1 \wedge \tau_2 > n\}^c \in \mathcal{F}_n$.

Soluzione dell'esercizio 3.8. Intuitivamente questa relazione è evidente: se il tempo τ_1 precede τ_2 , allora gli eventi antecedenti a τ_1 sono anche antecedenti a τ_2 . Verifichiamo formalmente la proprietà. Sia $A \in \mathcal{F}_{\tau_1}$. Vogliamo verificare che $A \in \mathcal{F}_{\tau_2}$, quindi che

$$A \cap \{\tau_2 \leq n\} \in \mathcal{F}_n, \quad \text{per ogni } n.$$

Ma, poiché $\tau_1 \leq \tau_2$, $\{\tau_2 \leq n\} \subset \{\tau_1 \leq n\}$. Dunque possiamo scrivere

$$A \cap \{\tau_2 \leq n\} = \underbrace{A \cap \{\tau_1 \leq n\}}_{\in \mathcal{F}_n} \cap \underbrace{\{\tau_2 \leq n\}}_{\in \mathcal{F}_n} \in \mathcal{F}_n.$$

Soluzione dell'esercizio 3.9. Infatti, $\tau_1 \wedge \tau_2$ è un tempo d'arresto tale che $\tau_1 \wedge \tau_2 \leq \tau_i$, $i = 1, 2$, dunque $\mathcal{F}_{\tau_1 \wedge \tau_2} \subset \mathcal{F}_{\tau_i}$, $i = 1, 2$, e quindi $\mathcal{F}_{\tau_1 \wedge \tau_2} \subset \mathcal{F}_{\tau_1} \cap \mathcal{F}_{\tau_2}$. Viceversa, dimostriamo che $\mathcal{F}_{\tau_1 \wedge \tau_2} \supset \mathcal{F}_{\tau_1} \cap \mathcal{F}_{\tau_2}$. Preso $A \in \mathcal{F}_{\tau_1} \cap \mathcal{F}_{\tau_2}$, allora $A \cap \{\tau_i \leq n\} \in \mathcal{F}_n$, per ogni n e $i = 1, 2$. Quindi,

$$\begin{aligned} A \cap \{\tau_1 \wedge \tau_2 \leq n\} &= A \cap \left(\{\tau_1 \leq n\} \cup \{\tau_2 \leq n\} \right) \\ &= \left(A \cap \{\tau_1 \leq n\} \right) \cup \left(A \cap \{\tau_2 \leq n\} \right) \in \mathcal{F}_n, \end{aligned}$$

cioè $A \in \mathcal{F}_{\tau_1 \wedge \tau_2}$.

Soluzione dell'esercizio 3.10. Osserviamo dapprima che

$$\begin{aligned} \{\tau_1 < \tau_2\} &= \cup_k \underbrace{\{\tau_1 < k\} \cap \{\tau_2 = k\}}_{\in \mathcal{F}_k \subset \mathcal{F}_\infty} \in \mathcal{F}_\infty \quad \text{e} \\ \{\tau_1 = \tau_2\} &= \cup_k \underbrace{\{\tau_1 = k\} \cap \{\tau_2 = k\}}_{\in \mathcal{F}_k \subset \mathcal{F}_\infty} \in \mathcal{F}_\infty. \end{aligned}$$

Inoltre, per ogni n ,

$$\begin{aligned} \{\tau_1 < \tau_2\} \cap \{\tau_1 = n\} &= \{\tau_1 = n\} \cap \{\tau_2 \leq n\}^c \in \mathcal{F}_n \quad \text{e} \\ \{\tau_1 < \tau_2\} \cap \{\tau_2 = n\} &= \{\tau_1 < n\} \cap \{\tau_2 = n\}^c \in \mathcal{F}_n. \end{aligned}$$

Dunque, $\{\tau_1 < \tau_2\} \in \mathcal{F}_{\tau_1}$ e $\{\tau_1 < \tau_2\} \in \mathcal{F}_{\tau_2}$, cioè $\{\tau_1 < \tau_2\} \in \mathcal{F}_{\tau_1} \cap \mathcal{F}_{\tau_2}$.

Capitolo 4

Modelli probabilistici per la finanza

4.1 Introduzione

In questo capitolo introdurremo un modello probabilistico utile per lo studio di alcuni problemi di finanza matematica, a cui abbiamo già accennato nel primo capitolo. La trattazione segue per lo più il testo di Lamberton e Lapeyre [6].

Considereremo i problemi legati al calcolo del prezzo delle *opzioni*, che in passato è stata la motivazione principale per la costruzione della teoria della finanza matematica e tuttora rappresenta un esempio rilevante di utilizzo della teoria della probabilità e, in particolare, del calcolo stocastico.

Le opzioni sono state brevemente introdotte nel Capitolo 1, riprendiamo ora gli aspetti più importanti. Il nostro obiettivo è la costruzione di un modello matematico adatto a studiarle.

Una *opzione* dà a colui che la detiene il diritto, ma non l'obbligo, di comprare o vendere un bene ad una data futura prefissata e ad un prezzo prefissato oggi. Un po' di terminologia:

- il bene cui si riferisce l'opzione è detto *bene o attivo sottostante*;
- l'istante in cui si può esercitare l'opzione è detto *maturità* o anche *data d'esercizio*;
- il prezzo prefissato è il *prezzo d'esercizio*.

Dunque, per caratterizzare un'opzione su un attivo finanziario occorre specificare

- il tipo di opzione: ad esempio, l'opzione a comprare si chiama *call*, l'opzione a vendere *put*;

- le modalità di esercizio: se l'opzione può essere esercitata esclusivamente alla data di maturità, l'opzione si dice *europea*; esistono anche opzioni, dette *americane*, che possono essere esercitate in un qualsiasi istante prima della maturità;
- il prezzo di esercizio K e la data di maturità T .

In questo capitolo ci occuperemo delle opzioni europee, rimandando al prossimo la trattazione delle opzioni americane.

Un'opzione costa. È evidente infatti che il detentore dell'opzione acquisisce un vantaggio, mentre chi ha ceduto l'opzione si è assunto dei rischi, che è giusto che vengano remunerati. Qual è il giusto compenso? Il prezzo dell'opzione è anche detto *premio*. Comunque, di solito, l'opzione è trattata in un mercato finanziario organizzato ed il prezzo è determinato dal mercato stesso. È comunque un problema interessante costruire un modello teorico per calcolare il prezzo, anche quando il premio è quotato dal mercato. Ciò permette infatti di studiare ed eventualmente individuare anomalie nelle quotazioni.

Ci si può aspettare che il prezzo di una data opzione debba dipendere dal comportamento dell'attivo sottostante (dal fatto che esso tenda a crescere o a decrescere oppure che sia soggetto o meno a forti oscillazioni). Vedremo che il prezzo di una opzione dipende dal comportamento dell'attivo sottostante unicamente per una sola quantità σ (la volatilità).

Il detentore dell'opzione è anche detto il *compratore dell'opzione*. Colui che invece cede l'opzione al detentore è detto *venditore dell'opzione*.

Per esempio, vediamo un po' più in dettaglio cosa succede nel caso della call europea (opzione a comprare). Indichiamo con t il tempo, supponendo $t \geq 0$ e $t = 0$ ha il significato di "oggi". Indichiamo ora con S_t il valore (=prezzo) di mercato all'istante t del bene sottostante l'opzione. Siano poi T l'istante di maturità e K il prezzo di esercizio della call. Il valore del bene sottostante a maturità è dato da S_T e la decisione di esercitare o meno l'opzione segue dal confronto tra S_T e K . Infatti,

se $S_T > K$, conviene comprare il bene al prezzo di esercizio K piuttosto che al prezzo di mercato S_T : *l'opzione è esercitata*;

se $S_T < K$, conviene comprare il bene al prezzo di mercato S_T piuttosto che al prezzo di esercizio K : *l'opzione non è esercitata*.

Dunque, la perdita a cui si espone il venditore di una opzione call, detto anche *valore o payoff della call*, è dato da

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

Si noti che S_T è il valore a maturità (cioè a $t = T$) del bene sottostante. In particolare è un valore futuro che oggi ($t = 0$) non si conosce.

I problemi di interesse, come vedremo sono due:

1. Quanto deve pagare il detentore il suo diritto di opzione?
2. Che tipo di investimenti deve fare il venditore dell'opzione per far fronte al contratto? In altre parole, come si fa a generare un ammontare di denaro che a maturità T dev'essere pari a $(S_T - K)_+$? Questo è il problema della *copertura dell'opzione*.

Considerazioni analoghe si possono fare nel caso di opzione put (opzione a vendere). Il valore della put europea è quindi dato da

$$(K - S_T)_+ = \max(K - S_T, 0) = (S_T - K)_-$$

dove, come sopra, K denota il prezzo di esercizio, T la maturità e S_T il valore del bene sottostante a maturità. Infatti, nel caso della put,

se $S_T < K$, conviene vendere il bene al prezzo di esercizio K piuttosto che al prezzo di mercato S_T : *l'opzione è esercitata*;

se $S_T > K$, conviene vendere il bene al prezzo di mercato S_T piuttosto che al prezzo di esercizio K : *l'opzione non è esercitata*.

Ricordiamo che sono in vigore le ipotesi sul mercato, che abbiamo fatto nel Capitolo 1 e in particolare nel paragrafo 1.2.3. Ad esse bisogna aggiungere quella di assenza di arbitraggio. Ricordiamo che una operazione di arbitraggio è un'operazione finanziaria tale che

- non occorre in alcun momento investire del capitale;
- non si può avere in alcun caso una perdita e si ha, con probabilità positiva, un guadagno.

Nel Capitolo 1, abbiamo già visto alcuni esempi di arbitraggio ed abbiamo anche osservato che, poiché i mercati finanziari sono molto liquidi (ovvero si possono effettuare numerose transazioni), la richiesta di assenza di arbitraggio è comunemente accettata nei modelli che descrivono i mercati finanziari.

In questo capitolo, daremo all'arbitraggio una definizione matematicamente rigorosa. Osserviamo che l'ipotesi di assenza di arbitraggio si traduce in condizioni che bisogna imporre al *modello matematico* che costruiremo. In altre parole, quando si costruisce un modello matematico per questo tipo di problemi sarà sempre necessario accertarsi che esso soddisfi alla proprietà di assenza di arbitraggio.

4.2 Modelli discreti per la descrizione dei mercati finanziari

In questo paragrafo introduciamo un modello discreto, abbastanza semplice, che però permetterà di descrivere l'evoluzione dei prezzi in un mercato finanziario.

Sia $(\Omega, \mathcal{F}, \mathbb{P})$ uno spazio di probabilità finito: Ω è un insieme di cardinalità finita e $\mathcal{F} = \mathcal{P}(\Omega)$. Supporremo anche che $\mathbb{P}(\{\omega\}) > 0$ per ogni $\omega \in \Omega$. Su $(\Omega, \mathcal{F}, \mathbb{P})$ supponiamo sia definita una filtrazione $(\mathcal{F}_n)_{n=0,1,\dots,N}$, con

$$\{\emptyset, \Omega\} = \mathcal{F}_0 \subset \mathcal{F}_1 \subset \dots \subset \mathcal{F}_N = \mathcal{F} = \mathcal{P}(\Omega).$$

L'indice n ha il significato di tempo (discreto), mentre la σ -algebra $\mathcal{F}_n$ rappresenta l'informazione data dal mercato fino all'istante n . Il tempo (finale) N rappresenterà invece la maturità dell'opzione.

Il mercato finanziario che consideriamo consiste di $d + 1$ titoli finanziari, i cui prezzi al tempo n sono dati dalle $d + 1$ componenti del vettore aleatorio

$$S_n = (S_n^0, S_n^1, \dots, S_n^d).$$

Poché gli investitori conoscono presente e passato del mercato (ma ovviamente non il futuro), supporremo che, per ogni n , la v.a. S_n sia $\mathcal{F}_n$ -misurabile. In particolare, si ha S_0 deterministico¹.

Supporremo che il titolo di prezzo S_n^0 sia il *titolo non rischioso*. Se si suppone che il rendimento r su un periodo unitario sia costante, abbiamo già visto nel primo capitolo che

$$S_n^0 = S_0^0(1 + r)^n$$

Normalmente, per convenzione, si pone $S_0^0 = 1$, dunque

$$S_n^0 = (1 + r)^n$$

Per $n = 0, 1, \dots, N$ sia

$$\beta_n = \frac{1}{S_n^0} = (1 + r)^{-n}$$

il *coefficiente di sconto* al tempo n : se al tempo 0 si investe a tasso costante una quantità di denaro pari a x , al tempo n si riceverà una quantità pari a

$$x_n = x(1 + r)^n = \frac{x}{\beta_n}.$$

Quindi, in particolare β_n , è l'ammontare (in Euro) che occorre investire nel titolo non rischioso per avere disporre al tempo n di 1 Euro.

I rimanenti titoli indicizzati con $1, 2, \dots, d$ sono detti *titoli rischiosi*. Ci troveremo spesso a considerare i valori scontati dei prezzi $\tilde{S}_n^1, \dots, \tilde{S}_n^d$, definiti da

$$\tilde{S}_n^1 = (1 + r)^{-n} S_n^1, \quad \dots \quad \tilde{S}_n^d = (1 + r)^{-n} S_n^d$$

¹Ricordiamo infatti che una v.a. misurabile rispetto alla σ -algebra banale $\mathcal{F}_0 = \{\emptyset, \Omega\}$ è necessariamente costante.

$\tilde{S}_n$ indicherà il vettore dei prezzi scontati, cioè

$$\tilde{S}_n = (1+r)^{-n} S_n = (1, \tilde{S}_n^1, \dots, \tilde{S}_n^d)$$

4.3 Strategie ed arbitraggio

Indicheremo con il termine *portafoglio* un insieme di titoli presi tra i $d+1$ titoli presenti sul mercato, ciascuno dei quali in una certa quantità. Vediamo meglio come si definisce matematicamente un portafoglio.

Definizione 4.3.1. Chiameremo *strategia di gestione*, o più semplicemente una *strategia*, è una sequenza predicibile di v.a. $\phi = (\phi_n)_n$ a valori in $\mathbb{R}^{d+1}$, dove $\phi_n = (\phi_n^0, \phi_n^1, \dots, \phi_n^d)$.

Le quantità ϕ_n^i della Definizione 4.3.1 rappresentano la quantità (numero di quote) dell' i -esimo titolo presente nel portafoglio. In particolare il valore del portafoglio al tempo n per una strategia ϕ è pari a

$$\phi_n^0 S_n^0 + \phi_n^1 S_n^1 + \dots + \phi_n^d S_n^d$$

Ricordiamo che una successione di v.a. $\phi = (\phi_n)_n$ si dice *predicibile* se

$$\phi_0 \text{ è } \mathcal{F}_0\text{-misurabile} \quad \text{e} \quad \text{per ogni } n \geq 1, \phi_n \text{ è } \mathcal{F}_{n-1}\text{-misurabile}$$

In particolare una strategia ϕ è adattata alla filtrazione $(\mathcal{F}_n)_n$, cioè, ϕ_n è $\mathcal{F}_n$ -misurabile, per ogni n .

La richiesta di predicibilità è abbastanza naturale. Imporre che ϕ_n sia $\mathcal{F}_{n-1}$ -misurabile significa tenere conto del fatto, che succede nella realtà, che l'investitore deve stabilire le quantità $\phi_n^0, \phi_n^1, \dots, \phi_n^d$ *prima* di conoscere i valori dei prezzi al tempo n .

Ad esempio, se il tempo indica i giorni e $S_n^0, S_n^1, \dots, S_n^d$ sono i prezzi allo n -esimo giorno, allora in chiusura dello n -esimo giorno, l'investitore dispone di un portafoglio $\phi_n^0, \phi_n^1, \dots, \phi_n^d$ ed il suo portafoglio vale

$$\phi_n^0 S_n^0 + \phi_n^1 S_n^1 + \dots + \phi_n^d S_n^d \tag{4.1}$$

Egli può decidere di vendere delle quote, acquistarne altre, eventualmente aggiungere nuovo denaro o toglierne, modificando la ripartizione del capitale investito. Per prendere queste decisioni egli dispone però solo delle informazioni fino al tempo n . Indichiamo con $\phi_{n+1}^0, \phi_{n+1}^1, \dots, \phi_{n+1}^d$ le quote dei vari attivi presenti nel portafoglio, dopo questa operazione di aggiustamento. Queste quantità devono quindi essere $\mathcal{F}_n$ -misurabili. Il suo capitale sarà ora pari a

$$\phi_{n+1}^0 S_n^0 + \phi_{n+1}^1 S_n^1 + \dots + \phi_{n+1}^d S_n^d \tag{4.2}$$

Osservazione 4.3.2. Nella Definizione 4.3.1, si richiede $\phi_n \in \mathbb{R}^{d+1}$, quindi $\phi_n^i \in \mathbb{R}$ per ogni $i = 0, 1, \dots, d$. In particolare si assume che queste v.a. possano prendere valori negativi oppure non interi. Osserviamo che questa condizione riflette le ipotesi che sono state fatte sul mercato nel paragrafo 1.2.3. In particolare è possibile effettuare vendite o acquisti allo scoperto. Poiché, come abbiamo visto, ϕ_n^i rappresenta la quantità investita nel titolo i , queste quantità ϕ_n^i possono prendere valori negativi. Inoltre è possibile acquistare frazioni di un bene finanziario, quindi le quantità ϕ_n^i possono prendere valori non interi.

Data una strategia ϕ , ad essa si possono associare alcune quantità d'interesse.

- Il valore, al tempo n , del portafoglio associato:

$$V_n(\phi) = \langle \phi_n, S_n \rangle = \sum_{i=0}^d \phi_n^i S_n^i.$$

- Il valore scontato, al tempo n , del portafoglio associato:

$$\tilde{V}_n(\phi) = \langle \phi_n, \tilde{S}_n \rangle = \sum_{i=0}^d \phi_n^i \tilde{S}_n^i.$$

dove

$$\tilde{S}_n = (1+r)^{-n} S_n = (1, \tilde{S}_n^1, \dots, \tilde{S}_n^d)$$

è il vettore dei valori scontati dei prezzi. Da notare che la prima coordinata di quest'ultimo è sempre uguale a 1. Da notare anche che, evidentemente, $\tilde{V}_0(\phi) = V_0(\phi)$

Nella classe di tutte le possibili strategie, particolare rilevanza hanno quelle autofinanzianti, soddisfacenti alla seguente

Definizione 4.3.3. Una strategia ϕ è detta *autofinanziante* se per ogni $n = 0, 1, \dots, N-1$ si ha

$$\langle \phi_n, S_n \rangle = \langle \phi_{n+1}, S_n \rangle.$$

Il significato di una strategia autofinanziante è abbastanza semplice. La quantità $\langle \phi_n, S_n \rangle$ è la stessa cosa che la (4.1), cioè il capitale investito al tempo n . Invece $\langle \phi_{n+1}, S_n \rangle$ è la stessa cosa che la (4.2). Dunque una strategia è autofinanziante se ogni aggiustamento delle quote viene fatto senza aggiunta o ritiro di capitale.

Osservazione 4.3.4. L'uguaglianza $\langle \phi_n, S_n \rangle = \langle \phi_{n+1}, S_n \rangle$ è equivalente a

$$\langle \phi_{n+1}, S_{n+1} - S_n \rangle = \langle \phi_{n+1}, S_{n+1} \rangle - \langle \phi_n, S_n \rangle$$

o anche

$$V_{n+1}(\phi) - V_n(\phi) = \langle \phi_{n+1}, S_{n+1} - S_n \rangle.$$

Questa seconda uguaglianza è particolarmente significativa: una strategia ϕ è autofinanziante se la variazione del portafoglio (guadagno se > 0 , perdita se < 0) $V_{n+1}(\phi) - V_n(\phi)$ tra n e $n + 1$ dipende dalla variazione dei prezzi $S_{n+1} - S_n$ e dalle loro quote ϕ_n detenute al tempo n .

La proposizione che segue dà un'ulteriore caratterizzazione delle strategie autofinanzianti.

Proposizione 4.3.5. *Le seguenti affermazioni sono equivalenti.*

i) *La strategia ϕ è autofinanziante.*

ii) *Per ogni $n = 0, 1, \dots, N$,*

$$V_n(\phi) = V_0(\phi) + \sum_{j=1}^n \langle \phi_j, \Delta S_j \rangle,$$

dove $\Delta S_j = S_j - S_{j-1}$ è il vettore delle variazioni dei prezzi al tempo j .

iii) *Per ogni $n = 0, 1, \dots, N$,*

$$\tilde{V}_n(\phi) = V_0(\phi) + \sum_{j=1}^n \langle \phi_j, \Delta \tilde{S}_j \rangle = V_0(\phi) + \sum_{j=1}^n \left(\phi_j^1 \Delta \tilde{S}_j^1 + \dots + \phi_j^d \Delta \tilde{S}_j^d \right) \quad (4.3)$$

dove $\Delta \tilde{S}_j = \tilde{S}_j - \tilde{S}_{j-1} = \beta_j S_j - \beta_{j-1} S_{j-1}$ è il vettore delle variazioni dei prezzi scontati al tempo j .

Dimostrazione. (i) $\Rightarrow$ (ii). Se ϕ è autofinanziante allora, per l'Osservazione 4.3.4,

$$V_n(\phi) = V_0(\phi) + \sum_{j=1}^n (V_j(\phi) - V_{j-1}(\phi)) = V_0(\phi) + \sum_{j=1}^n \langle \phi_j, \Delta S_j \rangle.$$

(ii) $\Rightarrow$ (i). Se vale (ii) allora

$$V_{n+1}(\phi) - V_n(\phi) = \langle \phi_n, \Delta S_{n+1} \rangle$$

e ancora l'Osservazione 4.3.4 garantisce che ϕ è autofinanziante.

(i) $\Leftrightarrow$ (iii). Poiché $\langle \phi_n, S_n \rangle = \langle \phi_{n+1}, S_n \rangle$ se e solo se $\langle \phi_n, \tilde{S}_n \rangle = \langle \phi_{n+1}, \tilde{S}_n \rangle$, basta procedere come sopra. La seconda uguaglianza in (iii) segue immediatamente dal fatto che $\Delta \tilde{S}_j^0 = \beta_j S_j^0 - \beta_{j-1} S_{j-1}^0 = 0$.

□

Osserviamo che la relazione (4.3) afferma, in particolare, che per una strategia autofinanziante, il valore del portafoglio al tempo n non dipende dalle quantità investite nel titolo non rischioso ϕ_n^0 . Più precisamente, si ha

Proposizione 4.3.6. Siano $((\phi_n^1, \dots, \phi_n^d))_{0 \leq n \leq N}$ un processo predicibile e V_0 una v.a. $\mathcal{F}_0$ -misurabile. Allora esiste un unico processo predicibile $(\phi_n^0)_{0 \leq n \leq N}$ tale che $\phi = ((\phi_n^0, \phi_n^1, \dots, \phi_n^d))_{0 \leq n \leq N}$ sia una strategia autofinanziante ed il valore del portafoglio associato a ϕ abbia valore iniziale V_0 : $V_0(\phi) = V_0$.

Dimostrazione. Il valore scontato del portafoglio associato a ϕ è

$$\tilde{V}_n(\phi) = \phi_n^0 + \phi_n^1 \tilde{S}_n^1 + \dots + \phi_n^d \tilde{S}_n^d$$

e abbiamo visto che ϕ è una strategia autofinanziante se e solo se

$$\tilde{V}_n(\phi) = V_0(\phi) + \sum_{j=1}^n \left(\phi_j^1 \Delta \tilde{S}_j^1 + \dots + \phi_j^d \Delta \tilde{S}_j^d \right).$$

Se $V_0(\phi) = V_0$, ϕ_n^0 è dunque univocamente individuato da

$$\begin{aligned} \phi_n^0 &= V_0 + \sum_{j=1}^n \left(\phi_j^1 \Delta \tilde{S}_j^1 + \dots + \phi_j^d \Delta \tilde{S}_j^d \right) - \left(\phi_n^1 \tilde{S}_n^1 + \dots + \phi_n^d \tilde{S}_n^d \right) = \\ &= V_0 + \sum_{j=1}^{n-1} \left(\phi_j^1 \Delta \tilde{S}_j^1 + \dots + \phi_j^d \Delta \tilde{S}_j^d \right) - \left(\phi_n^1 \tilde{S}_{n-1}^1 + \dots + \phi_n^d \tilde{S}_{n-1}^d \right). \end{aligned}$$

Basta quindi mostrare che $(\phi_n^0)_{0 \leq n \leq N}$ è predicibile, ma questo è immediato, poiché si vede che ϕ_n^0 è una funzione che dipende solo dalle quantità ϕ_k^i con $k \leq n$, S_j^i con $j \leq n-1$, che sono tutte quantità $\mathcal{F}_{n-1}$ -misurabili. $\square$

Osservazione 4.3.7. Osserviamo che, se i prezzi scontati dei titoli rischiosi, $(\tilde{S}_n^i)_n$, $i = 1, \dots, d$ fossero delle martingale, allora anche ogni portafoglio scontato sarebbe una martingala. Infatti, per $i = 1, \dots, d$, il processo

$$\sum_{j=1}^n \phi_j^i \Delta \tilde{S}_j^i = \sum_{j=1}^n \phi_j^i (\tilde{S}_j^i - \tilde{S}_{j-1}^i)$$

risulterebbe essere una martingala trasformata.

Abbiamo visto che un portafoglio può contenere anche quantità negative di un titolo. Sarà però ragionevole imporre talvolta che esso debba comunque avere un valore che sia globalmente positivo.

Definizione 4.3.8. Una strategia ϕ è detta *ammissibile* se è autofinanziante e se $V_n(\phi) \geq 0$ per ogni $n = 0, 1, \dots, N$.

Possiamo ora formalizzare il concetto di arbitraggio, ossia la possibilità di effettuare profitti senza rischio:

Definizione 4.3.9. Una *strategia di arbitraggio* è una strategia ammissibile ϕ , tale che il valore iniziale del portafoglio associato è nullo e il valore finale è non nullo.

Una strategia ϕ quindi è d'arbitraggio se è autofinanziante e se

$$\begin{aligned} V_0(\phi) &= 0 \\ V_n(\phi) &\geq 0 \text{ per ogni } n = 1, 2, \dots, N \\ \mathbb{P}(V_N(\phi) > 0) &> 0. \end{aligned}$$

4.4 Arbitraggio e martingale

È naturale porre la seguente

Definizione 4.4.1. Si dice che si è in *assenza di arbitraggio* o anche che il mercato è *privo di arbitraggio* se non è possibile costruire strategie di arbitraggio.

In questo paragrafo vedremo che la condizione di assenza di arbitraggio è equivalente ad una condizione molto interessante da un punto di vista matematico.

Osservazione 4.4.2. La Definizione 4.4.1 si può riscrivere in termini matematici nel modo seguente. Sia Γ la famiglia di v.a.

$$\begin{aligned} \Gamma &= \{X : \Omega \rightarrow \mathbb{R} : \mathbb{P}(X \geq 0) = 1 \text{ e } \mathbb{P}(X > 0) > 0\} \\ &= \{X : \Omega \rightarrow \mathbb{R} : X(\omega) \geq 0 \text{ per ogni } \omega \text{ ed esiste } \bar{\omega} \text{ t.c. } X(\bar{\omega}) > 0\}. \end{aligned}$$

(ricordiamo infatti che Ω è finito e $\mathbb{P}(\{\omega\}) > 0$ per ogni $\omega \in \Omega$). Allora si può dire che il mercato è privo di arbitraggio se

per ogni strategia ammissibile ϕ tale che $V_0(\phi) = 0$ allora $V_N(\phi) \notin \Gamma$.

Osserviamo inoltre che Γ è un *cono convesso*: se $X \in \Gamma$, allora anche $\lambda X \in \Gamma$ per ogni $\lambda > 0$ (dunque è un cono); in più, se $0 \leq t \leq 1$ e $X, Y \in \Gamma$, allora $tX + (1-t)Y \in \Gamma$ (dunque Γ è convesso).

Lemma 4.4.3. Sia $((\rho_n^1, \dots, \rho_n^d))_n$ una sequenza predicibile a valori in $\mathbb{R}^d$ e sia

$$\tilde{G}_n(\rho^1, \dots, \rho^d) = \sum_{j=1}^n \left(\rho_j^1 \Delta \tilde{S}_j^1 + \dots + \rho_j^d \Delta \tilde{S}_j^d \right), \quad n = 0, 1, \dots, N.$$

Se il mercato è privo di arbitraggio allora $G_N(\rho^1, \dots, \rho^d) \notin \Gamma$.

Dimostrazione Supponiamo per assurdo che $G_N(\rho^1, \dots, \rho^d) \in \Gamma$. Poiché $((\rho_n^1, \dots, \rho_n^d))_n$ è predicibile, per la Proposizione 4.3.6 esiste un processo predicibile $(\rho_n^0)_n$ tale che $\rho = ((\rho_n^0, \rho_n^1, \dots, \rho_n^d))_n$ è una strategia autofinanziante. Inoltre si ha $\tilde{G}_n(\rho^1, \dots, \rho^d) = \tilde{V}_n(\rho)$, con la condizione $V_0(\rho) = 0$.

Dunque, se fosse $\tilde{G}_n(\rho^1, \dots, \rho^d) \geq 0$ per ogni n , ρ sarebbe una strategia non solo autofinanziante ma anche ammissibile tale che $\mathbb{P}(V_N(\rho)) > 0) > 0$, dunque una strategia di arbitraggio, il che non è possibile. Dunque, dev'essere

$$\mathbb{P}(\tilde{G}_n(\rho^1, \dots, \rho^d) < 0) > 0 \text{ per qualche } n < N.$$

Mostriamo che allora è possibile costruire una strategia di arbitraggio. Poniamo

$$n_0 = \sup\{n : \mathbb{P}(\tilde{G}_n(\rho^1, \dots, \rho^d) < 0) > 0\}.$$

È chiaro che $n_0 \leq N - 1$, e inoltre si ha

$$\begin{aligned} \mathbb{P}(\tilde{G}_{n_0}(\rho^1, \dots, \rho^d) < 0) &> 0 \\ \tilde{G}_n(\rho^1, \dots, \rho^d) &\geq 0 \quad \text{per ogni } n = n_0 + 1, \dots, N \end{aligned}$$

Poniamo $A = \{\tilde{G}_{n_0}(\rho^1, \dots, \rho^d) < 0\}$ e definiamo il processo ϕ nel modo seguente: per $i = 1, \dots, d$,

$$\phi_n^i = \begin{cases} 0 & \text{se } n \leq n_0 \\ \rho_n^i 1_A & \text{se } n > n_0. \end{cases}$$

Verifichiamo che $(\phi_n^i)_n$ è predicibile: per $n \leq n_0$ si ha $\phi_n^i = 0$, dunque $\phi_n^i = 0$ è $\mathcal{F}_{n-1}$ -misurabile; se invece $n > n_0$, ϕ_n^i è una funzione misurabile di ρ_n^i , che è $\mathcal{F}_{n-1}$ -misurabile, e di $\tilde{G}_{n_0}(\rho^1, \dots, \rho^d)$, che è $\mathcal{F}_{n_0}$ -misurabile e quindi anche $\mathcal{F}_{n-1}$ -misurabile. Per la Proposizione 4.3.6, esiste un processo $(\phi_n^0)_n$, predicibile, tale che $\phi = ((\phi_n^0, \phi_n^1, \dots, \phi_n^d))_n$ sia una strategia autofinanziante di portafoglio scontato

$$\tilde{V}_n(\phi) = \sum_{j=1}^n \left(\phi_j^1 \Delta \tilde{S}_j^1 + \dots + \phi_j^d \Delta \tilde{S}_j^d \right)$$

(abbiamo scelto $V_0 = 0$). Sostituendo ϕ , otteniamo facilmente

$$\tilde{V}_n(\phi) = \begin{cases} 0 & \text{se } n \leq n_0 \\ (\tilde{G}_n(\rho) - \tilde{G}_{n_0}(\rho)) 1_A & \text{se } n > n_0. \end{cases}$$

Poiché $\tilde{G}_n(\rho) \geq 0$ per ogni $n \geq n_0$, mentre $\tilde{G}_{n_0}(\rho) < 0$ su A , ne segue che ogni n , dunque ϕ è ammissibile e che $\mathbb{P}(V_N(\phi) > 0) = \mathbb{P}(\tilde{G}_{n_0}(\rho^1, \dots, \rho^d) < 0) > 0$. Poiché $V_0(\phi) = 0$, ϕ è una strategia di arbitraggio. $\square$

Useremo nella dimostrazione del prossimo teorema il classico *teorema di separazione dei convessi*, nella forma seguente:

Teorema 4.4.4. *Siano $K \subset \mathbb{R}^m$ un insieme compatto e convesso e $\mathcal{V} \subset \mathbb{R}^m$ un sottospazio tali che $K \cap \mathcal{V} = \emptyset$. Allora, esiste $\lambda \in \mathbb{R}^m$ tale che*

$$\begin{aligned} \text{per ogni } x \in K & \quad \sum_{\ell=1}^m \lambda_\ell x_\ell > 0 \\ \text{per ogni } x \in \mathcal{V} & \quad \sum_{\ell=1}^m \lambda_\ell x_\ell = 0. \end{aligned}$$

In altre parole, il sottospazio $\mathcal{V}$ è contenuto in un iperpiano chiuso che non contiene K .

Per la dimostrazione e maggiori dettagli sul significato intuitivo di questo risultato, si rimanda all'Appendice a questo capitolo (cfr. Paragrafo 4.9.1). Vediamo ora il risultato principale di questo paragrafo. Esso prende il nome di *primo teorema fondamentale dell'asset pricing*.

Teorema 4.4.5. *Un mercato è privo di arbitraggio se e solo se esiste una misura di probabilità $\mathbb{P}^*$ equivalente a $\mathbb{P}$ tale che il vettore dei prezzi scontati dei titoli rischiosi è una $\mathbb{P}^*$ -martingala.*

Ricordiamo che due misure di probabilità $\mathbb{P}$ e $\mathbb{P}^*$ su $(\Omega, \mathcal{F})$ sono *equivalenti* se hanno gli stessi insiemi di misura nulla cioè se $\mathbb{P}(A) = 0$ se e solo se $\mathbb{P}^*(A) = 0$, per $A \in \mathcal{F}$. Ora, poiché supponiamo Ω finito, $\mathcal{F} = \mathcal{P}(\Omega)$ e $\mathbb{P}(\{\omega\}) > 0$ per ogni $\omega \in \Omega$, in questo caso ovviamente $\mathbb{P}^*$ è equivalente a $\mathbb{P}$ se e solo se $\mathbb{P}^*(\{\omega\}) > 0$ per ogni $\omega \in \Omega$.

Inoltre, sempre perché Ω è finito, possiamo identificare ogni v.a. X su Ω con il vettore $(X(\omega))_{\omega \in \Omega}$. Indicheremo con $\mathbb{R}^\Omega$ l'insieme di questi vettori².

Dimostrazione Prima parte: supponiamo che esista una misura di probabilità $\mathbb{P}^*$ equivalente a $\mathbb{P}$, tale che il vettore dei prezzi scontati $((\tilde{S}_n^1, \dots, \tilde{S}_n^d))_n$ sia una $\mathbb{P}^*$ -martingala e mostriamo che in tal caso il mercato è privo di arbitraggio.

Sia $\phi = (\phi_n)_n$ una strategia ammissibile (ovvero, autofinanziante e tale che $V_n(\phi) \geq 0$ per ogni n). Allora, dalla Proposizione 4.3.5 segue che

$$\tilde{V}_n(\phi) = V_0(\phi) + \sum_{j=1}^n \langle \phi_j, \Delta \tilde{S}_j \rangle.$$

Poiché ϕ_n è predicibile e $(\tilde{S}_n^i)_n$ è una $\mathbb{P}^*$ -martingala per ogni $i = 1, \dots, d$, $(\tilde{V}_n(\phi))_n$ è una $\mathbb{P}^*$ -martingala (si veda l'Osservazione 4.3.7), dunque, in particolare,

$$\mathbb{E}^*(\tilde{V}_N(\phi)) = \mathbb{E}^*(V_0(\phi)),$$

dove con $\mathbb{E}^*$ indichiamo l'aspettazione rispetto alla misura $\mathbb{P}^*$. Se supponiamo $V_0(\phi) = 0$, si ha

$$\mathbb{E}^*(\tilde{V}_N(\phi)) = 0.$$

²Osserviamo che $\mathbb{R}^\Omega$ è (nel senso che sono isomorfi) $\mathbb{R}^m$, essendo $m = \text{card}(\Omega)$. Infatti, indicando ad esempio $\Omega = \{\omega_1, \dots, \omega_m\}$, sia $T : \mathbb{R}^\Omega \rightarrow \mathbb{R}^m$, $T(X) = (X(\omega_1), \dots, X(\omega_m))$. È immediato verificare che T è un isomorfismo, con inversa $T^{-1} : \mathbb{R}^m \rightarrow \mathbb{R}^\Omega$, $T^{-1}(x_1, \dots, x_m) = X$ con X tale che $X(\omega_i) = x_i$, $i = 1, \dots, m$.

Poiché per ipotesi ϕ è ammissibile, in particolare si ha $V_N(\phi) \geq 0$. Dunque, otteniamo $V_N(\phi) \geq 0$ e $\mathbb{E}^*(\tilde{V}_N(\phi)) = 0$: ciò significa che necessariamente $\mathbb{P}^*(\tilde{V}_N(\phi) > 0) = 0$. Poiché $\mathbb{P}$ è equivalente a $\mathbb{P}^*$, si ha anche $\mathbb{P}(\tilde{V}_N(\phi) > 0) = 0$ ed otteniamo che ogni strategia ammissibile di valore iniziale nullo è necessariamente tale che $\mathbb{P}(\tilde{V}_N(\phi) > 0) = 0$, dunque non esistono strategie di arbitraggio.

Seconda parte: supponiamo che il mercato sia privo di arbitraggio e dimostriamo l'esistenza di una misura di probabilità $\mathbb{P}^*$ sotto la quale i prezzi scontati sono una martingala.

A tale scopo, utilizziamo la Proposizione 3.2.7: mostriamo che esiste una misura $\mathbb{P}^*$ equivalente a $\mathbb{P}$ tale che per ogni processo predicibile $(\rho_n^i)_n$ si ha che $\mathbb{E}^*(\sum_{n=1}^N \rho_n^i \Delta \tilde{S}_n^i) = 0$, per ogni $i = 1, \dots, d$. È immediato vedere che ciò equivale a dimostrare che

$$\mathbb{E}^*(\tilde{G}_N(\rho)) = 0, \text{ essendo } \tilde{G}_N(\rho) = \sum_{n=1}^N \left(\rho_n^1 \Delta \tilde{S}_n^1 + \dots + \rho_n^d \Delta \tilde{S}_n^d \right) \quad (4.4)$$

per ogni processo predicibile $((\rho_n^1, \dots, \rho_n^d))_n$ a valori in $\mathbb{R}^d$. Il Lemma 4.4.3 garantisce che $\tilde{G}_N(\rho) \notin \Gamma$. Ora, l'insieme

$$\mathcal{V} = \{ \tilde{G}_N(\rho); ((\rho_n^1, \dots, \rho_n^d))_n \text{ processo predicibile} \},$$

è un sottospazio di $\mathbb{R}^\Omega$ (cioè dell'insieme delle v.a. su Ω) e, per il Lemma 4.4.3, $\mathcal{V} \cap \Gamma = \emptyset$. Dunque, a maggior ragione, se poniamo

$$K = \{ X \in \Gamma; \sum_{\omega \in \Omega} X(\omega) = 1 \}$$

si ha

$$\mathcal{V} \cap K = \emptyset.$$

È immediato verificare che $K \subset \mathbb{R}^\Omega$ è convesso e compatto, (è chiuso e limitato) Possiamo dunque applicare il Teorema 4.4.4: esiste $\lambda \in \mathbb{R}^\Omega$ tale che

(a) se $X \in K$ allora $\sum_{\omega \in \Omega} \lambda(\omega) X(\omega) > 0$;

(b) se ρ è un processo predicibile allora $\sum_{\omega \in \Omega} \lambda(\omega) \tilde{G}_N(\rho)(\omega) = 0$.

Da (a) otteniamo che $\lambda(\omega) > 0$ per ogni $\omega \in \Omega$. Infatti, il vettore X definito $\bar{X}(\bar{\omega}) = 1$ e $\bar{X}(\omega) = 0$ appartiene a K e la condizione (a) implica subito che deve essere $\lambda(\omega) > 0$. Definiamo allora la seguente misura di probabilità $\mathbb{P}^*$ su Ω :

$$\mathbb{P}^*(\{\omega\}) = c_\lambda^{-1} \cdot \lambda(\omega), \text{ dove } c_\lambda = \sum_{\omega' \in \Omega} \lambda(\omega').$$

Per ipotesi $\mathbb{P}^*(\{\omega\}) > 0$ per ogni ω , dunque $\mathbb{P}^*$ è equivalente a $\mathbb{P}$. Inoltre, usando (b), otteniamo

$$\mathbb{E}^*(\tilde{G}_N(\rho)) = \sum_{\omega \in \Omega} \tilde{G}_N(\rho)(\omega) \mathbb{P}^*(\{\omega\}) = c_\lambda^{-1} \sum_{\omega \in \Omega} \tilde{G}_N(\rho)(\omega) \lambda(\omega) = 0.$$

Riassumendo: abbiamo determinato una misura di probabilità $\mathbb{P}^*$ equivalente a $\mathbb{P}$ tale che (4.4) è vera; poiché abbiamo visto che (4.4) garantisce che $((\tilde{S}_n^1, \dots, \tilde{S}_n^d))_n$ è una $\mathbb{P}^*$ -martingala, la tesi è dimostrata. $\square$

Dunque, l'assenza di arbitraggio equivale all'esistenza di una misura $\mathbb{P}^*$ equivalente a $\mathbb{P}$ sotto la quale i prezzi scontati sono $\mathbb{P}^*$ -martingale. Le misure $\mathbb{P}^*$ che verificano tali proprietà sono dette *misure equivalenti di martingala*.

Osservazione 4.4.6. È utile osservare che, in assenza di arbitraggio, ogni strategia autofinanziante di valore finale ≥ 0 è necessariamente ammissibile, cioè autofinanziante e tale che il portafoglio associato è sempre non negativo. Infatti, se $V_N(\phi) \geq 0$, allora anche $\tilde{V}_N(\phi) = (1+r)^{-N} V_N(\phi) \geq 0$. Inoltre, in assenza di arbitraggio esiste $\mathbb{P}^* \sim \mathbb{P}$ tale che $(\tilde{S}_n)_n$ è una $\mathbb{P}^*$ -martingala, quindi anche $(\tilde{V}_n(\phi))_n$ è una $\mathbb{P}^*$ -martingala (si veda l'Osservazione 4.3.7). Allora, per ogni $n \leq N$,

$$\tilde{V}_n(\phi) = \mathbb{E}^*(\tilde{V}_N(\phi) \mid \mathcal{F}_n) \geq 0,$$

poiché la speranza condizionata di una v.a. ≥ 0 è sempre ≥ 0 . Dunque ϕ è anche ammissibile.

Esempio 4.4.7. [La parità call-put] Vediamo ora che, in assenza di arbitraggio, tra il prezzo dell'opzione put e quello della call (per uno stesso strike e per una stessa maturità) vale una relazione fissa, che permette di dedurre l'uno dall'altro.

Indichiamo con C_n e P_n rispettivamente il prezzo al tempo n di una call e di una put, scritte sul bene sottostante di valore $(S_n)_{n \leq N}$, aventi stessa maturità N e uguale prezzo di esercizio K . In assenza di arbitraggio, vale allora la seguente relazione

$$C_n - P_n = S_n - K (1+r)^{-(N-n)}, \quad (4.5)$$

nota con il nome di *formula di parità call-put*.

Per dimostrare la (4.5), supponiamo che il nostro mercato sia formato, oltre che dal titolo non rischioso, dai seguenti titoli: il bene sottostante le opzioni call e put, l'opzione call e l'opzione put. In altre parole, supponiamo che

$$S_n^1 = S_n, \quad S_n^2 = C_n, \quad S_n^3 = P_n, \quad n = 0, 1, \dots, N.$$

Se il mercato è privo di arbitraggio, allora $((\tilde{S}_n^1, \tilde{S}_n^2, \tilde{S}_n^3))_n$ è una $\mathbb{P}^*$ -martingala, dunque anche $(\tilde{S}_n^2 - \tilde{S}_n^3)_n \equiv ((1+r)^{-n}(C_n - P_n))_n$ è una $\mathbb{P}^*$ -martingala:

$$(1+r)^{-n}(C_n - P_n) = \mathbb{E}^* \left((1+r)^{-N}(C_N - P_N) \mid \mathcal{F}_n \right).$$

Ora, a maturità N dev'essere necessariamente $C_N = (S_N - K)_+$ e $P_N = (K - S_N)_+$, dunque $C_N - P_N = S_N - K$ e si ha

$$(1+r)^{-n}(C_n - P_n) = \mathbb{E}^* \left((1+r)^{-N}(S_N - K) \mid \mathcal{F}_n \right).$$

Ora, poiché $(\tilde{S}_n)_n$ è anch'essa una $\mathbb{P}^*$ -martingala, $\mathbb{E}^*((1+r)^{-N}S_N \mid \mathcal{F}_n) = (1+r)^{-n}S_n$, dunque

$$(1+r)^{-n}(C_n - P_n) = (1+r)^{-n}S_n - (1+r)^{-N}K,$$

da cui segue immediatamente la (4.5).

La formula di parità (4.5) si può anche dimostrare direttamente usando argomenti di arbitraggio: se la (4.5) fosse falsa, cioè se si avesse

$$C_n - P_n > S_n - K(1+r)^{-(N-n)} \quad \text{oppure} \quad C_n - P_n < S_n - K(1+r)^{-(N-n)},$$

allora ci sarebbe arbitraggio.

Cominciamo col supporre $C_n - P_n > S_n - K(1+r)^{-(N-n)}$. Potremmo allora, all'istante n , vendere una call e comprare una quota di bene sottostante ed una put. Al netto, il bilancio dell'operazione è

$$C_n - P_n - S_n.$$

Se questa quantità è positiva la investiamo nel titolo senza rischio. Se è negativa ciò significa che abbiamo preso in prestito la quantità corrispondente. In entrambi i casi questo capitale all'istante N vale

$$(C_n - P_n - S_n)(1+r)^{N-n}.$$

Ora, all'istante N ,

- se $S_N > K$: la call è esercitata, dunque cediamo l'unità di sottostante che possediamo e riceviamo un compenso pari a K (la put non ha valore, dato che il prezzo di esercizio è inferiore a quello di mercato). Dunque il bilancio finale dell'operazione è

$$\begin{aligned} & K + (C_n - P_n - S_n)(1+r)^{N-n} \\ &= (1+r)^{N-n}(C_n - P_n - S_n + K(1+r)^{-(N-n)}) > 0; \end{aligned}$$

- se $S_T \leq K$: esercitiamo la put, cedendo l'unità di sottostante al prezzo K (la call non ha valore). Ci ritroviamo quindi con un capitale uguale ancora a

$$\begin{aligned} & K + (C_n - P_n - S_n)(1+r)^{N-n} \\ &= (1+r)^{N-n}(C_n - P_n - S_n + K(1+r)^{-(N-n)}) > 0; \end{aligned}$$

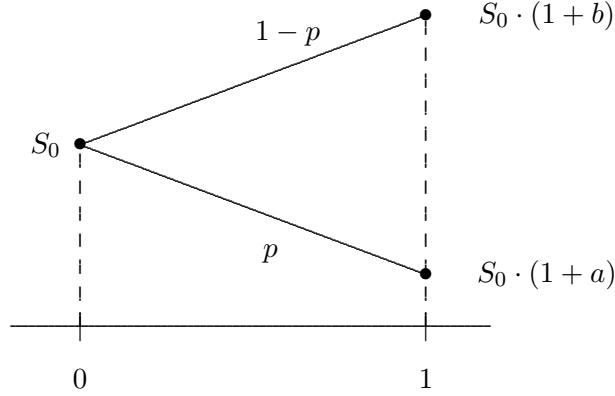


Figura 4.1: Il modello binomiale ad un periodo

Dunque, in ogni caso, all'istante N si ha un profitto sicuro. Queste strategie sono autofinanziate e di valore iniziale nullo. Se non ci fosse arbitraggio, sarebbero anche ammissibili per l'Osservazione 4.4.6. Avremmo dunque costruito una strategia di arbitraggio. Dunque, supponendo l'assenza di arbitraggio, non può essere $C_n - P_n > S_n - K(1+r)^{-(N-n)}$. In maniera simile si mostra che l'assenza di arbitraggio implica che non può essere $C_n - P_n < S_n - K(1+r)^{-(N-n)}$.

Esempio 4.4.8. [Modello binomiale a un periodo: mancanza di arbitraggio] Consideriamo il modello più semplice. Supponiamo $d = 1$ (cioè un solo attivo) e che vi sia un solo periodo temporale, durante il quale il prezzo dell'attivo possa crescere di un fattore $1+a$ oppure $1+b$, $-1 < a < b$, cioè (detto S il prezzo del titolo rischioso):

$$S_1 = \begin{cases} S_0 \cdot (1+a) & \text{con probabilità } = p \\ S_0 \cdot (1+b) & \text{con probabilità } = 1-p \end{cases}$$

dove S_0 (deterministico) è il prezzo iniziale (si veda la Figura 4.1). Si tratta di un modello senza arbitraggio?

Possiamo supporre $\Omega = \{1+a, 1+b\}$, $\mathcal{F}_1 = \mathcal{F} = \mathcal{P}(\Omega)$ e $S_1(\omega) = S_0 \cdot \omega$. Ovviamente, $\mathbb{P}(\{1+a\}) = p = 1 - \mathbb{P}(\{1+b\})$, dunque il modello verifica le condizioni che abbiamo sempre richiesto se supponiamo $p \in (0, 1)$. Una misura di probabilità $\mathbb{P}^* \sim \mathbb{P}$ dev'essere tale che $\mathbb{P}^*(\{1+a\}) = p^* = 1 - \mathbb{P}^*(\{1+b\})$, con $p^* \in (0, 1)$. Grazie al Teorema 4.4.5, possiamo dire che l'assenza di arbitraggio equivale a determinare, se esiste, p^* tale che il prezzo scontato dell'attivo di base sia una martingala, cioè

$$\mathbb{E}^*((1+r)^{-1}S_1 | \mathcal{F}_0) = S_0 \quad (4.6)$$

Poiché $\mathcal{F}_0$ è la σ -algebra banale, fare la speranza condizionale è lo stesso

che fare la speranza matematica. Si ha

$$\mathbb{E}^*((1+r)^{-1}S_1) = (1+r)^{-1}(p^*S_0(1+a) + (1-p^*)S_0(1+b))$$

Dunque perché (4.6) sia soddisfatta, deve essere

$$p^*(1+a) + (1-p^*)(1+b) = 1+r$$

ovvero

$$p^* = \frac{b-r}{b-a}.$$

Questo valore di p^* risulta > 0 se $b > r$ e < 1 se $a < r$. Dunque $0 < p^* < 1$ se e solo se $a < r < b$.

Riassumendo, possiamo dire che questo (semplice) modello è privo di arbitraggio se e solo se $a < r < b$, e in tal caso $\mathbb{P}^*$ è la probabilità definita tramite $\mathbb{P}^*({1+a}) = p^* = 1 - \mathbb{P}^*({1+b})$, con $p^* = (b-r)/(b-a)$. Osserviamo che esiste un solo valore p^* che rende (4.6) vera, cioè che determina una misura di martingala equivalente $\mathbb{P}^*$. Vedremo (cfr Esempio 4.5.6) che tale proprietà sarà particolarmente utile (dà infatti quella che chiameremo *completezza del mercato*).

4.5 Mercati completi e prezzo delle opzioni europee

Abbiamo già parlato di opzioni call e put europee. Ad esempio, una call europea sul bene sottostante di prezzo S^1 dà il diritto, ma non l'obbligo, di acquistare il bene alla data N (maturità) ad un prezzo (di esercizio) K fissato oggi ($n = 0$). Il valore della call è dato da $(S_N^1 - K)_+$. Il valore di una put con uguale maturità e prezzo di esercizio, è invece $(K - S_N)_+$. Esistono opzioni più generali che le call e le put. Per citarne solo alcune, un'opzione digital europea è un'opzione di valore

$$1_{\{S_N^1 \geq K\}},$$

con N la maturità e K lo strike: il valore a maturità è 1 se si osserva $S_N^1 \geq K$ e nullo in caso contrario.

Un'opzione call sulla media aritmetica dei titoli di prezzo $S^1, \dots, S^d$ ha valore

$$(\bar{S}_N - K)_+, \quad \text{essendo } \bar{S}_N = \frac{1}{d} \sum_{i=1}^d S_N^i.$$

In questi esempi, il valore a maturità dipende solo da $(S_N^1, \dots, S_N^d)$, cioè dal valore dei prezzi a maturità. Esistono anche opzioni più complicate, il cui valore è funzione di tutta l'evoluzione (traiettoria) del processo dei prezzi

fino a maturità, cioè di $((S_n^1, \dots, S_n^d))_{n \leq N}$. Questo tipo di opzioni è detto *path dependent*, cioè dipendente dalla traiettoria. Un esempio tipico sono le opzioni asiatiche: un'opzione asiatica call con maturità N e prezzo di esercizio K sul sottostante di prezzo S^1 ha valore

$$\left(\frac{1}{N+1} \sum_{j=0}^N S_j^1 - K \right)_+.$$

Dunque, si tratta di un'opzione call sulla media temporale del prezzo del sottostante $\frac{1}{N} \sum_{j=1}^N S_j^1$. Un altro tipo di opzioni path dependent è dato dalle opzioni con barriere. Ad esempio, una opzione call up-and-in sul primo bene ha valore

$$(S_N^1 - K)_+ 1_{\{\text{esiste } n \leq N \text{ tale che } S_n^1 \geq U\}}.$$

Qui, oltre alla maturità N e al prezzo di esercizio K , occorre specificare la quantità U : essa è contrattualmente specificata (quindi, deterministica) e rappresenta una barriera superiore. Una call up-and-in è quindi una normale call purché però esista $n \leq N$ tale che $S_n \geq U$, cioè purché il processo S^1 tocchi la barriera U entro l'istante di maturità N . In caso contrario, l'opzione ha valore nullo. Dunque, “up” significa che si lavora con una barriera superiore e “in” significa che l'opzione si attiva se e solo se la barriera viene toccata entro la maturità. Le altre opzioni con barriera sono quindi up-and-out, down-and-in, down-and-out, e per ciascuna si può proporre sia una call che una put, ma anche un'opzione diversa (una digital, un'asiatica etc.). Le cose si possono anche complicare con la doppia barriera, introducendo cioè due barriere U (superiore) e L (inferiore).

Più in generale, possiamo dare la seguente

Definizione 4.5.1. Un'opzione europea di maturità N è caratterizzata dal suo valore h , detto anche *payoff*, che si suppone essere una v.a. non negativa e $\mathcal{F}_N$ -misurabile.

A titolo di esempio, scriviamo il payoff h delle opzioni sopra riportate:

$$\begin{aligned} \text{call:} \quad h &= (S_N^1 - K)_+; \\ \text{put:} \quad h &= (K - S_N^1)_+; \\ \text{digital:} \quad h &= 1_{S_N^1 \geq K}; \\ \text{call su media aritmetica:} \quad h &= \left(\frac{1}{d} \sum_{i=1}^d S_N^i - K \right)_+; \\ \text{call asiatica:} \quad h &= \left(\frac{1}{N+1} \sum_{j=0}^N S_j^1 - K \right)_+; \\ \text{call up-and-in:} \quad h &= (S_N^1 - K)_+ 1_{\{\text{esiste } n \leq N \text{ tale che } S_n \geq U\}}; \\ \text{put up-and-out:} \quad h &= (K - S_N^1)_+ 1_{\{\text{per ogni } n \leq N, S_n < U\}}. \end{aligned}$$

In ogni caso, h è una v.a. non negativa e $\mathcal{F}_N$ -misurabile, perché è funzione dei prezzi fino a maturità N e quest'ultimi sono $\mathcal{F}_N$ -misurabili.

L'obiettivo ora è quello di trovare il giusto prezzo delle opzioni. A tale scopo, di fondamentale importanza è la seguente

Definizione 4.5.2. Un'opzione europea h si dice *replicabile* se esiste una strategia ammissibile ϕ tale che $V_N(\phi) = h$, cioè il cui valore finale del portafoglio è pari al payoff dell'opzione.

Osservazione 4.5.3. Se il mercato è privo di arbitraggio, perché un'opzione europea sia replicabile basta richiedere $V_N(\phi) = h$ per qualche strategia autofinanziante ϕ . In sostanza, in assenza di arbitraggio se $V_N(\phi) = h$ per una strategia autofinanziante ϕ allora ϕ è ammissibile. Infatti, ripetendo le argomentazioni dell'Osservazione 4.4.6, in assenza di arbitraggio $(\tilde{V}_n(\phi))_n$ è una $\mathbb{P}^*$ -martingala, per ogni misura equivalente di martingala $\mathbb{P}^*$ e per ogni strategia autofinanziante ϕ . Quindi, se $V_N(\phi) = h \geq 0$ allora

$$\tilde{V}_n(\phi) = \mathbb{E}^*(\tilde{V}_N(\phi) | \mathcal{F}_n) = \mathbb{E}^*((1+r)^{-N} h | \mathcal{F}_n) \geq 0$$

perché $(1+r)^{-N} h \geq 0$. Quindi anche $\tilde{V}_n(\phi) \geq 0$ e $V_n(\phi) \geq 0$ per ogni n e dunque ϕ è anche ammissibile.

Definizione 4.5.4. Un mercato si dice *completo* se ogni opzione è replicabile.

L'ipotesi di completezza del mercato è piuttosto restrittiva ed inoltre, a differenza della nozione di arbitraggio, non ha un chiaro significato finanziario. Tuttavia, studieremo nel prosieguo mercati completi, perché questa richiesta garantisce lo sviluppo di una teoria matematica piuttosto semplice che consente di studiare e calcolare il prezzo e la copertura delle opzioni (e questa è senz'altro una buona motivazione!).

Il teorema che segue, noto come il *secondo teorema fondamentale dell'asset pricing*, caratterizza il concetto di completezza del mercato quando ci si trovi in assenza di arbitraggio:

Teorema 4.5.5. *Un mercato privo di arbitraggio è completo se e solo se esiste un'unica misura di martingala equivalente $\mathbb{P}^*$, ovvero un'unica misura di probabilità $\mathbb{P}^*$ equivalente a $\mathbb{P}$ e tale che i prezzi scontati dei titoli rischiosi sono $\mathbb{P}^*$ -martingale.*

Dimostrazione. Prima parte: assumiamo che il mercato sia privo di arbitraggio e completo e dimostriamo che esiste un'unica misura equivalente di martingala.

Il Teorema 4.4.5 garantisce l'esistenza di una misura di martingala equivalente. Supponiamo ne esistano due: $\mathbb{P}_1$ e $\mathbb{P}_2$. Indichiamo con $\mathbb{E}_i$ la media valutata sotto $\mathbb{P}_i$, $i = 1, 2$. L'Osservazione 4.4.6 garantisce che, per ogni strategia autofinanziante ϕ , il valore scontato del portafoglio $\tilde{V}_n(\phi)$ è una

$\mathbb{P}_i$ -martingala. In particolare, poiché la v.a. $V_0(\phi)$ è $\mathcal{F}_0$ -misurabile e dunque costante,

$$\mathbb{E}_1(\tilde{V}_N(\phi)) = V_0(\phi) = \mathbb{E}_2(\tilde{V}_N(\phi)).$$

Poiché il mercato è anche completo, per ogni v.a. $h \geq 0$ esiste una strategia autofinanziante ϕ tale che $\tilde{V}_N(\phi) = (1+r)^{-N} V_N(\phi) = (1+r)^{-N} h$. Dunque

$$\mathbb{E}_1((1+r)^{-N} h) = \mathbb{E}_2((1+r)^{-N} h)$$

per ogni v.a. h non negativa, che sia $\mathcal{F}_N$ -misurabile. Ora, poiché per ipotesi $\mathcal{F}_N = \mathcal{F}$, possiamo scegliere $h = (1+r)^N 1_A$ per ogni $A \in \mathcal{F}$, ottenendo

$$\mathbb{P}_1(A) = \mathbb{P}_2(A) \quad \text{per ogni } A \in \mathcal{F},$$

da cui segue che $\mathbb{P}_1 = \mathbb{P}_2$.

Seconda parte: supponiamo l'assenza di arbitraggio e, per assurdo, che il mercato non sia completo; dimostriamo che non esiste una sola misura di martingala equivalente $\mathbb{P}^*$.

Se il mercato non è completo, esiste almeno un'opzione di payoff h che non è replicabile. Inoltre, poiché il mercato è privo di arbitraggio, esiste una misura di martingala equivalente $\mathbb{P}^*$.

Al variare di $((\phi_n^1, \dots, \phi_n^d))_n$ processo predicibile su $\mathbb{R}^d$ e di V_0 v.a. $\mathcal{F}_0$ -misurabile, sia $\mathcal{U}$ l'insieme delle v.a. della forma

$$\tilde{U}_N(\phi^1, \dots, \phi^d, V_0) = V_0 + \sum_{j=1}^N \left(\phi_n^1 \Delta \tilde{S}_n^1 + \dots + \phi_n^d \Delta \tilde{S}_n^d \right).$$

Per ogni $((\phi_n^1, \dots, \phi_n^d))_n$ processo predicibile su $\mathbb{R}^d$ e per ogni V_0 v.a. $\mathcal{F}_0$ -misurabile, la Proposizione 4.3.6 garantisce l'esistenza di un processo predicibile $(\phi_n^0)_n$ tale che $\phi = ((\phi_n^0, \phi_n^1, \dots, \phi_n^d))_n$ è una strategia autofinanziante e $U_N(\phi^1, \dots, \phi^d, V_0) = 1/\beta_n \tilde{U}_N(\phi^1, \dots, \phi^d, V_0) = V_N(\phi)$. Abbiamo quindi le seguenti conseguenze.

- $h \cdot (1+r)^{-N} \notin \mathcal{U}$ e quindi $\mathcal{U}$ è un sottoinsieme stretto dell'insieme $\mathbb{R}^\Omega$ delle v.a. su $(\Omega, \mathcal{F})$. Infatti, per ipotesi h non è replicabile: non esiste alcuna strategia autofinanziante ϕ tale che $V_N(\phi) = h$ (si veda anche l'Osservazione 4.5.3). Poiché abbiamo visto che $V_N(\phi) = U_N(\phi^1, \dots, \phi^d, V_0)$, segue che h non è della forma $U_N(\phi^1, \dots, \phi^d, V_0)$, cioè $h \notin \mathcal{U}$.
- Poiché $(\tilde{V}_n(\phi))_n$ è una $\mathbb{P}^*$ -martingala (Osservazione 4.4.6), $\mathbb{E}^*(\tilde{V}_N(\phi)) = V_0$, quindi possiamo dire

$$\mathbb{E}^*(\tilde{U}) = V_0 \quad \text{per ogni } \tilde{U} = \tilde{U}_N(\phi^1, \dots, \phi^d, V_0) \in \mathcal{U}. \quad (4.7)$$

Ora, definiamo su $\mathbb{R}^\Omega$ il seguente prodotto scalare³

$$\mathbb{R}^\Omega \times \mathbb{R}^\Omega \ni (X, Y) \mapsto \mathbb{E}^*(XY),$$

dove $\mathbb{E}^*$ denota, al solito, l'aspettazione sotto $\mathbb{P}^*$. Allora, essendo $\widetilde{\mathcal{U}}$ un sottospazio stretto di $\mathbb{R}^\Omega$, esiste un elemento $\bar{X} \in \mathbb{R}^\Omega$ non nullo ortogonale a $\widetilde{\mathcal{U}}$:

$$\mathbb{E}^*(\bar{X} \tilde{U}) = 0 \text{ per ogni } \tilde{U} \in \widetilde{\mathcal{U}}. \quad (4.8)$$

Poiché in particolare $\tilde{U} = 1 \in \widetilde{\mathcal{U}}$ (basta prendere $\phi_n^i = 0$ per ogni $i = 1, \dots, d$ e $n = 0, 1, \dots, N$ e $V_0 = 1$), da (4.8) si ha

$$\mathbb{E}^*(\bar{X}) = 0. \quad (4.9)$$

Ora, definiamo

$$\mathbb{P}^{**}(\{\omega\}) = \left(1 + \frac{\bar{X}(\omega)}{2\|\bar{X}\|_\infty}\right) \mathbb{P}^*(\{\omega\}), \quad \omega \in \Omega,$$

dove si è posto $\|\bar{X}\|_\infty = \sup_{\omega \in \Omega} |\bar{X}(\omega)|$, e osserviamo che:

1. $\mathbb{P}^{**}$ è una misura di probabilità: $\mathbb{P}^{**}(\{\omega\}) > 0$ per ogni ω (perché $\mathbb{P}^*(\{\omega\}) > 0$ e $1 + \bar{X}(\omega)/(2\|\bar{X}\|_\infty) > 0$) e, ricordando la (4.9),

$$\begin{aligned} \sum_{\omega \in \Omega} \mathbb{P}^{**}(\{\omega\}) &= \sum_{\omega \in \Omega} \mathbb{P}^*(\{\omega\}) + \frac{1}{2\|\bar{X}\|_\infty} \sum_{\omega \in \Omega} \bar{X}(\omega) \mathbb{P}^*(\{\omega\}) \\ &= 1 + \frac{1}{2\|\bar{X}\|_\infty} \mathbb{E}^*(\bar{X}) = 1; \end{aligned}$$

2. $\mathbb{P}^{**}$ è equivalente a $\mathbb{P}^*$: abbiamo visto in 1. che $\mathbb{P}^{**}(\{\omega\}) > 0$ per ogni $\omega \in \Omega$;
3. $\mathbb{P}^{**}$ è diversa da $\mathbb{P}^*$, perché $\bar{X}$ è non nulla;
4. posto $\mathbb{E}^{**}$ la media sotto $\mathbb{P}^{**}$, per ogni $\tilde{U} \in \widetilde{\mathcal{U}}$ si ha

$$\mathbb{E}^{**}(\tilde{U}) = V_0 \text{ per ogni } \tilde{U} = \tilde{U}_N(\phi^1, \dots, \phi^d, V_0) \in \widetilde{\mathcal{U}}.$$

Infatti, tenendo presente (4.7) e (4.8),

$$\begin{aligned} \mathbb{E}^{**}(\tilde{U}) &= \sum_{\omega \in \Omega} \tilde{U}(\omega) \mathbb{P}^{**}(\{\omega\}) \\ &= \sum_{\omega \in \Omega} \tilde{U}(\omega) \mathbb{P}^*(\{\omega\}) + \frac{1}{2\|\bar{X}\|_\infty} \sum_{\omega \in \Omega} \tilde{U}(\omega) \bar{X}(\omega) \mathbb{P}^*(\{\omega\}) \\ &= \mathbb{E}^*(\tilde{U}) + \frac{1}{2\|\bar{X}\|_\infty} \mathbb{E}^*(\tilde{U} \bar{X}) = V_0. \end{aligned}$$

³La verifica che si tratta effettivamente di un prodotto scalare è immediata.

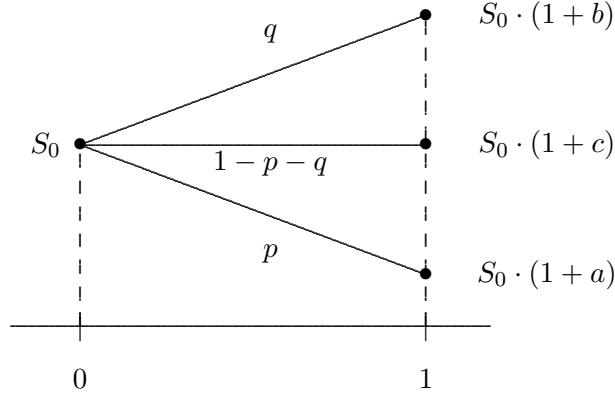


Figura 4.2: Il modello trinomiale ad un periodo

Riassumendo: abbiamo determinato una misura di probabilità $\mathbb{P}^{**}$ diversa da $\mathbb{P}^*$ ma equivalente a $\mathbb{P}^*$ tale che (si prenda nel punto 4. $\tilde{U} = \tilde{U}_N(\phi^1, \dots, \phi^d, V_0) \in \tilde{\mathcal{U}}$ con $V_0 = 0$)

$$\mathbb{E}^{**} \left(\sum_{j=1}^N \left(\phi_n^1 \Delta \tilde{S}_n^1 + \dots + \phi_n^d \Delta \tilde{S}_n^d \right) \right) = 0$$

per ogni processo predicibile $((\phi_n^1, \dots, \phi_n^d))_n$. La Proposizione 3.2.7 assicura che $(\tilde{S}_n^1, \dots, \tilde{S}_n^d)$ è una $\mathbb{P}^{**}$ -martingala. Ciò prova che esistono due diverse misure di martingala equivalenti, e la dimostrazione è completata. $\square$

Esempio 4.5.6. [Modello binomiale a un periodo: completezza] Riprendiamo l'Esempio 4.4.8: supponiamo $d = 1$ (cioè un solo attivo) e che vi sia un solo periodo temporale, durante il quale il prezzo dell'attivo può crescere di due fattori. Usando le notazioni dell'Esempio 4.4.8, si sta supponendo

$$S_1 = \begin{cases} S_0 \cdot (1 + a) & \text{con probabilità } = p \\ S_0 \cdot (1 + b) & \text{con probabilità } = 1 - p \end{cases}$$

dove S_0 (deterministico) è il prezzo iniziale. Si tratta di un modello completo? La risposta è sì: come sottolineato nell'Esempio 4.4.8, p^* dà l'unica misura di martingala equivalente.

Esempio 4.5.7. [Modello trinomiale a un periodo] Proviamo a complicare il modello binomiale, inserendo un'altra possibilità di salto per il prezzo dell'attivo. Supponiamo quindi che $d = 1$ e che vi sia un solo periodo

temporale durante il quale il prezzo dell'attivo possa crescere di tre fattori (si veda la Figura 4.2):

$$S_1 = \begin{cases} S_0 \cdot (1 + a) & \text{con probabilità } = p \\ S_0 \cdot (1 + c) & \text{con probabilità } = 1 - p - q \\ S_0 \cdot (1 + b) & \text{con probabilità } = q \end{cases}$$

dove S_0 (deterministico) è il prezzo iniziale, $p, q > 0$ e tali che $p + q < 1$. Problema: è un modello senza arbitraggio? È completo?

Come mostrato in figura, supponiamo a, b, c tali che $-1 < a < c < b$ e, analogamente al modello binomiale, è naturale prendere $\Omega = \{1+a, 1+b, 1+c\}$. Procediamo ancora come nel modello binomiale: occorre determinare una coppia p^* e q^* tali che $p^*, q^* > 0$ e $p^* + q^* < 1$ e tali che

$$\mathbb{E}^*((1+r)^{-1}S_1) = S_0,$$

essendo $\mathbb{E}^*$ la media sotto la probabilità $\mathbb{P}^*$ definita da: $\mathbb{P}^*(\{1+a\}) = p^*$, $\mathbb{P}^*(\{1+b\}) = q^*$ e $\mathbb{P}^*(\{1+c\}) = 1 - p^* - q^*$. Ora, $\mathbb{E}^*(S_1) = S_0(1+a)p^* + S_0(1+b)q^* + S_0(1+c)(1-p^*-q^*)$, dunque dev'essere

$$(1+a)p^* + (1+b)q^* + (1+c)(1-p^*-q^*) = 1+r,$$

o equivalentemente

$$p^*a + q^*b + (1-p^*-q^*)c = r.$$

Si tratta di una equazione lineare nelle incognite p^* e q^* soggette ai vincoli $p^*, q^* > 0$ e $p^* + q^* < 1$. In particolare, si tratta di una combinazione lineare convessa dei punti a, b e c : al variare dei coefficienti p^*, q^* e $1 - p^* - q^*$ si ottiene tutto l'intervallo di estremi (esclusi) $\min(a, b, c) = a$ e $\max(a, b, c) = b$. Dunque, tale equazione ha soluzione se e solo se $r \in (a, b)$: questo modello è privo di arbitraggio se e solo se $a < r < b$. In tal caso, possiamo anche dedurre che non è completo: se infatti esiste una soluzione dell'equazione sopra scritta, allora ne esistono infinite. In altre parole, esistono infinite misure di martingala equivalenti, cioè il modello non è completo.

Supponiamo ora che il nostro modello di mercato sia privo di arbitraggio e completo, dunque esiste un'unica misura di martingala equivalente $\mathbb{P}^*$. Sia h un'opzione, con maturità N , cioè una v.a. non negativa e $\mathcal{F}_N$ -misurabile. Allora esiste una strategia autofinanziante (ammissibile) ϕ che replica l'opzione, cioè tale che h è il valore finale del portafoglio associato a ϕ :

$$V_N(\phi) = h.$$

Inoltre, il valore scontato $\tilde{V}_n(\phi)$ di $V_N(\phi)$ è una $\mathbb{P}^*$ -martingala. Di conseguenza,

$$V_0(\phi) = \mathbb{E}^*(\tilde{V}_N(\phi)) = \mathbb{E}^*((1+r)^{-N} h),$$

e più in generale $\tilde{V}_n(\phi) = \mathbb{E}^*(\tilde{V}_N(\phi) | \mathcal{F}_n)$, che si può riscrivere, per $n = 0, 1, \dots, N$,

$$V_n(\phi) = \mathbb{E}^*\left((1+r)^{-(N-n)} h | \mathcal{F}_n\right). \quad (4.10)$$

La strategia ϕ prende il nome di *strategia replicante l'opzione* ed il portafoglio $V_n(\phi)$ ad essa associata è detto *portafoglio replicante*.

La (4.10) dice che il valore del portafoglio replicante è univocamente determinato, in ogni istante, dalla maturità N e dal payoff h . Inoltre, la (4.10) quantifica la ricchezza che necessita avere al tempo n per avere (replicare) h a maturità N . È quindi naturale stabilire il *prezzo dell'opzione* al tempo n pari alla quantità definita in (4.10).

In altre parole, se all'istante iniziale il compratore dà al venditore un ammontare pari a

$$\mathbb{E}^*\left((1+r)^{-N} h\right)$$

(cioè, al valore in (4.10) con $n = 0$) e se il venditore segue la strategia replicante ϕ , allora al tempo N è in grado di generare una quantità di denaro pari ad h , ossia al valore dell'opzione. In altre parole, seguendo questa strategia il venditore è in grado di coprire esattamente l'opzione.

Chi compra l'opzione deve quindi spendere un ammontare di denaro pari a $\mathbb{E}^*((1+r)^{-N} h)$. Colui che invece vende l'opzione, per coprirsi da eventuali perdite, deve possedere un portafoglio dato da $\mathbb{E}^*((1+r)^{-(N-n)} h | \mathcal{F}_n)$, il quale, a sua volta, si costruisce seguendo una strategia di copertura ϕ .

Se si interviene nel contratto ad un istante n qualsiasi, facendo considerazioni analoghe possiamo dire che il giusto prezzo dell'opzione al tempo n è

$$\mathbb{E}^*\left((1+r)^{-(N-n)} h \mid \mathcal{F}_n\right)$$

Ora, la strategia di copertura è un oggetto fondamentale. Al momento, sappiamo che esiste ma fino ad ora non è stato possibile determinarla, né caratterizzarla in qualche modo. Comunque, nel prossimo paragrafo introdurremo un modello (discreto) in cui sarà possibile calcolare esplicitamente la copertura delle opzioni call e put.

Infine, osserviamo che il calcolo del prezzo di un'opzione coinvolge esclusivamente la misura di martingala equivalente $\mathbb{P}^*$, e non quella originaria $\mathbb{P}$. Avremmo quindi potuto sviluppare la teoria esclusivamente considerando uno spazio misurabile $(\Omega, \mathcal{F})$ su cui sia definita una filtrazione $(\mathcal{F}_n)_n$. In altre parole, si può non citare mai la misura $\mathbb{P}$, che poi rappresenta la vera probabilità sul mercato finanziario.

Concludiamo il paragrafo ritrovando (cfr Esempio 4.4.7), a partire dalla formula del prezzo, la formula di parità call/put: se C_n e P_n sono rispettivamente il prezzo al tempo n di una call e di una put con uguale maturità N e stesso prezzo di esercizio K , allora, in un mercato privo di arbitraggio

e completo, vale la seguente relazione:

$$C_n - P_n = S_n^1 - K (1 + r)^{-(N-n)}, \quad (4.11)$$

Infatti, dalla (4.10) segue che

$$\begin{aligned} C_n &= \mathbb{E}^*((1+r)^{-(N-n)}(S_N^1 - K)_+ | \mathcal{F}_n) \quad \text{e} \\ P_n &= \mathbb{E}*((1+r)^{-(N-n)}(K - S_N^1)_+ | \mathcal{F}_n), \end{aligned}$$

dove $\mathbb{P}^*$ denota l'unica misura equivalente di martingala. Allora,

$$\begin{aligned} C_n - P_n &= (1+r)^{-(N-n)} \mathbb{E}^*[(S_N^1 - K)_+ - (K - S_N^1)_+ | \mathcal{F}_n] = \\ &= (1+r)^{-(N-n)} \mathbb{E}^*(S_N^1 - K | \mathcal{F}_n) = \\ &= (1+r)^{-(N-n)} ((1+r)^N \mathbb{E}^*(\tilde{S}_N^1 | \mathcal{F}_n) - K) = \\ &= (1+r)^{-(N-n)} ((1+r)^N \tilde{S}_n^1 - K) = S_n^1 - K (1+r)^{-(N-n)}, \end{aligned}$$

cioè la formula di parità.

4.6 Il modello di Cox, Ross e Rubinstein

In questo paragrafo studieremo il modello di Cox, Ross e Rubinstein, nel seguito scriveremo brevemente modello CRR, versione discreta del ben più famoso modello di Black e Scholes.

In questo modello, si suppone $d = 1$. Sul mercato sono quindi presenti due titoli: il titolo non rischioso, di prezzo $S_n^0 = (1+r)^n$, e quello rischioso, di prezzo S_n . Si suppone che S_n si evolva nel modo seguente: tra due istanti consecutivi, n e $n+1$, la variazione percentuale del prezzo assume due soli possibili valori, a e b , con $-1 < a < b$. In formule,

$$\begin{aligned} \frac{\Delta S_{n+1}}{S_n} &\text{ a valori in } \{a, b\}, \quad \text{o equivalentemente} \\ S_{n+1} &\text{ a valori in } \{S_n(1+a), S_n(1+b)\} \end{aligned}$$

Il valore iniziale S_0 è dato (dunque, deterministico). Se poniamo $T_0 = S_0$ e per $n \geq 1$, $T_n = S_n/S_{n-1}$, allora

$$T_n \text{ a valori in } \{(1+a), (1+b)\} \text{ ed inoltre } S_n = S_0 T_1 \cdots T_n.$$

Un modo di prendere l'insieme di tutti i possibili stati è allora quello di porre

$$\Omega = \{1+a, 1+b\}^N.$$

Ovviamente, per $\omega = (\omega_1, \dots, \omega_N) \in \Omega$, si ha

$$T_n(\omega) = \omega_n \quad \text{e} \quad S_n(\omega) = S_0 T_1(\omega) \cdots T_n(\omega) = S_0 \omega_1 \cdots \omega_n.$$

Prendiamo $\mathcal{F} = \mathcal{P}(\Omega)$, cosicché le funzioni fin qui definite sono variabili aleatorie. Costruiamo ora la filtrazione. Poniamo $\mathcal{F}_0 = \{\emptyset, \Omega\}$ e per $n \geq 1$ definiamo⁵

$$\mathcal{F}_n = \sigma(S_1, \dots, S_n) = \sigma(T_1, \dots, T_n).$$

Osserviamo che si ha⁶ $\mathcal{F}_N = \mathcal{F}$.

Riassumendo, abbiamo costruito uno spazio misurabile $(\Omega, \mathcal{F})$, una filtrazione $(\mathcal{F}_n)_{0 \leq n \leq N}$ e due processi $(T_n)_{0 \leq n \leq N}$ e $(S_n)_{0 \leq n \leq N}$ adattati alla filtrazione. Sottolineiamo che non supponiamo che le v.a. T_i siano indipendenti, né che siano equidistribuite.

Ora occorrerebbe definire una misura di probabilità $\mathbb{P}$ su $(\Omega, \mathcal{F})$. Come abbiamo già osservato alla fine del paragrafo precedente, allo scopo di prezzare opzioni sul sottostante S non interessa quale sia la misura di probabilità $\mathbb{P}$. Piuttosto, è importante conoscere la misura di martingala equivalente $\mathbb{P}^*$, una volta che avremo dimostrato che non c'è arbitraggio e che il mercato è completo. Quindi, per il momento non specifichiamo chi sia $\mathbb{P}$. Diciamo solo che $\mathbb{P} \in \mathcal{Q}$, dove

$$\mathcal{Q} = \{\mathbb{Q} : \mathbb{Q} \text{ è una probabilità su } (\Omega, \mathcal{F}) \text{ e } \mathbb{Q}(\{\omega\}) > 0 \text{ per ogni } \omega \in \Omega\}$$

Presa $\mathbb{P} \in \mathcal{Q}$, ovviamente è possibile calcolare la distribuzione congiunta di $(T_1, \dots, T_N)$ e di $(S_1, \dots, S_N)$. Infatti,

$$\begin{aligned} \mathbb{P}(T_1 = t_1, \dots, T_N = t_N) &= \mathbb{P}(\{\omega\}), \\ \text{con } \omega &= (t_1, \dots, t_N) \in \{1+a, 1+b\}^N, \end{aligned} \quad (4.12)$$

e

$$\begin{aligned} \mathbb{P}(S_1 = s_1, \dots, S_N = s_N) &= \mathbb{P}\left(T_1 = \frac{s_1}{s_0}, \dots, T_N = \frac{s_N}{s_{N-1}}\right) = \mathbb{P}(\{\omega\}) \\ \text{con } \omega &= (s_1/s_0, \dots, s_N/s_{N-1}) \in \{1+a, 1+b\}^N. \end{aligned} \quad (4.13)$$

La figura 4.6 mostra l'insieme di tutte le possibili traiettorie $(S_n)_{n \leq N}$, con $N = 3$. Si noti che esse formano un albero; un cammino percorribile sull'albero rappresenta ovviamente una possibile traiettoria.

4.6.1 Assenza di arbitraggio e completezza

Abbiamo già osservato che è di fondamentale importanza trovare quelle che abbiamo chiamato misure equivalenti di martingala. Nel lemma che segue, ne diamo una caratterizzazione.

⁵Osserviamo che effettivamente $\sigma(S_1, \dots, S_n) = \sigma(T_1, \dots, T_n)$. Tale asserzione segue dalla Proposizione 2.1.14 o anche, e immediatamente, dall'Esercizio 2.11. Infatti, poiché $(T_1, \dots, T_n) = (S_1/S_0, \dots, S_n/S_{n-1})$ è funzione misurabile di $(S_1, \dots, S_n)$, si ha $\sigma(T_1, \dots, T_n) \subset \sigma(S_1, \dots, S_n)$. Analogamente, da $(S_1, \dots, S_n) = (S_0 T_1, S_0 T_1 T_2, \dots, S_0 T_1 \cdots T_n)$ segue che $\sigma(S_1, \dots, S_n) \subset \sigma(T_1, \dots, T_n)$.

⁶Infatti, $\mathcal{F}_N = \sigma(T_1, \dots, T_N) = \sigma(\{\omega\} : \omega \in \Omega) = \mathcal{P}(\Omega) = \mathcal{F}$.

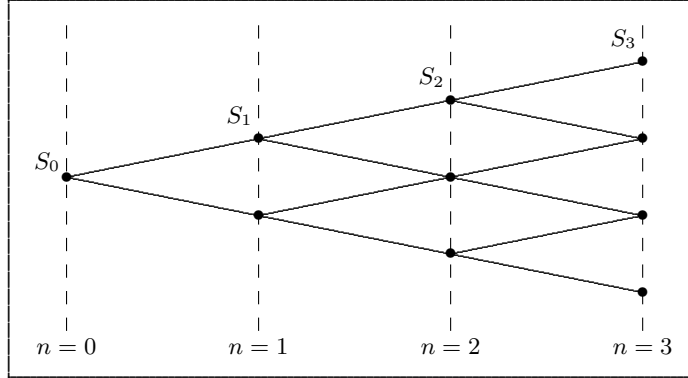


Figura 4.4 L'albero formato da tutti i possibili cammini di $(S_n)_{0 \leq n \leq 3}$.

Lemma 4.6.1. *Sia $\mathbb{P}^* \in \mathcal{Q}$. Il processo di prezzo scontato $(\tilde{S}_n)_n$ è una $\mathbb{P}^*$ -martingala se e solo se*

$$\mathbb{E}^*(T_{n+1} | \mathcal{F}_n) = 1 + r, \quad \text{per ogni } n = 0, 1, \dots, N-1.$$

Dimostrazione. $(\tilde{S}_n)_n = ((1+r)^{-n} S_n)_n$ è una $\mathbb{P}^*$ -martingala se e solo se

$$\mathbb{E}^*((1+r)^{-(n+1)} S_{n+1} | \mathcal{F}_n) = (1+r)^{-n} S_n, \quad \text{per ogni } n = 0, 1, \dots, N-1,$$

cioè

$$\frac{1}{S_n} \mathbb{E}^*(S_{n+1} | \mathcal{F}_n) = 1 + r, \quad \text{per ogni } n = 0, 1, \dots, N-1.$$

Ora, poiché S_n è $\mathcal{F}_n$ -misurabile, si ha

$$\frac{1}{S_n} \mathbb{E}^*(S_{n+1} | \mathcal{F}_n) = \mathbb{E}^*\left(\frac{S_{n+1}}{S_n} | \mathcal{F}_n\right) = \mathbb{E}^*(T_{n+1} | \mathcal{F}_n),$$

da cui la tesi. □

Vediamo ora quali sono le condizioni da imporre perché il mercato descritto da questo modello sia privo di arbitraggio.

Proposizione 4.6.2. *Perché il modello CRR sia privo di arbitraggio è necessario che a e b siano tali che*

$$a < r < b.$$

Dimostrazione. Supponiamo che il mercato sia privo di arbitraggio. Allora (cfr. Teorema 4.4.5) esiste una misura di probabilità $\mathbb{P}^*$ equivalente a $\mathbb{P}$, cioè $\mathbb{P}^* \in \mathcal{Q}$, tale che il processo del prezzo scontato del titolo rischioso è una $\mathbb{P}^*$ -martingala. Quindi, per il Lemma 4.6.1, dev'essere $\mathbb{E}^*(T_{n+1} | \mathcal{F}_n) = 1 + r$, da cui segue che

$$\mathbb{E}^*(T_{n+1}) = \mathbb{E}^*[\mathbb{E}^*(T_{n+1} | \mathcal{F}_n)] = 1 + r.$$

Ora, $\mathbb{E}^*(T_{n+1}) = (1+a)p_a^* + (1+b)p_b^*$, dove $p_a^* = \mathbb{P}^*(T_{n+1} = 1+a) = 1-p_b^* = 1 - \mathbb{P}^*(T_{n+1} = 1+b)$. Osserviamo che poiché $\mathbb{P}^* \in \mathcal{Q}$, dev'essere $p_a^* > 0$ e $p_b^* > 0$. Ma $(1+a)p_a^* + (1+b)p_b^*$ è una combinazione lineare convessa di $1+a$ e $1+b$, dunque assume valori tra il minimo, $1+a$, e il massimo, $1+b$. Inoltre gli estremi non possono essere raggiunti, perché $p_a^* = 1-p_b^* \in (0,1)$. Dunque deve essere

$$1+r = (1+a)p_a^* + (1+b)p_b^* \in (1+a, 1+b).$$

Dunque, $r \in (a, b)$.

Alternativamente, supponiamo per assurdo $r \notin (a, b)$ e mostriamo che è possibile costruire una strategia di arbitraggio.

Supponiamo $r \leq a$. In questo caso, al tempo 0 prendiamo in prestito un ammontare pari a S_0 , così da poter comprare il titolo rischioso, e al tempo N restituiamo il denaro preso in prestito. Questa strategia (ϕ_n^0, ϕ_n) si può riassumere così:

$$\phi_n^0 = -S_0 \quad \phi_n = 1.$$

In ogni istante n , il valore del nostro investimento è

$$V_n = S_n - S_0(1+r)^n.$$

In particolare, $V_0 = 0$ e $V_N = S_N - S_0(1+r)^N$. Ora, poiché $S_n \geq S_0(1+a)^n \geq S_0(1+r)^n$, $V_n \geq 0$ per ogni n , la strategia è ammissibile, ed inoltre⁷ $\mathbb{P}(V_N > 0) > 0$, cioè è una strategia di arbitraggio. Dunque dev'essere $r > a$. Analogamente si prova che⁸ $r < b$, da cui l'assurdo.

□

D'ora in poi supporremo a e b fissati in modo che si abbia $a < r < b$. Dunque, d'ora in poi supporremo che il modello CRR sia privo di arbitraggio.

Proposizione 4.6.3. *Supponiamo $a < r < b$ e poniamo*

$$p = \frac{b-r}{b-a}.$$

Allora, se $\mathbb{P}^ \in \mathcal{Q}$, $(\tilde{S}_n)_n$ è una $\mathbb{P}^*$ -martingala se e solo se sotto $\mathbb{P}^*$ le v.a. $T_1, \dots, T_N$ sono i.i.d. e tali che*

$$\mathbb{P}^*(T_n = 1+a) = p = 1 - \mathbb{P}^*(T_n = 1+b).$$

⁷Infatti, se $\bar{\omega} = (1+b, \dots, 1+b)$ allora $V_N(\bar{\omega}) = S_0(1+b)^N - S_0(1+r)^N > 0$, quindi $\mathbb{P}(V_N > 0) \geq \mathbb{P}(\{\bar{\omega}\}) > 0$.

⁸Basta ripetere lo stesso ragionamento invertendo i ruoli al titolo rischioso e a quello non rischioso.

Dimostrazione. Supponiamo che sotto $\mathbb{P}^*$ le v.a. $T_1, \dots, T_N$ siano i.i.d. e tali che $\mathbb{P}^*(T_n = 1 + a) = p = 1 - \mathbb{P}^*(T_n = 1 + b)$. In tal caso, T_{n+1} è indipendente da $\sigma(T_1, \dots, T_n) = \mathcal{F}_n$, quindi

$$\mathbb{E}^*(T_{n+1} | \mathcal{F}_n) = \mathbb{E}^*(T_{n+1}) = (1 + a)p + (1 + b)(1 - p) = 1 + r$$

e dalla Proposizione 4.6.1 segue che $(\tilde{S}_n)_n$ è una $\mathbb{P}^*$ -martingala. Mostriamo che $\mathbb{P}^* \in \mathcal{Q}$: da (4.12), per ogni $\omega \in \Omega = \{1 + a, 1 + b\}^N$,

$$\begin{aligned} \mathbb{P}^*(\{\omega\}) &= \mathbb{P}^*(T_1 = \omega_1, \dots, T_N = \omega_N) = \prod_{n=1}^N \mathbb{P}^*(T_n = \omega_n) \\ &= p^k (1 - p)^{N-k} > 0 \end{aligned}$$

dove $k = \#\{n; \omega_n = 1 + a\}$, poiché $0 < p < 1$. Quindi $\mathbb{P}^* \in \mathcal{Q}$.

Viceversa, supponiamo che $(\tilde{S}_n)_n$ sia una $\mathbb{P}^*$ -martingala, con $\mathbb{P}^* \in \mathcal{Q}$. Allora, dalla Proposizione 4.6.1 segue che $\mathbb{E}^*(T_{n+1} | \mathcal{F}_n) = 1 + r$, quindi

$$\begin{aligned} 1 + r &= \mathbb{E}^*(T_{n+1} | \mathcal{F}_n) = \\ &= (1 + a) \mathbb{P}^*(T_{n+1} = 1 + a | \mathcal{F}_n) + (1 + b) \mathbb{P}^*(T_{n+1} = 1 + b | \mathcal{F}_n) = \\ &= (1 + a) p_n^* + (1 + b) (1 - p_n^*) \end{aligned}$$

dove abbiamo posto $p_n^* = \mathbb{P}^*(T_{n+1} = 1 + a | \mathcal{F}_n)$, da cui ricaviamo

$$p_n^* = p = \frac{b - r}{b - a}.$$

Ora,

$$\begin{aligned} \mathbb{P}^*(T_{n+1} = 1 + a) &= \mathbb{E}^*(1_{\{T_{n+1}=1+a\}}) = \mathbb{E}^*[\mathbb{E}^*(1_{\{T_{n+1}=1+a\}} | \mathcal{F}_n)] \\ &= \mathbb{E}^*[\mathbb{P}^*(T_{n+1} = 1 + a | \mathcal{F}_n)] = p, \end{aligned}$$

da cui segue che $\mathbb{P}^*(T_n = 1 + a) = p = 1 - \mathbb{P}^*(T_n = 1 + b)$ per ogni $n = 1, \dots, N$. Mostriamo ora che $T_1, \dots, T_N$ sono indipendenti. Poniamo $p_t = p$ se $t = 1 + a$ e $p_t = 1 - p$ quando $t = 1 + b$. Allora, per $(t_1, \dots, t_N) \in \{1 + a, 1 + b\}^N$, si ha

$$\begin{aligned} \mathbb{P}^*(T_1 = t_1, \dots, T_N = t_N) &= \mathbb{E}^*[\mathbb{E}^*(1_{\{T_1=t_1, \dots, T_N=t_N\}} | \mathcal{F}_{N-1})] \\ &= \mathbb{E}^*[1_{\{T_1=t_1, \dots, T_{N-1}=t_{N-1}\}} \mathbb{P}^*(T_N = t_N | \mathcal{F}_{N-1})] = \\ &= p_{t_N} \mathbb{E}^*(1_{\{T_1=t_1, \dots, T_{N-1}=t_{N-1}\}}) = \\ &= \dots = \\ &= p_{t_N} p_{t_{N-1}} \dots p_{t_1} = \mathbb{P}^*(T_1 = t_1) \dots \mathbb{P}^*(T_N = t_N) \end{aligned}$$

da cui segue che, sotto $\mathbb{P}^*$, le v.a. $T_1, \dots, T_N$ sono indipendenti. $\square$

Riassumendo, la Proposizione 4.6.2 garantisce che, per $a < r < b$, il mercato è privo di arbitraggio, grazie al primo teorema fondamentale dell'asset pricing, Teorema 4.4.5; la Proposizione 4.6.3 garantisce che la misura di martingala equivalente è unica, dunque per il secondo teorema fondamentale dell'asset pricing, Teorema 4.5.5 il modello CRR è completo. Dunque

Teorema 4.6.4. *Se a, b sono tali che $a < r < b$, il modello CRR è privo di arbitraggio e completo.*

4.6.2 Prezzo e copertura delle opzioni nel modello CRR

Vediamo in questo paragrafo come si possono calcolare i prezzi delle opzioni e determinare i portafogli di copertura nel modello CRR. Consideriamo per cominciare un'opzione il cui payoff sia una funzione di S_N , cioè

$$h = F(S_N) \quad (4.14)$$

Si tratta di un caso che copre gli esempi delle opzioni put e call, per le quali si ha $F(x) = (K - x)_+$ e $F(x) = (x - K)_+$ rispettivamente.

Per la (4.10), il prezzo dell'opzione (4.14) è dato da

$$C_n = (1 + r)^{-(N-n)} \mathbb{E}^*[F(S_N) | \mathcal{F}_n],$$

dove $\mathbb{E}^*$ è la media fatta sotto la misura di martingala equivalente $\mathbb{P}^*$. La Proposizione 4.6.3 assicura che, sotto $\mathbb{P}^*$, $T_1, \dots, T_N$ sono indipendenti, quindi $T_{n+1} \cdots T_N$ è indipendente da $\mathcal{F}_n$. Ricordando che $S_N = S_n T_{n+1} \cdots T_N$ e che S_n è $\mathcal{F}_n$ -misurabile, otteniamo

$$\begin{aligned} C_n(1 + r)^{N-n} &= \mathbb{E}^*[F(S_N) | \mathcal{F}_n] = \\ &= \mathbb{E}^*[F(S_n T_{n+1} \cdots T_N) | \mathcal{F}_n] = \mathbb{E}^*[F(x T_{n+1} \cdots T_N)]_{|_{x=S_n}}, \end{aligned} \quad (4.15)$$

da cui segue che $C_n = c(n, S_n)$ con

$$c(n, x) = (1 + r)^{-(N-n)} \mathbb{E}^*[F(x T_{n+1} \cdots T_N)].$$

Ora, la v.a. $T_{n+1} \cdots T_N$ può assumere i valori $(1 + a)^j (1 + b)^{N-n-j}$ con probabilità pari alla probabilità che j v.a. tra $T_{n+1}, \dots, T_N$ assumano il valore $(1 + a)$ e le rimanenti $N - n - j$ assumano il valore $(1 + b)$, al variare di $j \in \{0, 1, \dots, N - n\}$. Poiché le $T_{n+1}, \dots, T_N$ sono indipendenti, ciò accade con probabilità pari a $\binom{N-n}{j} p^j (1 - p)^{N-n-j}$. Dunque,

$$\begin{aligned} &\mathbb{E}^*[F(x T_{n+1} \cdots T_N)] = \\ &= \sum_{j=0}^{N-n} \binom{N-n}{j} p^j (1 - p)^{N-n-j} F(x(1 + a)^j (1 + b)^{N-n-j}), \end{aligned} \quad (4.16)$$

Cerchiamo ora di vedere come si può fare la copertura dell'opzione. La strategia replicante $((\phi_n^0, \phi_n))_n$ deve essere tale che $V_n(\phi) = C_n$, quindi

$$\phi_n^0 S_n^0 + \phi_n S_n = c(n, S_n).$$

Considerando le due possibilità $S_n = S_{n-1} (1 + a)$ oppure $S_n = S_{n-1} (1 + b)$, otteniamo che deve essere

$$\begin{aligned} \phi_n^0 (1 + r)^n + \phi_n S_{n-1} (1 + a) &= c(n, S_{n-1} (1 + a)), \\ \phi_n^0 (1 + r)^n + \phi_n S_{n-1} (1 + b) &= c(n, S_{n-1} (1 + b)). \end{aligned} \quad (4.17)$$

Si vede subito che questo sistema lineare nelle incognite ϕ_n^0, ϕ_n ha sempre soluzione (la matrice dei coefficienti ha evidentemente rango 2, almeno se $a \neq b$). Inoltre è chiaro che la soluzione è una funzione della sola v.a. S_{n-1} , per cui essa è $\mathcal{F}_{n-1}$ -misurabile. Sottraendo la seconda equazione alla prima, si trova

$$\phi_n = \Delta(n, S_{n-1}) \quad \text{con} \quad \Delta(n, x) = \frac{c(n, x(1 + b)) - c(n, x(1 + a))}{x(b - a)} \quad (4.18)$$

Come previsto, ϕ_n è funzione della sola v.a. S_{n-1} , ed è effettivamente $\mathcal{F}_{n-1}$ -misurabile. Quindi il processo $(\phi_n)_n$ è predicibile. Sostituendo, nella prima delle equazioni (4.17), il valore di ϕ_n appena ottenuto, si ha $(1 + r)^n \phi_n^0 = c(n, S_n) - \phi_n S_n$ e si ottiene facilmente

$$\begin{aligned} \phi_n^0 &= \Delta^0(n, S_{n-1}) \quad \text{con} \\ \Delta^0(n, x) &= (1 + r)^{-n} \left(\frac{1 + b}{b - a} c(n, x(1 + a)) - \frac{1 + a}{b - a} c(n, x(1 + b)) \right). \end{aligned} \quad (4.19)$$

Ritroviamo che anche $(\phi_n^0)_n$ è un processo predicibile. La funzione

$$x \rightarrow \Delta(n, x) = \frac{c(n, x(1 + b)) - c(n, x(1 + a))}{x(b - a)} \quad (4.20)$$

si chiama la *delta* e, ad uno sguardo attento, essa misura la sensibilità della variazione del prezzo dell'opzione al tempo n rispetto ad una variazione del prezzo del titolo di base.

Vediamo ora un altro aspetto del modello CRR, che porterà a formulare una procedura numerica per la determinazione del prezzo delle opzioni della forma (4.14). Con un po' di lavoro, questo metodo può portare a procedure di calcolo per opzioni più complicate.

Abbiamo visto, nella prima parte di questo paragrafo, che il prezzo al tempo n delle opzioni della tipo (4.14) è della forma $c(n, S_n)$. È immediato che il processo

$$((1 + r)^{-n} c(n, S_n))_n \quad (4.21)$$

è una $\mathbb{P}^*$ -martingala. Infatti il prezzo al tempo n coincide con il valore al tempo n del portafoglio di copertura, dunque il processo in (4.21) non è altro che il portafoglio attualizzato.

Vale dunque la relazione

$$\mathbb{E}^*[(1+r)^{-(n+1)}c(n+1, S_{n+1}) | \mathcal{F}_n] = (1+r)^{-n}c(n, S_n)$$

Ma

$$\mathbb{E}^*[c(n+1, S_{n+1}) | \mathcal{F}_n] = \mathbb{E}^*[c(n+1, S_n T_{n+1}) | \mathcal{F}_n]$$

La v.a. S_n è $\mathcal{F}_n$ -misurabile, mentre T_{n+1} è indipendente da $\mathcal{F}_n$. Inoltre $\mathbb{P}^*(T_{n+1} = 1+a) = p^*$, $\mathbb{P}^*(T_{n+1} = 1+b) = 1-p^*$; dunque si ha

$$\begin{aligned} \mathbb{E}^*[c(n+1, S_n T_{n+1}) | \mathcal{F}_n] &= \mathbb{E}^*[c(n+1, x T_{n+1}) | \mathcal{F}_n] \big|_{x=S_n} = \\ &= (1-p^*)c(n+1, S_n(1+b)) + p^*c(n+1, S_n(1+a)) \end{aligned}$$

Si trova dunque la formula di ricorrenza

$$c(n, x) = \frac{1}{1+r} ((1-p^*)c(n+1, x(1+b)) + p^*c(n+1, x(1+a))), \quad (4.22)$$

che permette di calcolare “all’indietro” i valori di $c(n, x)$ lungo l’albero. Infatti il valore al tempo N del prezzo dell’opzione è ovviamente dato da

$$c(N, x) = F(x)$$

da questa relazione e dalla (4.22) si ricavano i valori di $c(N-1, x)$ e così via all’indietro nel tempo fino a trovare il valore del prezzo al tempo 0

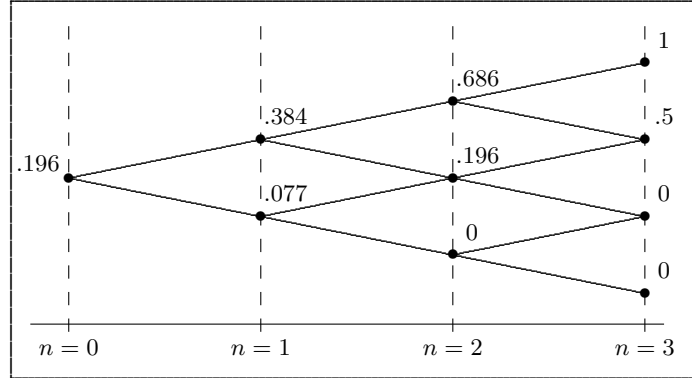


Figura 4.6 Esempio di calcolo all’indietro per un modello a tre periodi, per un’opzione che vale 1 per $S_3 = S_0(1+b)^3$, 0.5 per $S_3 = S_0(1+b)^2(1+a)$ e 0 per gli altri due possibili valori di S_3 . I valori numerici qui sono: $a = -0.1$, $b = 0.2$, $r = .02$. Ne segue $p^* = .6$.

La Figura 4.6 illustra il calcolo con lo schema appena descritto. Ad esempio il valore 0.686 che compare in uno dei nodi, viene ottenuto da

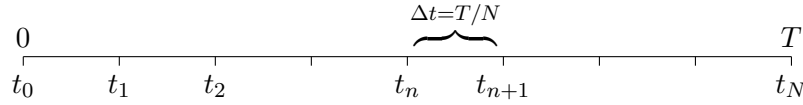
$$0.686 = \frac{1}{1+r} ((1-p^*) \cdot 1 + p^* \cdot 0.5) = \frac{1}{1.02} (1 \cdot 0.4 + 0.5 \cdot 0.6)$$

A titolo di verifica, calcoliamo il prezzo dell'opzione al tempo 0 con la formule (4.15) e (4.16):

$$C_0 = (1+r)^{-3} \left((1-p^*)^3 \cdot 1 + \binom{3}{1} p^* (1-p^*)^2 \right) = 0.196$$

4.6.3 Passaggio al limite e la formula di Black e Scholes

Le formule precedenti sono abbastanza soddisfacenti, ma richiedono comunque di essere aggiustate al mercato. In particolare per essere usate in pratica è necessario stabilire i valori di a e b . Ciò si può fare in concreto, vedremo però ora che, facendo crescere il numero di periodi N all'infinito, ed aggiustando convenientemente i valori di r , a e b , si può giungere ad espressioni per il prezzo delle opzioni e per i portafogli di copertura che sono semplici e sono facilmente interpretabili. Troveremo in particolare la formula di Black-Scholes, che ha costituito un po' il punto di partenza di tutta la moderna teoria dei modelli finanziari. In pratica supponiamo di suddividere l'intervallo temporale in N sotto periodi, nei quali il comportamento dell'attivo di base segue un modello CRR.



Ipotesi. 1) Supporremo $r = r_N$, con

$$r_N = \frac{RT}{N}. \quad (4.23)$$

Dunque il prezzo del titolo non rischioso $S_n^{0,N}$ è

$$S_n^{0,N} = (1 + r_N)^n = \left(1 + \frac{RT}{N}\right)^n.$$

Ciò significa che il denaro viene ricapitalizzato alla fine di ogni intervallo $[t_k, t_{k+1}] = [t_k, t_k + \Delta t]$ ad un tasso di interesse costante pari a $r_N = RT/N = R \Delta t$. Inoltre, si ha

$$S_N^{0,N} = \left(1 + \frac{RT}{N}\right)^N \xrightarrow{N \rightarrow \infty} e^{RT}.$$

La costante $R > 0$ prende il nome di *tasso istantaneo di interesse*.

2) I parametri $a = a_N$ e $b = b_N$ del modello sono scelti in modo tale che:

$$\log \frac{1+a_N}{1+r_N} = -\sigma \sqrt{T/N} \quad \text{e} \quad \log \frac{1+b_N}{1+r_N} = \sigma \sqrt{T/N} \quad (4.24)$$

o equivalentemente

$$1 + a_N = \left(1 + \frac{RT}{N}\right) e^{-\sigma \sqrt{T/N}} \quad \text{e} \quad 1 + b_N = \left(1 + \frac{RT}{N}\right) e^{\sigma \sqrt{T/N}} \quad (4.25)$$

dove σ è un parametro > 0 , sul cui significato torneremo più tardi. In particolare stiamo assumendo che, per $N \rightarrow \infty$, $a_N, b_N \rightarrow 0$,

Osserviamo che, poiché $\sigma > 0$, sotto l'Ipotesi 4.6.3 si ha

$$a_N < r_N < b_N,$$

quindi il mercato è privo di arbitraggio e completo (cfr. Teorema 4.6.4).

Osservazione 4.6.5. Sia $p = (b - r)/(b - a)$ il parametro che caratterizza la misura di martingala equivalente. Sotto l'Ipotesi 4.6.3 si ha $p = p_N$, con

$$p_N = \frac{e^{\sigma \sqrt{T/N}} - 1}{e^{\sigma \sqrt{T/N}} - e^{-\sigma \sqrt{T/N}}} \quad \text{e} \quad 1 - p_N = \frac{1 - e^{-\sigma \sqrt{T/N}}}{e^{\sigma \sqrt{T/N}} - e^{-\sigma \sqrt{T/N}}}. \quad (4.26)$$

Uno sviluppo in serie dà facilmente

$$\frac{e^x - 1}{e^x - e^{-x}} = \frac{x + \frac{1}{2}x^2 + o(x^2)}{2x + o(x^2)} = \frac{1}{2} + \frac{1}{4}x + o(x)$$

Sostituendo $x = \sigma \sqrt{T/N}$, è immediato verificare che

$$\lim_{N \rightarrow \infty} p_N = \lim_{N \rightarrow \infty} (1 - p_N) = \frac{1}{2}, \quad \lim_{N \rightarrow \infty} \sqrt{N}(2p_N - 1) = \frac{1}{2} \sigma \sqrt{T}. \quad (4.27)$$

Studiamo ora il comportamento per $N \rightarrow \infty$ del prezzo del bene sottostante e di quello della call e della put. Ricordiamo che

$$\log \frac{S_N^N}{S_0} = \sum_{n=1}^N \log T_n^N$$

dove ciascuna T_n^N può assumere i valori $1 + a_N$ e $1 + b_N$ e, sotto $\mathbb{P}^*$, le v.a. $T_1^N, \dots, T_N^N$ sono indipendenti ed identicamente distribuite. Ricordando i valori stabiliti per $1 + a_N$ e $1 + b_N$, si ha

$$\log \frac{S_N^N}{S_0} = N \log(1 + \frac{RT}{N}) + \sum_{n=1}^N \log \tilde{T}_n^N$$

ovvero

$$\log \frac{\tilde{S}_N^N}{S_0} = \sum_{n=1}^N \log \tilde{T}_n^N$$

dove le v.a. $\log \tilde{T}_n^N$ sono i.i.d. e prendono i valori $-\sigma \sqrt{T/N}$ e $+\sigma \sqrt{T/N}$ con probabilità p_N e $1 - p_N$ rispettivamente.

Proposizione 4.6.6. *La successione di v.a.*

$$\left(\sum_{n=1}^N \log \tilde{T}_n^N \right)_N \quad (4.28)$$

converge in legge, per $N \rightarrow \infty$ ad una v.a. gaussiana $N(-\frac{1}{2}\sigma^2 T, \sigma^2 T)$.

Dimostrazione. La dimostrazione consisterà nel calcolo del limite delle funzioni caratteristiche. A questo scopo, studiamo prima il comportamento della media, μ_N , e del momento del second'ordine, m_N^2 della v.a. $\log \tilde{T}_n^N$. Si ha

$$\begin{aligned}\mu_N &= -p_N \sigma \sqrt{\frac{T}{N}} + (1 - p_N) \sigma \sqrt{\frac{T}{N}} \\ m_N^2 &= \sigma^2 \frac{T}{N}\end{aligned}$$

Dunque, ricordando le (4.27), si ha

$$\begin{aligned}\lim_{n \rightarrow \infty} N \mu_N &= -\frac{1}{2} \sigma^2 T \\ \lim_{n \rightarrow \infty} N m_N^2 &= \sigma^2 T\end{aligned}\tag{4.29}$$

Ora, se indichiamo con φ_N la funzione caratteristica delle v.a. $U_n^N = \log \tilde{T}_n^N$, siamo ricondotti al calcolo del limite di

$$\left(\varphi_N(t) \right)^N, \text{ per } N \rightarrow \infty.$$

Ora sappiamo che

$$\varphi_N(t) = 1 + \varphi'_N(0)t + \frac{1}{2} \varphi''_N(0)t^2 + \frac{1}{3!} t^3 \varphi'''_N(\tau)\tag{4.30}$$

dove τ è compreso tra 0 e t . Ricordiamo le relazioni $\varphi'_N(0) = i\mu_N$, $\varphi''_N(0) = -m_N^2$, $\varphi'''_N(\tau) = \mathbb{E}[(iU_n^N)^3 e^{i\tau U_n^N}]$. Inoltre, poiché $|U_n^N| = \sigma \sqrt{T/N}$, si ha

$$|\varphi'''_N(\tau)| \leq \mathbb{E}[|iU_n^N|^3] = \sigma^3 \left(\frac{T}{N}\right)^{3/2}.$$

Riprendendo la (4.30), otteniamo

$$\varphi_N(t) = 1 + i\mu_N t - \frac{1}{2} m_N^2 t^2 + o\left(\frac{1}{N}\right)\tag{4.31}$$

Dunque la funzione caratteristica delle v.a. (4.28) vale

$$\begin{aligned}\left(\varphi_N(t) \right)^N &= \left(1 + i\mu_N t - \frac{1}{2} m_N^2 t^2 + o\left(\frac{1}{N}\right) \right)^N = \\ &= \left(1 + \frac{1}{N} \cdot N(i\mu_N t - \frac{1}{2} m_N^2 t^2 + o\left(\frac{1}{N}\right)) \right)^N.\end{aligned}$$

Per le (4.29),

$$N \times (i\mu_N t - \frac{1}{2} m_N^2 t^2 + o\left(\frac{1}{N}\right)) \xrightarrow{N \rightarrow \infty} -i\frac{1}{2} \sigma^2 T t - \frac{1}{2} \sigma^2 T t^2$$

e dunque

$$\left(\varphi_N(t) \right)^N \xrightarrow{N \rightarrow \infty} e^{-i\frac{1}{2} \sigma^2 T t} e^{-\frac{1}{2} \sigma^2 T t^2}$$

che è appunto la funzione caratteristica di una v.a. $N(-\frac{1}{2} \sigma^2 T, \sigma^2 T)$. $\square$

Abbiamo dunque dimostrato che, rispetto alla probabilità di rischio neutro $\mathbb{P}^*$, $\log(\tilde{S}_N^N/S_0)$ converge in legge verso una v.a. $Y \sim N(-\frac{1}{2}\sigma^2 T, \sigma^2 T)$, mentre, evidentemente $\log(S_N^N/S_0)$ converge in legge verso $RT + Y$. Dunque il prezzo attualizzato $\tilde{S}_N$ converge in legge verso una v.a. della forma $S_0 e^Y$, dove $Y \sim N(-\frac{1}{2}\sigma^2 T, \sigma^2 T)$. Una v.a. della forma e^Z con $Z \sim N(b, a)$ si dice che segue una *legge lognormale* di parametri b e a .

Osservazione 4.6.7. Nelle righe precedenti abbiamo usato due proprietà della convergenza in legge.

- a) Se $(X_n)_n$ converge in legge verso una v.a. X e ϕ è una funzione monotona strettamente crescente, allora $(\phi(X_n))_n$ converge in legge verso $\phi(X)$.
- b) Se $(X_n)_n$ converge in legge verso una v.a. X e $(a_n)_n$ è una successione di numeri reali con $a \rightarrow_{n \rightarrow \infty} a$, allora $(a_n + X_n)_n$ converge in legge verso $a + X$.

Sareste in grado di provare queste due affermazioni?

Passiamo ora a considerare il prezzo di una opzione put a tempo 0:

$$P_0^N = (1 + \frac{RT}{N})^{-N} \mathbb{E}^*[(K - S_N^N)_+] = \mathbb{E}^*[(1 + \frac{RT}{N})^{-N} K - \tilde{S}_N^N]_+,$$

Sappiamo che $\tilde{S}_N^N$ converge in legge alla v.a. $S_0 e^Y$ con $Y \sim N(-\frac{1}{2}\sigma^2 T, \sigma^2 T)$, dunque (si usa ancora il punto b) della osservazione precedente) si ha che

$$(1 + \frac{RT}{N})^{-N} K - \tilde{S}_N^N \xrightarrow[N \rightarrow \infty]{\mathcal{L}} e^{-RT} K - S_0 e^Y \quad (4.32)$$

Abbiamo bisogno ora di un risultato che verrà visto in un corso più avanzato

Teorema 4.6.8. *La successione $(X_n)_n$ converge in legge verso una v.a. X se e solo se, per ogni funzione reale f , continua e limitata, si ha*

$$\lim_{n \rightarrow \infty} \mathbb{E}[f(X_n)] = \mathbb{E}[f(X)]$$

Dal Teorema 4.6.8 segue che

$$\lim_{N \rightarrow \infty} \mathbb{E}^*[(1 + \frac{RT}{N})^{-N} K - \tilde{S}_N^N]_+ = \mathbb{E}^*[(e^{-RT} K - S_0 e^Y)_+] \quad (4.33)$$

Infatti le v.a. $(1 + \frac{RT}{N})^{-N} K - \tilde{S}_N^N$ e $e^{-RT} K - S_0 e^Y$ prendono comunque valori più piccoli di K . Dunque, posto $f(x) = x_+ \wedge K = \min(x_+, K)$, si può scrivere

$$\begin{aligned} ((1 + \frac{RT}{N})^{-N} K - \tilde{S}_N^N)_+ &= f((1 + \frac{RT}{N})^{-N} K - \tilde{S}_N^N) \\ (e^{-RT} K - S_0 e^Y)_+ &= f(e^{-RT} K - S_0 e^Y) \end{aligned}$$

Si può quindi applicare il Teorema 4.6.8 alla funzione f che è continua e limitata.

Teorema 4.6.9. *Nell'Ipotesi 4.6.3, si ha*

$$\lim_{N \rightarrow \infty} P_0^N = P_0$$

dove

$$P_0 = K e^{-RT} \Phi(-d_2) - S_0 \Phi(-d_1) \quad (4.34)$$

essendo Φ la funzione di ripartizione di una legge gaussiana standard e

$$\begin{aligned} d_1 &= \frac{1}{\sigma \sqrt{T}} \left(\log \frac{S_0}{K} + RT + \frac{1}{2} \sigma^2 T \right) \\ d_2 &= d_1 - \sigma \sqrt{T} = \frac{1}{\sigma \sqrt{T}} \left(\log \frac{S_0}{K} + RT - \frac{1}{2} \sigma^2 T \right). \end{aligned}$$

Dimostrazione. Per la (4.33):

$$\lim_{N \rightarrow \infty} P_0^N = \mathbb{E}^*[(e^{-RT} K - S_0 e^Y)_+],$$

dove $Y \sim N(-\frac{1}{2} \sigma^2 T, \sigma^2 T)$. Ci resta da mostrare che la speranza matematica a destra nella relazione precedente vale in effetti come precisato nella (4.34). Osserviamo che possiamo scrivere $Y = -\frac{1}{2} \sigma^2 T + \sigma \sqrt{T} Z$, dove $Z \sim N(0, 1)$. Indicando con x il valore di S_0 , si ha

$$\begin{aligned} P_0 &= \mathbb{E}[f(-\frac{1}{2} \sigma^2 T + \sigma \sqrt{T} Z)] \\ &= \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{+\infty} (e^{-RT} K - x e^{-\frac{1}{2} \sigma^2 T + \sigma \sqrt{T} z})_+ e^{-z^2/2} dz. \end{aligned}$$

Ora la disuguaglianza

$$e^{-RT} K - x e^{-\frac{1}{2} \sigma^2 T + \sigma \sqrt{T} z} \geq 0$$

è soddisfatta se e solo se

$$-RT + \log K \geq \log x - \frac{1}{2} \sigma^2 T + \sigma \sqrt{T} z$$

ovvero

$$z \leq \frac{1}{\sigma \sqrt{T}} \left(-\log \frac{x}{K} - RT + \frac{1}{2} \sigma^2 T \right) = -d_2.$$

Quindi P_0 è uguale a

$$\begin{aligned} P_0 &= \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{-d_2} \{e^{-RT} K - x e^{-\frac{1}{2} \sigma^2 T + \sigma \sqrt{T} z}\} e^{-z^2/2} dz = \\ &= \frac{e^{-RT} K}{\sqrt{2\pi}} \int_{-\infty}^{-d_2} e^{-z^2/2} dz - \frac{x}{\sqrt{2\pi}} \int_{-\infty}^{-d_2} e^{-\frac{1}{2} \sigma^2 T + \sigma \sqrt{T} z} e^{-z^2/2} dz \\ &= e^{-RT} K \Phi(-d_2) - \frac{x}{\sqrt{2\pi}} \int_{-\infty}^{-d_2} e^{-\frac{1}{2} (z - \sigma \sqrt{T})^2} dz \end{aligned}$$

Con un semplice cambio di variabile, si ha

$$\frac{1}{\sqrt{2\pi}} \int_{-\infty}^{-d_2} e^{-\frac{1}{2} (z - \sigma \sqrt{T})^2} dz = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{-d_2 + \sigma \sqrt{T}} e^{-t^2/2} dt = \Phi(-d_1)$$

da cui segue la tesi

□

Teorema 4.6.10. *Nell'Ipotesi 4.6.3, esiste*

$$\lim_{N \rightarrow \infty} C_0^N = C_0$$

dove

$$C_0 = S_0 \Phi(d_1) - K e^{-RT} \Phi(d_2).$$

Dimostrazione. Per dimostrare la convergenza della call, non possiamo ricalcare la dimostrazione dell'analogo risultato della put. Infatti, la v.a. $(S_0 e^Y - e^{-RT} K)_+$ non è limitata e non si può applicare il Teorema 4.6.8. Useremo un altro metodo, in particolare la formula di parità call/put: da (4.11), si ha infatti

$$C_0^N = P_0^N + S_0 - K \left(1 + \frac{RT}{N}\right)^{-N}$$

quindi C_0^N converge, per $N \rightarrow \infty$ a

$$\begin{aligned} P_0 + S_0 - K e^{-RT} &= S_0(1 - \Phi(-d_1)) - K e^{-RT} (1 - \Phi(-d_2)) \\ &= S_0 \Phi(d_1) - K e^{-RT} \Phi(d_2) = C_0. \end{aligned}$$

□

Ci si può chiedere anche quale debba essere il valore del prezzo ad un tempo t intermedio tra la data di emissione e la maturità T . È chiaro che, perché non ci siano opportunità di arbitraggio, occorre che il prezzo di un'opzione call emessa al tempo t debba avere lo stesso prezzo, C_t , di un'opzione call emessa al tempo 0 ed avente lo stesso prezzo di esercizio e la stessa maturità. Dunque l'espressione per C_t deve essere, $C_t = c(t, S_t)$, dove

$$c(t, x) = x \Phi(d_1(t)) - K e^{-R(T-t)} \Phi(d_2(t)). \quad (4.35)$$

dove

$$\begin{aligned} d_1(t) &= \frac{1}{\sigma \sqrt{T-t}} \left(\log \frac{x}{K} + R(T-t) + \frac{1}{2} \sigma^2 (T-t) \right) \\ d_2(t) &= d_1 - \sigma \sqrt{T-t} = \frac{1}{\sigma \sqrt{T-t}} \left(\log \frac{x}{K} + R(T-t) - \frac{1}{2} \sigma^2 (T-t) \right). \end{aligned} \quad (4.36)$$

cioè la stessa espressione di C_0 , con $T - t$ al posto di T .

4.6.4 Studio empirico della velocità di convergenza

In questo paragrafo proponiamo il seguente

Esercizio 4.1. *Studiare empiricamente, con l'ausilio del calcolatore, la velocità di convergenza del prezzo CRR al prezzo Black e Scholes della call e della put.*

Commentiamo meglio cosa si chiede. Consideriamo, a titolo di esempio, il caso della call.

Nel Teorema 4.6.10, abbiamo dimostrato che

$$\lim_{N \rightarrow \infty} C_0^N = C_0$$

dove C_0^N è il prezzo CRR della call su N suddivisioni dell'intervallo $[0, T]$ e C_0 è il prezzo BS dell'analoga opzione. Problema: qual è la velocità di convergenza? In altre parole, si chiede di determinare α affinché

$$C_0^N = C_0 + \frac{c}{N^\alpha} + o\left(\frac{1}{N^\alpha}\right),$$

dove c rappresenta una costante non nulla opportuna e, al solito, $o(1/N^\alpha)$ denota un infinitesimo di ordine superiore a $1/N^\alpha$: $\lim_{N \rightarrow \infty} N^\alpha o(1/N^\alpha) = 0$. Per risolvere, almeno numericamente, questo problema, si può procedere come segue. Innanzitutto osserviamo che, passando al logaritmo, α si cerca affinché

$$\log |C_0^N - C_0| = c_1 - \alpha \log N + \log(1 + \gamma_N) \simeq c_1 - \alpha \log N,$$

perché $\gamma_N = N^\alpha \log(1 + N^\alpha o(1/N^\alpha)) \rightarrow 0$ quando $N \rightarrow \infty$. Dunque, in scala logaritmica, si cerca il coefficiente angolare che regola la relazione lineare tra $\log N$ e $\log |C_0^N - C_0|$. Questo si può fare in vari modi, lasciamo aperta ogni possibilità di soluzione. Nella pratica, però, occorre tener conto di alcune cose. Innanzitutto, non si può ad esempio pensare di plottare i dati per ogni N : occorre fissare un buon numero di valori per N grande ma non troppo (bisogna anche tener conto dei tempi di calcolo!). Ad esempio, si può scegliere

$$N = N_0 + i\delta, \quad i = 0, 1, \dots, I$$

dove N_0 e δ sono, rispettivamente, un valore minimo e un passo prefissati. Poi, si procede a calcolare C_0^N . Il modo più opportuno è quello di usare la formula “all'indietro” (4.22). Infine, una volta calcolato C_0 , si passa alla scala logaritmica e si procede allo studio empirico.

Prima di passare ad altro, un ulteriore commento. Per calcolare C_0 occorre la funzione di ripartizione della normale standard: nel Paragrafo 4.9.2, è proposto un codice⁹ in C.

Una volta determinato il valore di α , è possibile costruire una successione che converge più velocemente a C_0 con il metodo di estrapolazione di Romberg. Vale la pena di citarlo e mostrarlo per la sua semplicità ed efficacia, oltre al fatto che è del tutto generale. Supponiamo infatti che la successione (numerica) $\{x_N\}_N$ sia tale che

$$x_N = x + \frac{c}{N^\alpha} + o\left(\frac{1}{N^\alpha}\right),$$

⁹All'indirizzo <http://www.mat.uniroma2.it/~caramell/did.0607/pf.htm> è scaricabile il relativo file di testo.

e quindi

$$x_{2N} = x + \frac{c}{2^\alpha N^\alpha} + o\left(\frac{1}{N^\alpha}\right).$$

Dividendo allora entrambi i membri della prima uguaglianza e sottraendo la seconda si ottiene

$$\frac{2^\alpha x_{2N} - x_N}{2^\alpha - 1} = x + o\left(\frac{1}{N^\alpha}\right).$$

Abbiamo quindi determinato una nuova successione $\{y_N\}_N$, dove

$$y_N = \frac{2^\alpha x_{2N} - x_N}{2^\alpha - 1},$$

che converge a x per $N \rightarrow \infty$ più velocemente di quella di partenza.

Esercizio 4.2. *Una volta risolto l'Esercizio 4.1, implementare il metodo di estrapolazione di Romberg e verificare empiricamente che la nuova successione converge più velocemente a C_0 .*

A proposito: ovviamente il valore giusto per α si conosce, la dimostrazione però è complicata e pressoché impossibile da fare con i prerequisiti di questo corso. Dimenticavo: $\alpha = 1$.

4.7 La volatilità

La (4.35) del paragrafo precedente, è la formula di Black-Scholes. In questo paragrafo e nel prossimo studiamo un po' il comportamento del prezzo dell'opzione in funzione delle quantità da cui dipende.

Intanto osserviamo che essa dipende da un certo numero di parametri. Di questi, quasi tutti sono noti a priori: x , R , T e K sono quantità note al tempo 0. Il solo parametro che resta da determinare per applicare concretamente la formula è la volatilità σ .

Vediamo innanzitutto, in che modo il prezzo dipende dalla volatilità, calcolandone la derivata. Si trova, con un po' di lavoro, che

$$\frac{\partial C_0}{\partial \sigma} = \frac{x\sqrt{T}}{\sqrt{2\pi}} e^{-\frac{1}{2} d_1(t)^2}$$

In particolare il prezzo è, una funzione strettamente crescente (e quindi invertibile) della volatilità. Questa osservazione ha due conseguenze importanti.

Intanto gli attivi finanziari per i quali il prezzo dell'opzione è elevato sono quelli molto volatili (cioè quelli per i quali la volatilità σ è elevata). Tenendo conto del fatto che, più è grande σ e più sono grandi le oscillazioni del prezzo in un intervallo di tempo (cfr. la (4.25)), i titoli più volatili (e quindi le cui opzioni sono le più costose) sono quelli per i quali i prezzi presentano le maggiori oscillazioni del prezzo nel tempo.

Il fatto che il prezzo sia una funzione invertibile di σ fornisce un modo semplice di determinare la volatilità a partire dai dati del mercato. Se c^* è il prezzo di un'opzione quotata sul mercato, per una data maturità ed un dato prezzo di esercizio, invertendo l'applicazione che alla volatilità associa il prezzo, per K, R, T fissati, si può ottenere il valore della volatilità σ in funzione di c^* . Questo valore può essere utilizzato per determinare il prezzo di opzioni con altri valori del prezzo di esercizio e della maturità. Questa osservazione permette anche di verificare la bontà del modello che abbiamo sviluppato. Infatti i valori della volatilità ottenuti in questo modo invertendo la funzione prezzo, per delle opzioni con valori diversi del prezzo di esercizio e della maturità, devono dare lo stesso risultato. Delle discrepanze tra i valori ottenuti indicheranno che il modello spiega solo parzialmente il mercato reale. Nella realtà queste anomalie si osservano in concreto, ma il modello di Black-Scholes viene tuttora considerato un buon modello e comunque una tappa obbligata verso la costruzione di modelli più sofisticati.

4.8 Le greche

Una certa importanza hanno le derivate della funzione prezzo, $c(t, x)$, ottenuta nella (4.35). Queste derivate, che indicano la sensibilità alla variazione del parametro corrispondente, tradizionalmente si indicano con lettere dell'alfabeto greco, per cui vengono chiamate le *greche*. Abbiamo già incontrato una di queste derivate, quella rispetto alla volatilità, che si chiama la *vega* (che però non è una lettera greca...). Le altre greche sono

- la delta, cioè la derivata rispetto al prezzo dell'attivo sottostante

$$\Delta = \frac{\partial c}{\partial x}(x, t)$$

- la gamma, che è la derivata seconda, sempre rispetto a x

$$\Gamma = \frac{\partial^2 c}{\partial x^2}(x, t)$$

- la theta, cioè la derivata rispetto al tempo

$$\Theta = \frac{\partial c}{\partial t}(x, t)$$

Calcoleremo ora queste quantità, nel caso dell'opzione call. Per la delta, si ha

$$\frac{\partial c}{\partial x}(x, t) = \Phi(d_1(t)) + \underbrace{x \Phi'(d_1(t)) \frac{\partial d_1(t)}{\partial x} - K e^{-R(T-t)} \Phi'(d_2(t)) \frac{\partial d_2(t)}{\partial x}}_{(A)}$$

Osserviamo che

$$\frac{\partial d_1(t)}{\partial x} = \frac{\partial d_2(t)}{\partial x}$$

poiche' le due funzioni d_1 e d_2 differiscono per una quantità che non dipende da x . Mostriamo che la quantità indicata con (A) nella formula precedente si annulla. Infatti essa è uguale a

$$\frac{\partial d_1(t)}{\partial x} \Phi'(d_2(t)) \left\{ x \partial \Phi'(d_1(t)) \Phi'(d_2(t)) - K e^{-RT} \right\}$$

Ricordando che Φ' è la densità della gaussiana $N(0, 1)$,

$$\begin{aligned} \frac{\Phi'(d_1(t))}{\Phi'(d_2(t))} &= e^{-\frac{1}{2}(d_1(t)^2 - d_2(t)^2)} = e^{-\frac{1}{2}(d_1(t) - d_2(t))(d_1(t) + d_2(t))} = \\ &= e^{-\frac{1}{2} \cdot 2(\log \frac{x}{K} + RT)} = \frac{K}{x} e^{-RT} \end{aligned}$$

e sostituendo si trova l'espressione cercata per la delta

$$\frac{\partial c}{\partial x}(x, t) = \Phi(d_1(t))$$

Da notare che si tratta sempre di una quantità compresa tra 0 e 1. Possiamo ora facilmente calcolare la gamma della call.

$$\frac{\partial^2 c}{\partial x^2} = \Phi'(d_1(t)) \frac{\partial d_1(t)}{\partial x} = \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}d_1(t)^2} \frac{1}{\sigma\sqrt{\tau}x}$$

dove indichiamo $\tau = T - t$. Con un calcolo un po' più lungo si ottiene la theta:

$$\frac{\partial c}{\partial t}(x, t) = - \left[\frac{x\sigma}{2\sqrt{T-t}\sqrt{2\pi}} e^{-\frac{1}{2}d_1(t)^2} + R K e^{-R(T-t)} \Phi(d_2(t)) \right]$$

Mettendo insieme le relazioni ottenute, si vede che la funzione c è soluzione del problema alle derivate parziali

$$\frac{\partial c}{\partial t} + \frac{\partial^2 c}{\partial x^2} + R x \frac{\partial c}{\partial x} - R c = 0$$

con la condizione $c(x, T) = (x - K)_+$.

4.9 Appendice al Capitolo 4

4.9.1 Teorema di separazione dei convessi

In questo paragrafo dimostriamo il Teorema 4.4.4 (*teorema di separazione dei convessi*), di cui ricordiamo l'enunciato:

Sia $K \subset \mathbb{R}^m$ un insieme compatto e convesso e sia $\mathcal{V}$ un sottospazio di $\mathbb{R}^m$ tali che $K \cap \mathcal{V} = \emptyset$. Allora, esiste $\lambda \in \mathbb{R}^m$ tale che

$$\text{per ogni } x \in K, \sum_{\ell=1}^m \lambda_\ell x_\ell > 0;$$

$$\text{per ogni } x \in \mathcal{V}, \sum_{\ell=1}^m \lambda_\ell x_\ell = 0.$$

La dimostrazione è diretta conseguenza del seguente

Teorema 4.9.1. *Sia $C \subset \mathbb{R}^m$ un insieme chiuso e convesso che non contenga l'origine. Allora, esiste un funzionale lineare ξ su $\mathbb{R}^m$ ed un numero $\alpha > 0$ tali che*

$$\text{per ogni } x \in C \text{ si ha } \xi(x) \geq \alpha.$$

Dimostrazione. Sia R un numero positivo tale che la palla chiusa $\bar{B}_R$ centrata nell'origine e di raggio R interseca C . L'insieme $C \cap \bar{B}_R$ è un compatto (perché chiuso e limitato), dunque l'applicazione $C \cap \bar{B}_R \ni x \mapsto \|x\|$ assume il minimo. Sia x_0 il punto in cui il minimo è assunto. Segue immediatamente che

$$\text{per ogni } x \in C \text{ si ha } \|x\| \geq \|x_0\|.$$

Osserviamo che x_0 non è altro che la proiezione dell'origine su C . Ora, preso $x \in C$, tutti i punti del tipo $x_0 + t(x - x_0) \in C$ quando $t \in [0, 1]$, perché $x, x_0 \in C$ e C è convesso. Quindi, per ogni $t \in [0, 1]$, $\|x_0 + t(x - x_0)\| \geq \|x_0\|$ e sviluppando questa disuguaglianza¹⁰ si ottiene:

$$\langle x_0, x \rangle \geq \|x_0\|,$$

per ogni $x \in C$. Posto allora $\xi(x) = \langle x_0, x \rangle$ e $\alpha = \|x_0\|$, il teorema è dimostrato. □

Possiamo ora passare alla

Dimostrazione del Teorema 4.4.4. Poniamo

$$C = K - \mathcal{V} = \{x \in \mathbb{R}^m : x = y - z, \text{ per qualche } y \in K, z \in \mathcal{V}\}.$$

È immediato vedere che C è convesso, chiuso e non contiene l'origine. Per il Teorema 4.9.1 esistono un funzionale lineare ξ su $\mathbb{R}^m$ ed un numero positivo α tali che per ogni $x \in C$ si ha $\xi(x) \geq \alpha$. Quindi,

$$\text{per ogni } y \in K \text{ e } z \in \mathcal{V}, \xi(y) - \xi(z) \geq \alpha. \quad (4.37)$$

¹⁰Ricordiamo che $\|x + y\|^2 = \|x\|^2 + \|y\|^2 + 2\langle x, y \rangle$.

Ora, fissiamo $y \in K$ e $z \in \mathcal{V}$. Sia $c > 0$: applicando la disuguaglianza (4.37) a $cz \in \mathcal{V}$, si ottiene $\xi(z) \leq (\xi(y) - \alpha)/c \rightarrow 0$ per $c \rightarrow +\infty$, dunque $\xi(z) \leq 0$. Se invece prendiamo $c < 0$ e poi $c \rightarrow -\infty$, otteniamo $\xi(z) \geq 0$. Ma allora, $\xi(z) = 0$ per ogni $z \in \mathcal{V}$ e quindi $\xi(y) \geq \alpha > 0$ per ogni $y \in K$. Infine, osserviamo che un funzionale lineare su $\mathbb{R}^m$ è sempre del tipo $\xi(x) = \sum_{i=1}^m \lambda_i x_i$, per un opportuno vettore $\lambda \in \mathbb{R}^m$, da cui segue la tesi. $\square$

4.9.2 Un codice C per la funzione di ripartizione normale standard

```
/*One-Dimensional Normal Law.
Cumulative distribution function.
Abramowitz, Milton et Stegun,
Handbook of Mathematical Functions, 1968,
Dover Publication, New York, page 932 (26.2.18).
Precision 10-7*/

double N(double x) {
const double p= 0.2316419;
const double b1= 0.319381530;
const double b2= -0.356563782;
const double b3= 1.781477937;
const double b4= -1.821255978;
const double b5= 1.330274429;
const double one_over_twopi= 0.39894228;

double t;

if(x >= 0.0)
{
t = 1.0 / ( 1.0 + p * x );
return (1.0 - one_over_twopi * exp( -x * x / 2.0 ) *
t * ( t * ( t * ( t * ( t * b5 + b4 ) + b3 ) + b2 )
+ b1 ));
}
else
{
/* x < 0 */
t = 1.0 / ( 1.0 - p * x );
return ( one_over_twopi * exp( -x * x / 2.0 ) *
t * ( t * ( t * ( t * ( t * b5 + b4 ) + b3 ) + b2 )
+ b1 ));
}
}
```

Capitolo 5

Opzioni americane

5.1 Il modello

Consideriamo un modello di mercato finanziario così come descritto nel Paragrafo 4.1. Il mercato è quindi formato da $d+1$ titoli di prezzi $S_n^0, S_n^1, \dots, S_n^d$, dove

$$S_n^0 = (1+r)^n$$

denota il prezzo del titolo non rischioso e $S_n^1, \dots, S_n^d$ i prezzi dei d titoli rischiosi al tempo n . Indicheremo ancora con $(\mathcal{F}_n)_n$ la filtrazione associata, rispetto alla quale $(S_n)_n = ((S_n^0, S_n^1, \dots, S_n^d))_n$ è adattato, ricordando che $\mathcal{F}_n$ dà l'evoluzione del mercato fino al tempo n .

Supponiamo d'ora in poi che il mercato sia privo di arbitraggio e completo: esiste ed è unica la misura equivalente di martingala $\mathbb{P}^*$, tale che il processo dei prezzi scontati dei titoli rischiosi è una $\mathcal{F}_n$ -martingala. Denoteremo, al solito, con $\mathbb{E}^*$ l'aspettazione valutata sotto $\mathbb{P}^*$.

In questo capitolo intendiamo studiare il prezzo e la copertura delle opzioni americane. Ricordiamo che un'opzione americana è caratterizzata dal fatto di poter essere esercitata in un qualsiasi istante minore od uguale alla maturità N . Per caratterizzare un'opzione americana è quindi necessario precisare per ogni tempo n , il payoff, Z_n . Cioè la quantità di denaro che occorre sborsare se il detentore dell'opzione decide di esercitarla al tempo n . Poiché il payoff al tempo n deve, ragionevolmente, essere noto al tempo n , siamo condotti alla seguente definizione, nella quale identifichiamo un'opzione con il suo payoff.

Definizione 5.1.1. *Un'opzione americana è una famiglia di v.a. positive $(Z_n)_n$, adattata alla filtrazione $(\mathcal{F}_n)_n$.*

Ad esempio, per una call o una put americana, scritta sul titolo di prezzo $(S_n^1)_n$, si avrà rispettivamente

$$\begin{aligned} Z_n &= Z_n^{\text{call}} = (S_n^1 - K)_+ \\ Z_n &= Z_n^{\text{put}} = (K - S_n^1)_+. \end{aligned}$$

Nella problematica delle opzioni americane è opportuno anche considerare il problema dal punto di vista del detentore: quando conviene esercitare il diritto di opzione? In particolare occorre modellizzare l'istante (aleatorio), τ , in cui l'opzione viene esercitata. Poiché è ragionevole supporre che questo istante dipenda dall'evoluzione del mercato, dovrà essere $\{\tau = n\} \in \mathcal{F}_n$, per ogni n . Ovvero, ricordando il paragrafo 4.4, supporremo che τ sia un $\mathcal{F}_n$ -tempo d'arresto. Poniamo

$$\mathcal{T}_{0,N} = \{(\mathcal{F}_n)_n - \text{tempi d'arresto a valori in } \{0, 1, \dots, N\}\}.$$

l'insieme di tutti i tempi d'arresto della filtrazione $(\mathcal{F}_n)_n$. Vedremo che il prezzo di un'opzione americana è pari a

$$\sup_{\tau \in \mathcal{T}_{0,N}} \mathbb{E}^*[(1+r)^{-\tau} Z_\tau]. \quad (5.1)$$

Cioè il prezzo è uguale al sup, al variare della strategia d'esercizio, della media dei payoff scontati, media fatta rispetto alla probabilità di rischio neutro. Vedremo inoltre che il sup nella relazione precedente è in realtà un massimo e, dunque, esiste (almeno) una strategia di esercizio ottimale.

Come si vede lo studio delle opzioni americane è naturalmente legato a quello della determinazione del tempo di esercizio ottimale. Molti degli oggetti che introdurremo sono propri della teoria dell'*arresto ottimo*.

5.2 Il problema dell'arresto ottimo

Per stabilire il prezzo equo di un'opzione americana, cominciamo procedendo con un'induzione "all'indietro". Vediamo come. Consideriamo un'opzione americana $(Z_n)_n$ e indichiamo con U_n il suo prezzo al tempo n . Attenzione alla solita possibilità di confusione di questa terminologia: Z_n è il payoff dell'opzione, cioè quanto deve sborsare il venditore per onorare il contratto, qualora il detentore decida di esercitare l'opzione al tempo n ; U_n è invece il giusto prezzo che il venditore deve ricevere come compenso dell'opzione e che ora determineremo.

Tempo N . Il compratore dell'opzione decide di esercitare il suo diritto d'opzione: il venditore deve sborsare una quantità di denaro pari a

$$U_N = Z_N.$$

Tempo $N - 1$. Ci sono due possibilità.

1. Il compratore dell'opzione decide di esercitare il suo diritto d'opzione all'istante $N - 1$. In tal caso, il venditore deve pagare un ammontare pari a Z_{N-1} .

2. Il compratore decide di non esercitare l'opzione all'istante $N - 1$ ma all'istante N . In questo caso, il venditore deve essere pronto, all'istante $N - 1$, a pagare all'istante N una quantità di denaro pari a Z_N . Per onorare questo contratto, abbiamo visto che all'istante $N - 1$ il venditore deve possedere, in media,

$$(1 + r)^{-(N-(N-1))} \mathbb{E}^*(Z_N | \mathcal{F}_{N-1}) = \frac{1}{(1 + r)} \mathbb{E}^*(U_N | \mathcal{F}_{N-1}).$$

Dunque, all'istante $N - 1$ il venditore deve possedere una quantità di denaro U_{N-1} che deve coprire le due possibilità appena elencate, e cioè

$$U_{N-1} = \max \left(Z_{N-1}, \frac{1}{(1 + r)} \mathbb{E}^*(U_N | \mathcal{F}_{N-1}) \right).$$

Tempo $n = N - 2, N - 1, \dots, 0$. Si procede per ricorrenza in modo analogo a quanto già visto. Supponiamo quindi che la quantità U_{n+1} sia ben definita, essendo U_{n+1} l'ammontare di denaro che il venditore dell'opzione deve possedere all'istante (successivo) $n + 1$ per poter onorare il contratto d'opzione. Allora:

1. se il compratore dell'opzione decide di esercitare il suo diritto d'opzione all'istante n , allora il venditore deve possedere una quantità di denaro pari a Z_n ;
2. se invece il compratore decide di non esercitare l'opzione all'istante n , allora il venditore deve essere pronto, all'istante n , a pagare all'istante $n + 1$ una quantità di denaro pari a U_{n+1} , dunque deve possedere

$$(1 + r)^{-(n+1-n)} \mathbb{E}^*(U_{n+1} | \mathcal{F}_n) = \frac{1}{(1 + r)} \mathbb{E}^*(U_{n+1} | \mathcal{F}_n).$$

Dunque, all'istante n il venditore deve possedere una quantità di denaro U_n che deve coprire le due possibilità appena elencate, e cioè

$$U_n = \max \left(Z_n, \frac{1}{(1 + r)} \mathbb{E}^*(U_{n+1} | \mathcal{F}_n) \right).$$

U_n rappresenta quindi il prezzo dell'opzione all'istante n .

Ricapitolando, siamo condotti a considerare il processo $(U_n)_n$ definito da

$$\begin{aligned} U_N &= Z_N \quad \text{e per } n = N - 1, N - 2, \dots, 0, \\ U_n &= \max \left(Z_n, \frac{1}{(1 + r)} \mathbb{E}^*(U_{n+1} | \mathcal{F}_n) \right) \end{aligned} \tag{5.2}$$

Vedremo, nel prossimo paragrafo, che $(U_n)_n$ è effettivamente il prezzo dell'opzione al tempo n . In questo paragrafo studieremo le sue proprietà e

metteremo in evidenza come esso interviene nel problema dell'arresto ottimo, cioè nella determinazione dei tempi di arresto per i quali il sup nella (5.1) è raggiunto.

Abbiamo visto che, per le opzioni europee, i prezzi scontati costituiscono una $\mathbb{P}^*$ -martingala. Vale la stessa proprietà anche nel caso americano? La risposta è no. Vale però il risultato seguente.

Proposizione 5.2.1. *Sia U_n il prezzo di un'opzione americana di payoff $(Z_n)_n$ al tempo n , $n = 0, 1, \dots, N$, dato da (5.2). Siano $(\tilde{U}_n)_n$ e $(\tilde{Z}_n)_n$ rispettivamente il processo di prezzo e di payoff scontato:*

$$\tilde{U}_n = \frac{U_n}{S_n^0} = (1+r)^{-n} U_n \quad \text{e} \quad \tilde{Z}_n = \frac{Z_n}{S_n^0} = (1+r)^{-n} Z_n.$$

Allora $(\tilde{U}_n)_n$ è una $\mathbb{P}^*$ -supermartingala ed inoltre è la più piccola supermartingala che domina $(\tilde{Z}_n)_n$, cioè tale che $\tilde{U}_n \geq \tilde{Z}_n$.

Un po' di terminologia: dato un processo $(X_n)_n$, si chiama inviluppo di Snell di $(X_n)_n$ la più piccola supermartingala che domina $(X_n)_n$. Il problema del calcolo del prezzo delle opzioni americane si riconduce quindi al calcolo dell'inviluppo di Snell del processo di payoff $(\tilde{Z}_n)_n$.

Osserviamo che, da (5.2), segue immediatamente che

$$\begin{aligned} \tilde{U}_N &= \tilde{Z}_N \quad \text{e per } n = N-1, N-2, \dots, 0, \\ \tilde{U}_n &= \max(\tilde{Z}_n, \mathbb{E}^*(\tilde{U}_{n+1} | \mathcal{F}_n)) \end{aligned} \tag{5.3}$$

Dimostrazione della Proposizione 5.2.1. Da (5.3), si ha

$$\tilde{U}_n \geq \mathbb{E}^*(\tilde{U}_{n+1} | \mathcal{F}_n) \quad \text{e} \quad \tilde{U}_n \geq \tilde{Z}_n,$$

dunque $(\tilde{U}_n)_n$ è una $\mathbb{P}^*$ -supermartingala che domina $(\tilde{Z}_n)_n$. Mostriamo che è la più piccola. Sia $(\tilde{V}_n)_n$ un'altra $\mathbb{P}^*$ -supermartingala che domina $(\tilde{Z}_n)_n$: mostriamo che $\tilde{V}_n \geq \tilde{U}_n$. Intanto, si ha $\tilde{V}_N \geq \tilde{Z}_N = \tilde{U}_N$. Ma allora,

$$\tilde{V}_{N-1} \geq \mathbb{E}^*(\tilde{V}_N | \mathcal{F}_{N-1}) \geq \mathbb{E}^*(\tilde{U}_N | \mathcal{F}_{N-1})$$

e quindi, poiché deve essere $\tilde{V}_{N-1} \geq \tilde{Z}_{N-1}$,

$$\tilde{V}_{N-1} \geq \max(\tilde{Z}_{N-1}, \mathbb{E}^*(\tilde{U}_N | \mathcal{F}_{N-1})) = \tilde{U}_{N-1}.$$

Dunque, si ha anche $\tilde{V}_{N-1} \geq \tilde{U}_{N-1}$. Procedendo per ricorrenza, otteniamo che $\tilde{V}_n \geq \tilde{U}_n$ per ogni n , da cui la tesi. $\square$

Continuiamo a studiare le proprietà di martingala associate a $(U_n)_n$.

Proposizione 5.2.2. *Sia*

$$\nu_0 = \inf\{n \geq 0 \text{ t.c. } \tilde{U}_n = \tilde{Z}_n\}.$$

Allora ν_0 è un $\mathcal{F}_n$ -tempo d'arresto a valori in $\{0, 1, \dots, N\}$ e il processo arrestato $(\tilde{U}_{n \wedge \nu_0})_{0 \leq n \leq N}$ è una $\mathcal{F}_n$ -martingala.

Dimostrazione. Osserviamo che, poiché $U_N = Z_N$ e dunque $\tilde{U}_N = \tilde{Z}_N$, evidentemente dev'essere $\nu_0 \leq N$, cioè ν_0 è a valori in $\{0, 1, \dots, N\}$. ν_0 è evidentemente un tempo d'arresto. Infatti si può scrivere $\nu_0 = \inf\{n \geq 0; \tilde{U}_n - \tilde{Z}_n = 0\}$ e dunque ν_0 è un tempo d'ingresso (si veda Esempio 3.3.2). Poniamo $\tilde{U}_n^{\nu_0} = \tilde{U}_{n \wedge \nu_0}$ e osserviamo che

$$\tilde{U}_n^{\nu_0} = U_0 + \sum_{j=1}^{n \wedge \nu_0} 1_{\{\nu_0 \geq j\}} \Delta \tilde{U}_j = U_0 + \sum_{j=1}^n 1_{\{\nu_0 \geq j\}} \Delta \tilde{U}_j,$$

dove naturalmente $\Delta \tilde{U}_j = \tilde{U}_j - \tilde{U}_{j-1}$. Si noti che, per ogni j , la v.a. $1_{\{\nu_0 \geq j\}}$ è $\mathcal{F}_{j-1}$ -misurabile, perché $1_{\{\nu_0 \geq j\}} = 1 - 1_{\{\nu_0 \leq j-1\}}$ e $\{\nu_0 \leq j-1\} \in \mathcal{F}_{j-1}$, poiché ν_0 è un tempo d'arresto. Dunque

$$\tilde{U}_{n+1}^{\nu_0} - \tilde{U}_n^{\nu_0} = 1_{\{\nu_0 \geq n+1\}} (\tilde{U}_{n+1} - \tilde{U}_n),$$

se i ha

$$\mathbb{E}^*(\tilde{U}_{n+1}^{\nu_0} - \tilde{U}_n^{\nu_0} | \mathcal{F}_n) = \mathbb{E}^*(1_{\{\nu_0 \geq n+1\}} (\tilde{U}_{n+1} - \tilde{U}_n) | \mathcal{F}_n) \quad (5.4)$$

Ora, sull'insieme $\{\nu_0 \geq n+1\}$, si ha $\tilde{U}_n > \tilde{Z}_n$, dunque

$$1_{\{\nu_0 \geq n+1\}} \tilde{U}_n = 1_{\{\nu_0 \geq n+1\}} \mathbb{E}^*(\tilde{U}_{n+1} | \mathcal{F}_n)$$

e quindi, poiché $\{\nu_0 \geq n+1\} \in \mathcal{F}_n$,

$$\begin{aligned} & \mathbb{E}^*[1_{\{\nu_0 \geq n+1\}} (\tilde{U}_{n+1} - \tilde{U}_n) | \mathcal{F}_n] = \\ &= \mathbb{E}^*[1_{\{\nu_0 \geq n+1\}} (\tilde{U}_{n+1} - \mathbb{E}^*(\tilde{U}_{n+1} | \mathcal{F}_n)) | \mathcal{F}_n] = \\ &= 1_{\{\nu_0 \geq n+1\}} \mathbb{E}^*[\tilde{U}_{n+1} - \mathbb{E}^*(\tilde{U}_{n+1} | \mathcal{F}_n) | \mathcal{F}_n] = 0 \end{aligned}$$

da cui, sostituendo in (5.4) si ha la tesi. □

La Proposizione 5.2.2 ha una conseguenza importante.

Corollario 5.2.3. *Il tempo d'arresto ν_0 soddisfa le seguenti uguaglianze:*

$$U_0 = \tilde{U}_0 = \mathbb{E}^*(\tilde{Z}_{\nu_0}) = \sup_{\nu \in \mathcal{T}_{0,N}} \mathbb{E}^*(\tilde{Z}_\nu),$$

(ricordiamo che $\mathcal{T}_{0,N}$ è l'insieme degli $\mathcal{F}_n$ -tempi d'arresto che prendono valori in $\{0, 1, \dots, N\}$).

Dimostrazione. Per definizione di ν_0 , si ha $\tilde{U}_{\nu_0} = \tilde{Z}_{\nu_0}$. Poiché $(\tilde{U}_n^{\nu_0})_n$ è una martingala e $\tilde{U}_N^{\nu_0} = \tilde{U}_{\nu_0}$, si ha

$$U_0 = \tilde{U}_0 = \tilde{U}_0^{\nu_0} = \mathbb{E}^*(\tilde{U}_N^{\nu_0}) = \mathbb{E}^*(\tilde{U}_{\nu_0}) = \mathbb{E}^*(\tilde{Z}_{\nu_0}).$$

Ora, prendiamo un qualsiasi tempo d'arresto $\nu \in \mathcal{T}_{0,N}$. Per la Proposizione 5.2.1, $(\tilde{U}_n)_n$ è una $\mathbb{P}^*$ -supermartingala, quindi (cfr. Proposizione 3.3.4) anche $(\tilde{U}_n^\nu)_n = (\tilde{U}_{n \wedge \nu})_n$ è una $\mathbb{P}^*$ -supermartingala. Inoltre, essendo $\tilde{U}_n \geq \tilde{Z}_n$ per ogni n , si ha anche $\tilde{U}_\nu \geq \tilde{Z}_\nu$. Ma allora,

$$U_0 = \tilde{U}_0 \geq \mathbb{E}^*(\tilde{U}_N^\nu) = \mathbb{E}^*(\tilde{U}_\nu) \geq \mathbb{E}^*(\tilde{Z}_\nu),$$

da cui segue la tesi. □

Osserviamo che, con un'occhiata più attenta alla dimostrazione del Corollario 5.2.3, un tempo d'arresto τ è ottimale se e solo se succedono insieme due cose:

- 1) il processo arrestato $(\tilde{U}_n^\tau)_n$ è una martingala;
- 2) $\tilde{U}_\tau = \tilde{Z}_\tau$.

5.3 Prezzo e copertura delle opzioni americane

Vediamo ora come si può costruire un portafoglio di copertura di un'opzione americana di payoff $(Z_n)_n$. Abbiamo osservato che il processo $(U_n)_n$, costruito nel paragrafo precedente, è tale che i suoi valori scontati $(\tilde{U}_n)_n$ costituiscono una supermartingala. Sia allora

$$\tilde{U}_n = \tilde{M}_n - \tilde{A}_n$$

la sua decomposizione di Doob (si veda il paragrafo 3.2.3). Dunque $(\tilde{M}_n)_n$ è una martingala, mentre $(\tilde{A}_n)_n$ è un processo crescente predicibile tale che $A_0 = 0$. Consideriamo l'opzione europea di payoff $M_N = (1+r)^N \tilde{M}_N$ e sia $(V_n(\phi))_n$ il suo portafoglio replicante, di cui abbiamo studiato le proprietà nel paragrafo 4.6. Poiché il portafoglio attualizzato $(\tilde{V}_n(\phi))_n$ è una martingala rispetto a $\mathbb{P}^*$, si ha

$$\tilde{V}_n(\phi) = \mathbb{E}^*[\tilde{V}_N(\phi) | \mathcal{F}_n] = \mathbb{E}^*[\tilde{M}_N | \mathcal{F}_n] = \tilde{M}_n.$$

Dunque, moltiplicando per il fattore di sconto $(1+r)^n$,

$$V_n(\phi) = (1+r)^n \tilde{M}_n = U_n + (1+r)^n \tilde{A}_n \geq U_n \geq Z_n$$

Dunque con un capitale pari a

$$U_0 = V_0(\phi) = \mathbb{E}^*[\tilde{M}_N | \mathcal{F}_0],$$

è possibile costituire un portafoglio ammissibile che garantisce ad ogni istante n di coprire il payoff dell'opzione Z_n (infatti, $V_n(\phi) \geq Z_n$ per ogni n). La condizione di assenza di arbitraggio garantisce dunque che il prezzo dell'opzione, c_0 , non può essere superiore a $U_0 = V_0(\phi)$ (cioè, $c_0 \leq U_0$). Mostriamo anzi che la condizione di assenza di arbitraggio implica che $c_0 = U_0$.

Infatti, se il detentore dell'opzione usa la strategia ottimale di arresto per stabilire l'istante in cui esercitare l'opzione e quindi la esercita al tempo ν_0 , allora si ha

$$V_{\nu_0}(\phi) = (1+r)^{\nu_0} \widetilde{M}_{\nu_0} = U_{\nu_0} + (1+r)^{\nu_0} \widetilde{A}_{\nu_0}$$

Il Lemma 5.3.1 qui sotto, implica che $\widetilde{A}_{\nu_0} = 0$. Quindi

$$V_{\nu_0}(\phi) = (1+r)^{\nu_0} \widetilde{M}_{\nu_0} = U_{\nu_0} = Z_{\nu_0}.$$

Dunque in questo caso la strategia ϕ copre esattamente il payoff dell'opzione.

Lemma 5.3.1. *Siano $(X_n)_n$ una supermartingala e $(A_n)_n$ il suo processo crescente associato dato dalla decomposizione di Doob. Siano τ un tempo di arresto e $(X_n^\tau)_n$ la supermartingala ottenuta arrestando $(X_n)_n$ al tempo τ . Allora il processo crescente associato a $(X_n^\tau)_n$ è $(A_n^\tau)_n$, dove $A_n^\tau = A_{n \wedge \tau}$.*

Dimostrazione. Intanto il processo $(A_n^\tau)_n$ è predicibile, poiché si può scrivere

$$A_n^\tau = \sum_{k=0}^{n-1} A_k 1_{\{\tau \geq k\}} + A_n 1_{\{\tau > n-1\}}$$

e nel termine a destra figurano tutte quantità $\mathcal{F}_{n-1}$ -misurabili. Inoltre, posto $M_n = X_n + A_n$, si ha, grazie al teorema di arresto, che $(M_{n \wedge \tau})_n$ è una martingala. Poiché $A_0^\tau = A_0 = 0$ e

$$M_{n \wedge \tau} = X_{n \wedge \tau} + A_{n \wedge \tau} = X_n^\tau + A_n^\tau,$$

per l'unicità del processo crescente nella decomposizione di Doob, la tesi è dimostrata. $\square$

Tornando alla questione che ci interessava, il processo crescente $(A_n^{\nu_0})_n$ associato a $(U_n^{\nu_0})_n$ è dato da $A_n^{\nu_0} = A_{n \wedge \nu_0}$. Poiché, d'altra parte, $(U_n^{\nu_0})_n$ è una martingala il suo processo crescente associato è nullo. Dunque $A_{\nu_0} = 0$.

Ricapitolando, il prezzo all'istante 0 è dato dalla formula (5.1) o equivalentemente, in maniera più esplicita, dalla (5.2) con $n = 0$, oppure da

$$U_0 = \mathbb{E}^*((1+r)^{-\nu_0} Z_{\nu_0})$$

che poi è esattamente quanto ci aspettavamo, come descritto nell'introduzione a questo capitolo. Questo risultato si generalizza facilmente per

il calcolo del prezzo dell'opzione in ogni istante $n = 0, 1, \dots, N$. Riassumiamo qui di seguito il risultato generale, senza riportarne la dimostrazione (per la quale, basta semplicemente riadattare quelle appena viste).

Proposizione 5.3.2. *Il prezzo di un'opzione americana al tempo n , $n = 0, 1, \dots, N$ è dato dalla formula (5.2) o equivalentemente da*

$$U_n = \mathbb{E}^*[(1+r)^{-(\nu_n-n)} Z_{\nu_n} | \mathcal{F}_n] = \sup_{\nu \in \mathcal{T}_{n,N}} \mathbb{E}^*[(1+r)^{-(\nu-n)} Z_\nu | \mathcal{F}_n],$$

dove $\mathcal{T}_{n,N}$ denota l'insieme dei tempi d'arresto a valori in $\{n, \dots, N\}$ e $\nu_n \in \mathcal{T}_{n,N}$ è definito da

$$\nu_n = \inf\{j \geq n; U_j = Z_j\}.$$

Intuitivamente, un'opzione americana di processo di payoff $(Z_n)_n$ deve costare di più dell'opzione europea di payoff $h = Z_N$, poiché il detentore può esercitare l'opzione in un qualsiasi momento tra l'emissione e la maturità N . Quest'affermazione è precisata dalla seguente

Proposizione 5.3.3. *Siano rispettivamente C_n il prezzo di un'opzione americana di processo di payoff $(Z_n)_{0 \leq n \leq N}$ e c_n il prezzo dell'opzione europea associata, cioè di payoff $h = Z_N$ e maturità N . Allora,*

$$C_n \geq c_n, \quad \text{per ogni } n = 0, 1, \dots, N.$$

Inoltre, se $c_n \geq Z_n$ per ogni n allora

$$C_n = c_n, \quad \text{per ogni } n = 0, 1, \dots, N.$$

Dimostrazione. Osserviamo che $C_N = Z_N = c_N$. Siano $(\tilde{C}_n)_n$ e $(\tilde{c}_n)_n$ i processi scontati. Ricordando che $(\tilde{C}_n)_n$ è una $\mathbb{P}^*$ -supermartingala e $(\tilde{c}_n)_n$ una $\mathbb{P}^*$ -martingala, si ha

$$\tilde{C}_n \geq \mathbb{E}^*(\tilde{C}_N | \mathcal{F}_n) = \mathbb{E}^*(\tilde{c}_N | \mathcal{F}_n) = \tilde{c}_n$$

per ogni $n = 0, 1, \dots, N$. Dunque, $C_n \geq c_n$.

Supponiamo ora che si abbia $c_n \geq Z_n$. Allora $(\tilde{c}_n)_n$, che è una martingala, è una $\mathbb{P}^*$ -supermartingala che domina il payoff scontato $(\tilde{Z}_n)_n$. Poiché $(\tilde{C}_n)_n$ è l'involuppo di Snell di $(\tilde{Z}_n)_n$, necessariamente dev'essere $\tilde{c}_n \geq \tilde{C}_n$, da cui segue che $\tilde{c}_n = \tilde{C}_n$ e quindi $c_n = C_n$, per ogni possibile n . $\square$

Esempio 5.3.4. (*Prezzi delle call europea ed americana*) Consideriamo una opzione call sul titolo di prezzo $(S_n^1)_n$:

$$Z_n = Z_n^{\text{call}} = (S_n^1 - K)_+.$$

Indichiamo con C_n il prezzo della call americana al tempo n e con c_n il prezzo della call europea:

$$c_n = \mathbb{E}^*[(1+r)^{-(N-n)}(S_N - K)_+ | \mathcal{F}_n].$$

Mostriamo che in questo caso si ha

$$c_n = C_n.$$

Per la Proposizione 5.3.2, basta far vedere che $c_n \geq Z_n = (S_n^1 - K)_+$ per ogni $n = 0, 1, \dots, N$. Ricordiamo che il prezzo scontato $\tilde{c}_n = (1+r)^{-n}c_n$ è una martingala. Dunque, applicando la disuguaglianza di Jensen alla funzione convessa $x \rightarrow x^+$,

$$\begin{aligned} \tilde{c}_n &= \mathbb{E}^*[(1+r)^{-N}(S_N - K)_+ | \mathcal{F}_n] \geq \left(\mathbb{E}^*[(1+r)^{-N}(S_N - K) | \mathcal{F}_n] \right)_+ \\ &= \left(\mathbb{E}^*[\tilde{S}_N | \mathcal{F}_n] - K(1+r)^{-N} \right)_+ = (\tilde{S}_n - K(1+r)^{-N})_+ \end{aligned}$$

Moltiplicando per $(1+r)^n$, si ottiene

$$c_n \geq (1+r)^n(\tilde{S}_n - K(1+r)^{-N})_+ = (S_n - K(1+r)^{-(N-n)})_+ \geq (S_n - K)_+$$

dove abbiamo usato il fatto che $(1+r)^{-(N-n)} \leq 1$. Osserviamo infine che, dalla Proposizione 5.3.2, si ha

$$C_n = \sup_{\nu \in \mathcal{T}_{n,N}} \mathbb{E}^*[(1+r)^{-(\nu-n)} Z_\nu | \mathcal{F}_n] = \mathbb{E}^*[(1+r)^{-(N-n)} Z_N | \mathcal{F}_n] = c_n$$

dunque il tempo (ottimale) di esercizio è sempre, nel caso della call, $\nu_n = N$: la cosa migliore è esercitare direttamente a maturità.

Vedremo che la situazione pu' essere diversa per opzioni diverse dalla call e, in particolare, per le put, delle quali studieremo il comportamento qualitativo nell'ambito del modello CRR nel prossimo paragrafo.

5.4 La put americana nel modello CRR

In questo paragrafo, consideriamo il modello CRR. Ricordiamo che in tal caso il mercato consiste del titolo non rischioso, $S_n^0 = (1+r)^n$, e di un solo titolo rischioso, di prezzo

$$S_n = S_0 \cdot T_1 \cdots T_n.$$

Supponiamo che il mercato sia privo di arbitraggio e completo, cioè $r \in (a, b)$. Sotto la misura equivalente di martingala $\mathbb{P}^*$, le variabili aleatorie $T_1, \dots, T_n$ sono indipendenti ed identicamente distribuite e tali che

$$\mathbb{P}^*(T_1 = 1+a) = p = 1 - \mathbb{P}^*(T_1 = 1+b), \quad \text{dove } p = \frac{r-a}{b-a}.$$

Consideriamo una opzione put americana, di maturità N e prezzo d'esercizio K :

$$Z_n = Z_n^{\text{put}} = (K - S_n)_+.$$

Proposizione 5.4.1. *Il prezzo della put americana al tempo n è dato da*

$$P_n = P_{\text{am}}(n, S_n),$$

dove la funzione $P_{\text{am}}(n, x)$ è definita da

$$\begin{aligned} P_{\text{am}}(N, x) &= (K - x)_+ \quad \text{e per } n = N - 1, N - 2, \dots, 0 \\ P_{\text{am}}(n, x) &= \max \left((K - x)_+, \frac{1}{1+r} f(n+1, x) \right), \end{aligned}$$

essendo

$$f(n+1, x) = p P_{\text{am}}(n+1, x(1+a)) + (1-p) P_{\text{am}}(n+1, x(1+b)).$$

Dimostrazione. Da (5.2), si ha

$$P_N = Z_N = (K - S_N)_+ = P_{\text{am}}(N, S_N).$$

Al tempo $n = N - 1$, sempre da (5.2) si ha

$$P_{N-1} = \max \left((K - S_{N-1})_+, \frac{1}{1+r} \mathbb{E}^*(P_N | \mathcal{F}_{N-1}) \right).$$

Osserviamo che

$$\mathbb{E}^*(P_N | \mathcal{F}_{N-1}) = \mathbb{E}^*(P_{\text{am}}(N, S_N) | \mathcal{F}_{N-1}) = \mathbb{E}^*(P_{\text{am}}(N, S_{N-1} \cdot T_N) | \mathcal{F}_{N-1}).$$

Ora, poiché S_{N-1} è $\mathcal{F}_{N-1}$ -misurabile e T_N è indipendente da $\mathcal{F}_{N-1}$, l'ultima media condizionata vale

$$\mathbb{E}^*(P_{\text{am}}(N, S_{N-1} \cdot T_N) | \mathcal{F}_{N-1}) = \mathbb{E}^*(P_{\text{am}}(N, x \cdot T_N)) \Big|_{x=S_{N-1}}.$$

Posto allora

$$f(N, x) = \mathbb{E}^*(P_{\text{am}}(N, x \cdot T_N)),$$

segue che

$$\mathbb{E}^*(P_N | \mathcal{F}_{N-1}) = f(N, S_{N-1}),$$

da cui si ottiene $P_{N-1} = P_{\text{am}}(N-1, S_{N-1})$, dove

$$P_{\text{am}}(N-1, x) = \max \left((K - x)_+, \frac{1}{1+r} f(N, x) \right).$$

Ora,

$$f(N, x) = \mathbb{E}^*(P_{\text{am}}(N, x \cdot T_N)) = P_{\text{am}}(N, x \cdot (1+a)) p + P_{\text{am}}(N, x \cdot (1+b)) (1-p),$$

dunque la tesi è vera per $n = N - 1$. Procedendo analogamente per $n = N - 2, \dots, 0$, si ottiene il risultato finale

□

Osservazione 5.4.2. La funzione $P_{\text{am}}(0, x)$ introdotta nella Proposizione 5.4.1 si può anche rappresentare nel modo seguente:

$$P_{\text{am}}(0, x) = \sup_{\nu \in \mathcal{T}_{0,N}} \mathbb{E}^* \left((1+r)^{-\nu} (K - x V_\nu)_+ \right) \quad (5.5)$$

dove $V_0 = 1$ e $V_n = T_1 \cdots T_n$, $n = 1, \dots, N$.

Infatti, dalla Proposizione 5.3.2 si ha

$$\begin{aligned} P_{\text{am}}(0, S_0) = P_0 &= \sup_{\nu \in \mathcal{T}_{0,N}} \mathbb{E}^* \left((1+r)^{-\nu} (K - S_\nu)_+ \right) \\ &= \sup_{\nu \in \mathcal{T}_{0,N}} \mathbb{E}^* \left((1+r)^{-\nu} (K - S_0 V_\nu)_+ \right) \end{aligned}$$

per ogni scelta di S_0 .

La rappresentazione (5.5) consente di effettuare uno studio qualitativo del prezzo dell'opzione all'istante iniziale. Queste informazioni, riassunte nella Proposizione 5.4.3, consentono di tracciare il grafico della funzione $x \mapsto P_{\text{am}}(0, x)$, come riportato in Figura 5.1.

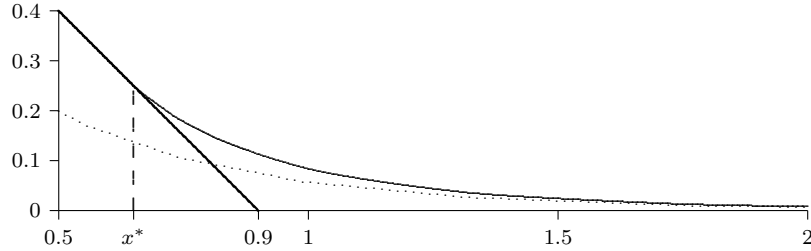


Figura 5.1 Andamento del prezzo al tempo 0 di una opzione put americana, al variare del prezzo del titolo di base (in ascisse). I valori sono $a = -.1$, $b = .2$, $r = .02$, $K = 0.9$, $N = 20$. Qui il valore x^* è dell'ordine di 0.65. Il prezzo si annulla per $x > K(1+a)^{-N} = 7.4$. A puntini è disegnato l'andamento delle corrispondente opzione europea.

Proposizione 5.4.3. 1. La funzione $x \mapsto P_{\text{am}}(0, x)$ è continua, non crescente e convessa.

2. Supponiamo $a < 0$. Esiste un punto $x^* \in [0, K)$ tale che

- (a) per $x \in [0, x^*]$, $P_{\text{am}}(0, x) = (K - x)_+$;
- (b) per $x \in (x^*, K/(1+a)^N)$, $P_{\text{am}}(0, x) > (K - x)_+$;
- (c) per $x \geq K/(1+a)^N$, $P_{\text{am}}(0, x) = 0$.

Dimostrazione. 1. Dalla Proposizione 5.4.1, segue immediatamente che $x \mapsto P_{\text{am}}(N, x) = (K - x)_+$ è continua. Dunque, anche $x \mapsto f(N, x) = pP_{\text{am}}(N, x(1+a)) + (1-p)P_{\text{am}}(N, x(1+b))$ è continua e quindi anche $x \mapsto P_{\text{am}}(N-1, x) = \max((K-x)_+, (1+r)^{-1}f(N, x))$ lo è (essendo il max tra funzioni continue). Procedendo analogamente, si prova che $x \mapsto P_{\text{am}}(n, x)$ è continua per ogni $n = N-2, \dots, 1, 0$.

Le altre due proprietà seguono direttamente dalla rappresentazione (5.5) e dal fatto che $x \mapsto (K - x V_\nu)_+$ è non crescente e convessa. Ad esempio, per quanto riguarda la convessità, preso $\alpha \in [0, 1]$ si ha:

$$\begin{aligned} & \mathbb{E}^* \left((1+r)^{-\nu} (K - (\alpha x_1 + (1-\alpha)x_2) V_\nu)_+ \right) \\ & \leq \alpha \mathbb{E}^* \left((1+r)^{-\nu} (K - x_1 V_\nu)_+ \right) + (1-\alpha) \mathbb{E}^* \left((1+r)^{-\nu} (K - x_2 V_\nu)_+ \right) \end{aligned}$$

e passando al sup per $\nu \in \mathcal{T}_{0,N}$ si ottiene la tesi.

2. Poniamo $g(x) = P_{\text{am}}(0, x) - (K - x)$. Poiché $P_{\text{am}}(0, x) \geq (K - x)_+$, $g(x) \geq 0$. In particolare $g(0) = 0$, perché, usando (5.5),

$$0 \leq g(0) = \sup_{\nu \in \mathcal{T}_{0,N}} \mathbb{E}^* \left(\underbrace{(1+r)^{-\nu}}_{\leq 1, \forall \nu \geq 0} K \right) - K \leq K - K = 0.$$

Invece, $g(K) > 0$. Infatti,

$$g(K) = K \sup_{\nu \in \mathcal{T}_{0,N}} \mathbb{E}^* \left((1+r)^{-\nu} (1 - V_\nu)_+ \right) \geq K \mathbb{E}^* \left((1+r)^{-1} (1 - V_1)_+ \right)$$

(nel secondo passaggio si è scelto $\nu = 1$). Ora, $V_1 = T_1$ e

$$\mathbb{E}^* \left((1+r)^{-1} (1 - V_1)_+ \right) = (1+r)^{-1} \left((-a)_+ p + (-b)_+ (1-p) \right) = (1+r)^{-1} ap > 0$$

quindi $g(K) > 0$.

Dividiamo ora lo studio della funzione g in due casi: (i) $x \in [0, K]$ e (ii) $x > K$.

(i) Se $x \in [0, K]$, allora $g(x) = P_{\text{am}}(0, x) - (K - x)$ è una funzione continua e convessa (essendo somma di una funzione convessa continua e di una funzione lineare). Poniamo allora

$$x^* = \inf \{ x > 0 : g(x) > 0 \}.$$

Ovviamente, $x^* < K$ (perché $g(K) > 0$) ed inoltre si ha: $g(x) = 0$ per ogni $x \in [0, x^*]$ e $g(x) > 0$ per ogni $x \in (x^*, K)$. Infatti, se per assurdo esistesse un $x \leq x^*$ tale che $g(x) > 0$ allora, dalla definizione di x^* , dovrebbe anche essere $x \geq x^*$, il che non è vero a meno che non si sia scelto $x = x^*$. Osservando però che, dal teorema della permanenza del segno (ricordiamo che g è continua), dev'essere $g(x^*) = 0$, segue allora che in effetti $g(x) = 0$

per ogni $x \in [0, x^*]$. Prendiamo ora un $x \in (x^*, K)$ e mostriamo che $g(x) > 0$. Osserviamo intanto che, sempre dalla definizione di x^* , deve esistere una successione $\{x_n^*\}_n$ tale che $x_n^* \downarrow x^*$ e $g(x_n^*) > 0$. Ora, se $x > x^*$ allora $x > x_n^*$ per qualche n , dunque $x_n^* = tx + (1-t)x^*$ per qualche $t \in (0, 1)$. Usando la convessità di g , si ha

$$0 < g(x_n^*) = g(tx + (1-t)x^*) \leq tg(x) + (1-t)g(x^*) = tg(x),$$

il che prova che $g(x) > 0$ per ogni $x \in (x^*, K]$.

(ii) Se invece $x > K$, allora $g(x) = P_{\text{am}}(0, x)$ è una funzione non negativa e, dal punto 1., anche non crescente. Poiché $g(K) > 0$, nell'intervallo in questione g è positiva, tende a decrescere e una volta raggiunto lo zero, si mantiene identicamente nulla.

Ricapitolando, quanto visto in (i) e (ii) prova l'asserzione (a). Per verificare (b) e (c), basta far vedere che $\bar{x} = K/(1+a)^N$ è il più piccolo punto $> K$ in cui $x \mapsto P_{\text{am}}(0, x)$ si annulla. Cerchiamo allora $\bar{x}$ tale che $P_{\text{am}}(0, \bar{x}) = 0$. Dalla (5.5) dev'essere

$$(K - \bar{x} V_\nu)_+ = 0 \quad \text{per ogni } \nu \in \mathcal{T}_{0,N},$$

dunque $\bar{x} \geq K/V_\nu$ per ogni $\nu \in \mathcal{T}_{0,N}$. Allora

$$\bar{x} = \max_{\Omega} \max_{\nu \in \mathcal{T}_{0,N}} \frac{K}{V_\nu} = \frac{K}{\min_{\Omega} \min_{\nu \in \mathcal{T}_{0,N}} V_\nu}.$$

Ora, poiché $V_\nu = T_1 \cdots T_\nu$ e $T_n \in \{1+a, 1+b\}$, allora

$$\min_{\Omega} \min_{\nu \in \mathcal{T}_{0,N}} V_\nu = \min_{n \in \{0, \dots, N\}} (1+a)^n = (1+a)^N$$

perché, ricordiamo $0 < 1+a < 1$. Dunque $\bar{x} = K/(1+a)^N$ e la tesi è dimostrata.

□

Esercizio 5.1. Implementare al calcolatore il principio di programmazione dinamica della Proposizione 5.4.1 e disegnare il grafico della funzione $x \mapsto P_{\text{am}}(0, x)$ che dà il prezzo della put americana al tempo 0, come in Figura 5.1, a pagina 115. Si prendano, a titolo di esempio, i parametri a, b e r come in Figura 5.1.

Capitolo 6

Simulazione e metodi Monte Carlo per la finanza

In questo capitolo, dopo aver introdotto i metodi Monte Carlo e averne visto possibili utilizzi alla soluzione di alcuni problemi di interesse in finanza, saranno proposti esempi di applicazione sotto forma di esercizi al calcolatore. In particolare, sarà richiesto di simulare il processo CRR. Come vedremo, lo strumento fondamentale è la possibilità di generare variabili aleatorie su $(0, 1)$. Molti linguaggi di programmazione hanno in libreria una procedura che genera tali v.a. Ad esempio, in Pascal esiste una procedura di impiego immediato, che si chiama `Random`: il comando

`Z:=Random`

assegna a `Z` un numero casuale su $(0, 1)$. In C, invece, non esiste un comando immediato ma occorre scrivere una procedura apposita. Ad esempio,

```
double nrandom(void) {  
    return ((double)rand()/(double)RAND_MAX);  
}
```

restituisce un numero che in buona approssimazione è a caso su $(0, 1)$. Esistono comunque altri generatori aleatori più sofisticati (soprattutto in grado di generare “tanti” numeri casuali: come vedremo, questo è piuttosto importante per i metodi Monte Carlo), ad esempio KNUTH¹.

I lettori interessati possono trovare in Glasserman [4] un ottimo testo di riferimento per i metodi numerici-Monte Carlo per le applicazioni in finanza.

6.1 Metodi Monte Carlo: generalità

Il metodo Monte Carlo è un metodo **stocastico** per il calcolo numerico di quantità **deterministiche**. L'idea è la seguente.

¹All'indirizzo <http://www.mat.uniroma2.it/~caramell/did.0607/pf.htm> è possibile scaricare file di testo contenenti queste ed altre procedure.

Supponiamo di voler calcolare una quantità α e per di più supponiamo che α si possa rappresentare come la media di una v.a. X :

$$\alpha = \mathbb{E}(X).$$

Ovviamente, stiamo supponendo anche di lavorare su un opportuno spazio di probabilità $(\Omega, \mathcal{F}, \mathbb{P})$, dove la v.a. $X : (\Omega, \mathcal{F}) \rightarrow (\mathbb{R}^d, \mathcal{B}(\mathbb{R}^d))$ è definita e il simbolo $\mathbb{E}$ denota l'aspettazione sotto $\mathbb{P}$. Va da sé che X deve essere integrabile. Nel seguito, indicheremo con Λ la legge di² X .

Immaginiamo che X sia anche di quadrato integrabile e sia $\{X^q\}_q$ una successione di v.a. indipendenti, tutte con legge Λ . Allora, per la Legge dei Grandi Numeri, si ha:

$$\hat{\alpha}_Q := \frac{1}{Q} \sum_{q=1}^Q X^q \xrightarrow{Q \rightarrow \infty} \mathbb{E}(X) = \alpha.$$

Dunque, un modo per approssimare α è quello di calcolare la media empirica di un buon numero di v.a. indipendenti con legge uguale alla legge di X . Questo è il metodo Monte Carlo. Osserviamo che:

- nella pratica, questo metodo si implementa in poche righe di programma al calcolatore se è possibile simulare v.a. con legge Λ ;
- la stima $\hat{\alpha}_Q$ di α è soggetta a due errori: il primo conseguente al fatto di fissare un valore (seppur grande) di Q , il secondo dovuto al fatto che $\hat{\alpha}_Q$ è una **variabile aleatoria** (il che, rozzamente, significa che facendo girare il programma più volte, anche con lo stesso valore per Q , si possono ottenere stime diverse $\hat{\alpha}_Q$ per α).

La “bontà” della stima (ovvero, la velocità di convergenza) si può studiare usando il Teorema del Limite Centrale, che riportiamo nella seguente, qui particolarmente utile, versione:

$$\sqrt{Q} (\hat{\alpha}_Q - \alpha) \xrightarrow{\mathcal{L}} N(0, C)$$

dove il simbolo $\xrightarrow{\mathcal{L}}$ denota la convergenza in legge, $N(0, C)$ è la legge gaussiana su $\mathbb{R}^d$ di media 0 e matrice di covarianza C , e C qui sta per la matrice di covarianza associata al vettore aleatorio $X = (X_1, \dots, X_d)$: $C_{ij} = \text{Cov}(X_i, X_j)$. In particolare, se Q è grande allora vale la seguente approssimazione:

$$\hat{\alpha}_Q \cong \alpha + \frac{1}{\sqrt{Q}} C^{1/2} Z, \quad (6.1)$$

con Z v.a. gaussiana standard su $\mathbb{R}^d$ e $C^{1/2}$ denota una matrice quadrata A di dimensione d tale che³ $AA^t = C$. Ciò significa che l'errore **aleatorio** che

²Cioè, $\Lambda(A) = \mathbb{P}(X \in A)$, per $A \in \mathcal{B}(\mathbb{R}^d)$.

³Poiché C è simmetrica e definita positiva (non-negativa), è noto che una tale matrice A esiste.

si compie approssimando α con $\hat{\alpha}_Q$ è dell'ordine $\frac{1}{\sqrt{Q}} C^{1/2} Z$, con Z gaussiana standard. In altre parole, l'ordine di convergenza di un metodo Monte Carlo è del tipo $1/\sqrt{Q}$, e cioè: è lento! Ciò nonostante, occorre sottolineare che il metodo è estremamente flessibile ed inoltre, come abbiamo appena visto, la velocità di convergenza è indipendente dalla dimensione d (e quest'ultima è, forse, la peculiarità del metodo).

Infine, grazie alla (6.1), è possibile individuare un intervallo di confidenza (IC) approssimato per α , a livello δ desiderato. Infatti, supponiamo per semplicità che $d = 1$, e quindi

$$\hat{\alpha}_Q \cong \alpha + \frac{\sigma}{\sqrt{Q}} Z, \quad (6.2)$$

con $Z \sim N(0, 1)$. Fissato un livello δ , un IC bilatero per α , a livello δ , ha quindi estremi

$$\hat{\alpha}_Q \pm \frac{\sigma}{\sqrt{Q}} \phi_{\frac{1+\delta}{2}}, \quad (6.3)$$

dove $\phi_{\frac{1+\delta}{2}}$ denota il quantile di ordine $(1 + \delta)/2$ della legge gaussiana standard. Ad esempio, per un IC al 95% (cioè $\delta = 0.95$), $\phi_{\frac{1+\delta}{2}} = \phi_{0.975} = 1.96$; se invece $\delta = 0.9$ allora $\phi_{\frac{1+\delta}{2}} = \phi_{0.95} = 1.645$.

Volendo usare un metodo Monte Carlo, nella pratica è cruciale conoscere sia la stima $\hat{\alpha}_Q$ di α sia consegnare all'utente un IC, tipicamente al 95% (cioè, $\delta = 0.95$). Osserviamo però che nella (6.3) compare σ^2 , cioè la varianza di X , che spesso non si conosce. In tal caso, si sostituisce nella (6.3) la tipica stima $\hat{\sigma}_Q^2$ (non distorta) per σ^2 :

$$\hat{\sigma}_Q^2 = \frac{1}{Q-1} \sum_{q=1}^Q (X^q - \hat{\alpha}_Q)^2 = \frac{1}{Q-1} \sum_{q=1}^Q (X^q)^2 - \hat{\alpha}_Q^2,$$

cosicché l'IC approssimato al 95% ha estremi

$$\hat{\alpha}_Q \pm 1.96 \sqrt{\frac{\hat{\sigma}_Q^2}{Q}}.$$

6.2 Simulazione del modello CRR

La simulazione del modello CRR è abbastanza semplice e si basa sulla generazione di v.a. bernoulliane. Nel seguito, assumeremo di essere sotto la misura neutrale al rischio $\mathbb{P}^*$, e indicheremo, al solito, con $\mathbb{E}^*$ l'aspettazione sotto $\mathbb{P}^*$.

Ricordiamo che il modello CRR assume un unico sottostante rischioso $(S_n)_n$ in aggiunta al sottostante non rischioso, che si evolve seguendo la legge classica

$$S_n^0 = (1+r)^n, \quad n = 0, \dots, N.$$

Ricordiamo che, nel modello CRR, si ha

$$S_n = S_0 \cdot T_1 \cdots T_n, \quad n = 0, \dots, N$$

(con l'ovvia condizione $T_1 \cdots T_n = 1$ quando $n = 0$), dove le v.a. (T_n) sono i.i.d. con legge

$$\mathbb{P}^*(T_n = 1+a) = p^* = 1 - \mathbb{P}^*(T_n = 1+b), \quad \text{con } p^* = \frac{b-r}{b-a}.$$

Ovviamente, stiamo assumendo $-1 < a < r < b$. Per ragioni legate al comportamento asintotico del modello CRR (si veda il Paragrafo 4.6.3), i parametri r , a e b si scelgono rispettivamente come in (4.23) e (4.25), e cioè:

$$\begin{aligned} r &\equiv r_N = \frac{RT}{N} \\ 1+a &\equiv 1+a_N = \left(1 + \frac{RT}{N}\right) e^{-\sigma \sqrt{T/N}} \\ 1+b &\equiv 1+b_N = \left(1 + \frac{RT}{N}\right) e^{+\sigma \sqrt{T/N}} \end{aligned}$$

dove, ricordiamo, R rappresenta il tasso di interesse uniperiodale (che per N grande si può approssimare con il tasso istantaneo), T è la maturità, N denota il numero di date in cui il prezzo del sottostante viene monitorato su $[0, T]$ ed infine σ è la volatilità del titolo.

Posto

$$U_n = \frac{1}{\sigma \sqrt{t/N}} \log \left(\left(1 + \frac{RT}{N}\right)^{-1} T_n \right)$$

si ottiene banalmente che

1. sotto $\mathbb{P}^*$, $U_1, \dots, U_N$ sono i.i.d. con legge

$$\mathbb{P}^*(U_n = -1) = p^* = 1 - \mathbb{P}^*(U_n = +1);$$

2. per ogni $n \leq N$, $T_n = (1 + \frac{RT}{N}) e^{\sigma \sqrt{t/N} U_n}$ e quindi

$$S_n = S_0 \left(1 + \frac{RT}{N}\right)^n \exp \left(\sigma \sqrt{t/N} \sum_{i=1}^n U_i \right) \quad (6.4)$$

Dalla (6.4), la generazione di una traiettoria di $(S_n)_{n \leq N}$ segue immediatamente dalla generazione di N v.a. i.i.d. che assumono valori -1 e $+1$ rispettivamente con probabilità p^* e $1 - p^*$. Per simulare una tale v.a., un algoritmo semplice è quello suggerito dalla seguente proposizione, la cui dimostrazione è immediata.

Proposizione 6.2.1. *Sia Z una v.a. uniforme su $(0, 1)$ e sia $p \in (0, 1)$. Allora*

$$U = 1 - 2 \mathbf{1}_{\{Z \leq p\}} = \begin{cases} -1 & \text{con probabilità } p \\ +1 & \text{con probabilità } 1 - p \end{cases}$$

Dunque, per simulare le v.a. i.i.d. $U_1, \dots, U_N$ basta generare N v.a. i.i.d. uniformi su $(0, 1)$ e applicare la Proposizione 6.2.1 con $p = p^*$. Richiamando successivamente una qualsiasi procedura in grado di simulare uniformi su $(0, 1)$, si ottengono generazioni indipendenti.

La simulazione del modello CRR è a questo punto immediata a partire dalla (6.4).

6.3 Prezzo di opzioni europee con metodi Monte Carlo

I metodi Monte Carlo trovano nella finanza una naturale applicazione per la risoluzione numerica dei problemi di prezzaggio e copertura delle opzioni. Infatti, supponiamo di essere in un modello di mercato completo, dove $\mathbb{E}^*$ denota l'aspettazione valutata sotto la misura di rischio neutro. La formula del prezzo di un'opzione europea di payoff h e maturità T è data da

$$\text{prezzo al tempo } n = (1 + r)^{-(N-n)} \mathbb{E}^*(h \mid \mathcal{F}_n),$$

e in particolare

$$\alpha = (1 + r)^{-N} \mathbb{E}^*(h)$$

è il prezzo dell'opzione al tempo iniziale. La rappresentazione del prezzo in termini di media suggerisce immediatamente l'uso di un metodo Monte Carlo. Infatti, come visto nel Paragrafo 6.1, il prezzo α si approssima con

$$\hat{\alpha}_Q := (1 + r)^{-N} \frac{1}{Q} \sum_{q=1}^Q h^q$$

dove $(h^q)_{1 \leq q \leq Q}$ sono Q copie indipendenti del payoff h sotto la misura neutrale al rischio $\mathbb{P}^*$. Nella pratica, le h^q vengono simulate al calcolatore, nel modo seguente.

Come prima cosa, occorre fissare un modello per i prezzi di mercato. Qui, supporremo di essere nel contesto del modello CRR. Inoltre, il payoff h è, per ipotesi, non negativo e $\mathcal{F}_N$ -misurabile, dunque nel nostro caso è funzione dei prezzi osservati fino a maturità. In altre parole, esiste una opportuna funzione (misurabile) $G : \mathbb{R}^N \rightarrow R_+$ tale che

$$h = G(S_1, \dots, S_N).$$

Ricordiamo infatti che, ad esempio, per una call si ha $G(x_1, \dots, x_N) \equiv F(x_N) = (x_N - K)_+$, quindi c'è dipendenza dall'evoluzione del prezzo sottostante solo a maturità. Ma non è sempre così: ad esempio, per una call asiatica (si veda anche il prossimo Paragrafo 6.3.2) si ha

$$G(x_1, \dots, x_N) = \left(\frac{1}{N+1} \sum_{n=0}^N x_n - K \right)_+.$$

Comunque, la cosa importante è che h è funzione del sottostante: h è quindi simulabile se è possibile simulare il sottostante. Ma questo è vero: detta $(S_n^q)_{0 \leq n \leq N}$ una generazione della traiettoria del sottostante come visto nel Paragrafo 6.2, una simulazione h^q del payoff si ottiene ponendo

$$h^q = G(S_1^q, \dots, S_N^q).$$

Quindi, riassumendo, il prezzo al tempo 0 si calcola numericamente con un metodo Monte Carlo tramite la quantità

$$(1+r)^{-N} \frac{1}{Q} \sum_{q=1}^Q G(S_1^q, \dots, S_N^q).$$

6.3.1 Call e put standard

Supponiamo di avere a che fare con un'opzione call oppure put, cioè

$$h = h_{\text{call}} = (S_T - K)_+ \quad \text{oppure} \quad h = h_{\text{put}} = (K - S_T)_+$$

In tal caso, sappiamo che la funzione P che dà il prezzo al tempo n è data dalla formula (4.16), che diviene

$$\begin{aligned} P_{\text{call}}(n, x) &= \sum_{j=0}^{N-n} \binom{N-n}{j} p^{*j} (1-p^*)^{N-n-j} \left(x(1+a)^j (1+b)^{N-n-j} - K \right)_+ \\ P_{\text{put}}(n, x) &= \sum_{j=0}^{N-n} \binom{N-n}{j} p^{*j} (1-p^*)^{N-n-j} \left(K - x(1+a)^j (1+b)^{N-n-j} \right)_+ \end{aligned} \quad (6.5)$$

Anche la copertura ha una formula esplicita, che si calcola a partire dalla delta definita in (4.20).

Se però N è grande, lavorare con coefficienti binomiali può essere complicato. In questo caso, è più opportuno usare la formula di ricorrenza “all'indietro” per la funzione prezzo, vista nella (4.22), e cioè

$$\begin{aligned} P_{\text{call}}(N, x) &= (x - K)_+ \text{ e per } n = N-1, N-2, \dots, 0, \\ P_{\text{call}}(n, x) &= \frac{1}{1+r} \left((1-p^*) P_{\text{call}}(n+1, x(1+b)) + p^* P_{\text{call}}(n+1, x(1+a)) \right) \end{aligned} \quad (6.6)$$

Per quanto riguarda la put, si ha analogamente

$$P_{\text{put}}(N, x) = (K - x)_+ \text{ e per } n = N - 1, N - 2, \dots, 0,$$

$$P_{\text{put}}(n, x) = \frac{1}{1 + r} \left((1 - p^*) P_{\text{put}}(n + 1, x(1 + b)) + p^* P_{\text{put}}(n + 1, x(1 + a)) \right) \quad (6.7)$$

In effetti, qui non occorre appellarsi a metodi numerici: le formule esistono, basta usarle! Ciò nonostante, proponiamo il seguente esercizio.

Esercizio 6.1. *Calcolare numericamente, con il metodo Monte Carlo, il prezzo di una call e di una put standard in 0, con IC al 95%, e studiare empiricamente la convergenza ai risultati forniti dalle formule esatte quando il numero di simulazioni Q tende a $+\infty$.*

A titolo di esempio, si considerino i seguenti parametri: $S_0 = K = 100$, $T = 1$ anno suddiviso in monitoraggi trimestrali, quadrimestrali, mensili e giornalieri; tasso d'interesse R pari al 5%; volatilità σ pari al 20%.

6.3.2 Opzioni call/put asiatiche

Un'opzione asiatica ha payoff h che dipende dalla media temporale del sottostante su tutto l'intervallo di interesse $[0, T]$:

$$h = F\left(\frac{1}{N+1} \sum_{n=0}^N S_n\right)$$

Ad esempio, una call asiatica a maturità T e con prezzo di esercizio K ha payoff

$$h = \left(\frac{1}{N+1} \sum_{n=0}^N S_n - K \right)_+.$$

Ricordiamo che le formule esplicite per call/put standard riassunte nel Paragrafo 6.3.1 si basano sul fatto che il payoff dell'opzione è funzione del prezzo del sottostante solo nell'istante finale ($h = F(S_N)$), mentre qui dipende dall'evoluzione dei prezzi durante tutto l'intervallo temporale di interesse. Dunque, il caso delle opzioni asiatiche è piuttosto interessante perché non esistono formule esatte per calcolarne prezzo e/o copertura, e i metodi numerici sono gli unici praticabili.

Esercizio 6.2. *Usando semplicemente il payoff e la rappresentazione in forma di aspettazione del prezzo, calcolare numericamente, con il metodo Monte Carlo, il prezzo in 0 di una call e di una put asiatica, con IC al 95%. Si considerino gli stessi parametri suggeriti nell'Esercizio 6.1. Si confronti poi il risultato ottenuto con il prezzo dell'analoga opzione call e put standard. Qual è più costosa all'aumentare di σ ?*

6.3.3 Opzioni call/put con barriere

Un'opzione call/put con barriere prevede l'annullamento o l'attivazione del payoff classico call/put qualora una o più barriere vengano raggiunte o meno dal prezzo del titolo sottostante entro la maturità. Per semplicità, consideriamo il caso di una sola barriera. Riassumiamo qui di seguito le quattro principali opzioni call/put con barriera. Prenderemo $F(x) = (x - K)_+$ oppure $F(x) = (K - x)_+$ a seconda che ci interessi una call oppure una put.

Opzione up-and-in: fissata una barriera superiore $U > S_0$, il raggiungimento della barriera corrisponde all'attivazione:

$$h_{\text{up\&in}} = F(S_N) 1_{\{\text{esiste } n \leq N \text{ tale che } S_n \geq U\}};$$

opzione down-and-in: fissata una barriera inferiore $L < S_0$, il raggiungimento della barriera corrisponde all'attivazione:

$$h_{\text{down\&in}} = F(S_N) 1_{\{\text{esiste } n \leq N \text{ tale che } S_n \leq L\}};$$

opzione up-and-out: fissata una barriera superiore $U > S_0$, il non raggiungimento della barriera corrisponde all'attivazione:

$$h_{\text{up\&out}} = F(S_N) 1_{\{S_n < U \text{ per ogni } n \leq N\}};$$

opzione down-and-out: fissata una barriera inferiore $L < S_0$, il non raggiungimento della barriera corrisponde all'attivazione:

$$h_{\text{down\&out}} = F(S_N) 1_{\{S_n > L \text{ per ogni } n \leq N\}}.$$

Anche qui, il payoff dipende dall'evoluzione del prezzo su tutto l'intervallo temporale, una complicazione che non consente la determinazione di una formula esatta per il prezzo.

Esercizio 6.3. Usando il payoff e la rappresentazione in forma di aspettazione del prezzo, calcolare numericamente con il metodo Monte Carlo e studiare al variare della barriera il prezzo in 0 di una call barriera up-and-in e di una put barriera down-and-out.

Si considerino gli stessi parametri suggeriti nell'Esercizio 6.1.

6.4 Copertura dinamica di opzioni europee

In questo paragrafo vedremo come vengono usate nella pratica le formule del prezzo e della copertura, rimanendo ovviamente nel contesto CRR.

Supponiamo di lavorare con un'opzione call standard, di maturità T e prezzo di esercizio K . Fissiamo N monitoraggi, agli istanti $0 = t_0 < t_1 <$

$\dots < t_N = T$, e denotiamo al solito $(S_n)_{0 \leq n \leq N}$ e $(V_n)_{0 \leq n \leq N}$ rispettivamente il prezzo del sottostante ed il portafoglio replicante osservati a $t_1, \dots, t_N$. Seguiremo il punto di vista del venditore dell'opzione, che quindi, una volta avuto il premio iniziale, alle date fissate gestirà le sue operazioni di riaggiustamento delle quote del titolo rischioso e non, in modo da costruirsi il portafoglio replicante.

Ricordiamo che il valore del portafoglio replicante corrisponde al valore dell'opzione ed è dato da

$$V_n = c(n, S_n) = \Delta(n, S_{n-1}) S_n + \Delta^0(n, S_{n-1}) S_n^0 \quad n = 1, \dots, N,$$

dove Δ e Δ^0 indicano le funzioni che danno la strategia replicante: $\phi_n = \Delta(n, S_{n-1})$ e $\phi_n^0 = \Delta^0(n, S_{n-1})$. Riportiamo qui per comodità le formule (4.18) e (4.19) che danno Δ e Δ^0 :

$$\begin{aligned} \Delta(n, x) &= \frac{c(n, x(1+b)) - c(n, x(1+a))}{x(b-a)} \\ \Delta^0(n, x) &= (1+r)^{-n} \left(\frac{1+b}{b-a} c(n, x(1+a)) - \frac{1+a}{b-a} c(n, x(1+b)) \right). \end{aligned} \quad (6.8)$$

Passo 0 All'istante iniziale $t_0 = 0$, si riceve il compenso della vendita dell'opzione, quindi

$$V_0 = P(0, S_0).$$

Occorre però anche calcolare le quote di titolo rischioso e non rischioso da investire al passo successivo, date rispettivamente da

$$\phi_1 = \Delta(1, S_0) \quad \text{e} \quad \phi_1^0 = \Delta^0(1, S_0).$$

Queste quote saranno utilizzate nell'istante successivo di riaggiustamento, cioè al tempo t_1 .

Passo 1 All'istante t_1 , dapprima si osserva il prezzo del sottostante rischioso, e dunque si aggiorna il portafoglio usando le quote determinate al tempo 0:

$$V_1 = \phi_1 S_1 + \phi_1^0 (1+r).$$

A questo punto, si possono aggiornare le quote:

$$\phi_1 = \Delta(2, S_1) \quad \text{e} \quad \phi_1^0 = \Delta^0(2, S_1).$$

$\vdots$

Passo k All'istante t_k , si conoscono le quote ϕ_k e ϕ_k^0 da investire rispettivamente nel titolo rischioso e nel titolo non rischioso. Si osserva il prezzo del titolo rischioso, ed il valore aggiornato del portafoglio è

$$V_k = \phi_k S_k + \phi_k^0 (1+r)^k.$$

Le quote con cui lavorare al passo $k+1 \leq N$ sono aggiornate nel modo solito:

$$\phi_{k+1} = \Delta(k+1, S_k) \quad \text{e} \quad \phi_{k+1}^0 = \Delta^0(k+1, S_k).$$

$\vdots$

Passo N All'istante finale $t_N = T$, il portafoglio è quindi

$$V_N = \phi_N S_N + \phi_N^0 (1+r)^N$$

e dev'essere

$$V_N = \text{payoff} = (S_N - K)_+$$

Osserviamo che se l'opzione di interesse fosse più sofisticata di una call standard, ad esempio una call asiatica, non potremmo appellarci a formule esatte per il calcolo del prezzo e/o della delta, e quindi si procede con ulteriori approssimazioni numeriche (che daranno alla fine non esattamente il payoff ma una sua approssimazione).

Esercizio 6.4. *Implementare al calcolatore l'algoritmo della copertura dinamica appena descritto, nel caso di una call. Ovviamente, occorre simulare una possibile traiettoria del processo di prezzo e, per verificare che l'uguaglianza finale vale per davvero, si sconsiglia (eufemismo!!!) di simulare il prezzo sotto la misura di rischio neutro.*

Si considerino gli stessi parametri suggeriti nell'Esercizio 6.1 e le formule (6.8) della delta di una call. Per quanto riguarda i prezzi contenuti nelle formule (6.8), si usi dapprima la formula esatta come in (6.6); ripetere poi l'esercizio calcolando i prezzi necessari con Monte Carlo (ovviamente, in quest'ultimo caso, poiché i prezzi sono approssimati non si avrà una uguaglianza "pulita": ci si aspetta infatti che il valore finale del portafoglio approssimi il payoff).

6.5 Put americana

I metodi Monte Carlo possono essere usati anche nel contesto (prezzo, copertura etc.) di opzioni americane. Essi divengono però piuttosto sofisticati, e non vogliamo qui entrare nel merito (rimandiamo i lettori interessati al bel libro di Glasserman [4]). Ci limiteremo quindi a suggerire alcuni esercizi a carattere didattico, allo scopo di maneggiare prezzo e copertura di un'opzione put americana.

6.5.1 Prezzo

Abbiamo studiato la formula del prezzo $P_{\text{am}}(n, S_n)$ di una put americana, come risultato del principio di programmazione dinamica dato nella Proposizione 5.4.1. D'altra parte, dal Corollario 5.2.3, si ha anche

$$P_{\text{am}}(0, S_0) = \mathbb{E}^*((1+r)^{-\nu_0}(K - S_{\nu_0})_+), \quad (6.9)$$

dove

$$\nu_0 = \inf\{n \geq 0 : P_{\text{am}}(n, S_n) = (K - S_n)_+\} \quad (6.10)$$

è l'istante di esercizio.

Come primo esercizio, si chiede di verificare proprio la formula (6.9):

Esercizio 6.5. *Implementare il principio di programmazione dinamica come nella Proposizione 5.4.1. Simulare quindi molte volte l'istante di esercizio ν_0 sotto la misura di rischio neutro e, usando il metodo Monte Carlo, verificare la rappresentazione del prezzo iniziale $P_{\text{am}}(0, S_0)$ come nella formula (6.10).*

In altre parole, si chiede di generare ν_0 tante volte, ad esempio Q volte, e, detto ν_0^q il generico valore simulato per ν_0 , verificare che in effetti si ha:

$$P_{\text{am}}(0, S_0) \simeq \frac{1}{Q} \sum_{q=1}^Q (1+r)^{-\nu_0^q} (K - S_{\nu_0^q})_+.$$

Osserviamo che la media in (6.9) è valutata sotto $\mathbb{P}^*$, dunque si procederà a simulare tante volte il titolo rischioso sotto la misura di rischio neutro e, da (6.10), a determinare il primo istante in cui si osserva l'uguaglianza tra la funzione prezzo e la funzione payoff.

Sottolineiamo ancora una volta che, nell'Esercizio 6.5, ν_0 va simulato sotto la misura di rischio neutro, che in generale è diversa dalla misura di mercato. Se invece ν_0 è generato sotto la misura di mercato, è possibile, tramite il metodo Monte Carlo, studiarne alcune proprietà statistiche, quali la media, la varianza e la legge empirica:

$$\begin{aligned} \mathbb{E}(\nu_0) &\simeq \hat{\nu}_{0,Q} = \frac{1}{Q} \sum_{q=1}^Q \nu_0^q, \\ \text{Var}(\nu_0) &\simeq \frac{1}{Q-1} \sum_{q=1}^Q (\nu_0^q)^2 - \hat{\nu}_{0,Q}^2, \\ \mathbb{P}(\nu_0 = n) &\simeq \frac{1}{Q} \sum_{q=1}^Q 1_{\{\nu_0^q = n\}}, \quad n = 0, 1, \dots, N. \end{aligned}$$

Esercizio 6.6. *Calcolare numericamente, con il metodo Monte Carlo, la media e la varianza dell'istante di esercizio ν_0 al variare di $\gamma = \mathbb{P}(T_n = 1+a) \in (0, 1)$. Disegnare inoltre l'istogramma della legge empirica di ν_0 .*

6.5.2 Copertura dinamica

Vediamo ora come, nella pratica, vengono coperte le opzioni put americane. Per ogni n , $P_{\text{am}}(n, S_n)$ denota il valore della put americana all'istante n . Vogliamo quindi trovare le quote $(\phi_n)_n$ e $(\phi_n^0)_n$ rispettivamente di titolo rischioso e non rischioso, tali che:

- $(\phi_n)_n$ e $(\phi_n^0)_n$ siano predicibili;
- fino all'istante di esercizio, ossia per $n \leq \nu_0$, si abbia

$$P_{\text{am}}(n, S_n) = \phi_n S_n + \phi_n^0 S_n^0$$

Osserviamo che se sono soddisfatte entrambe le condizioni scritte sopra, all'istante di esercizio si ha

$$\phi_{\nu_0} S_{\nu_0} + \phi_{\nu_0}^0 S_{\nu_0}^0 = P_{\text{am}}(\nu_0, S_{\nu_0}) \equiv (K - S_{\nu_0})_+$$

e quindi il payoff è esattamente replicato nel momento in cui l'opzione è esercitata.

Inizialmente, si riceve il compenso $P_{\text{am}}(0, S_0)$, dove ovviamente S_0 denota il prezzo iniziale del sottostante. Se $P_{\text{am}}(0, S_0) = (K - S_0)_+$, allora non ci sono quote da preparare: l'opzione si conclude immediatamente. Se invece $P_{\text{am}}(0, S_0) > (K - S_0)_+$ allora $\nu_0 \geq 1$ e si preparano le quote da investire all'istante 1. Procediamo come abbiamo fatto nel caso europeo: all'istante 1 dev'essere $\phi_1 S_1 + \phi_1^0 S_1^0 = P_{\text{am}}(1, S_1)$ e poiché si può avere solo $S_1 = S_0(1+a)$ oppure $S_1 = S_0(1+b)$, dev'essere

$$\begin{aligned} \phi_1 S_0(1+a) + \phi_1^0 S_1^0 &= P_{\text{am}}(1, S_0(1+a)) \\ \phi_1 S_0(1+b) + \phi_1^0 S_1^0 &= P_{\text{am}}(1, S_0(1+b)). \end{aligned}$$

Sottraendo, si ottiene $\phi_1 = \Delta_{\text{am}}(1, S_0)$, dove

$$\Delta_{\text{am}}(1, x) = \frac{P_{\text{am}}(1, x(1+b)) - P_{\text{am}}(1, x(1+a))}{x(b-a)}. \quad (6.11)$$

Sostituendo, si ottiene facilmente $\phi_1^0 = \Delta_{\text{am}}^0(1, S_0)$, dove

$$\Delta_{\text{am}}^0(1, x) = (1+r)^{-1} \left(\frac{1+b}{b-a} P_{\text{am}}(1, x(1+a)) - \frac{1+a}{b-a} P_{\text{am}}(1, x(1+b)) \right). \quad (6.12)$$

Ora, all'istante 1 si osserva il sottostante e si ha evidentemente $P_{\text{am}}(1, S_1) = \phi_1 S_1 + \phi_1^0 S_1^0$. Se $P_{\text{am}}(1, S_1) = (K - S_1)_+$, allora non ci sono quote da preparare: l'opzione si esercita. Se invece $P_{\text{am}}(1, S_1) > (K - S_1)_+$ allora $\nu_0 \geq 2$ e si preparano le quote da investire all'istante 2 con un ragionamento analogo a quello visto sopra. Volendo generalizzare le formule (6.11) e (6.12), si ottiene che per ogni $n \leq \nu_0$, il portafoglio replicante è dato da

$$\phi_n = \Delta_{\text{am}}(n, S_{n-1}) \quad \text{e} \quad \phi_n^0 = \Delta_{\text{am}}^0(n, S_{n-1}),$$

dove

$$\Delta_{\text{am}}(n, x) = \frac{P_{\text{am}}(n, x(1+b)) - P_{\text{am}}(n, x(1+a))}{x(b-a)} \quad (6.13)$$

e

$$\Delta_{\text{am}}^0(n, x) = (1+r)^{-n} \left(\frac{1+b}{b-a} P_{\text{am}}(n, x(1+a)) - \frac{1+a}{b-a} P_{\text{am}}(n, x(1+b)) \right). \quad (6.14)$$

Sulla base di quanto appena visto, possiamo riassumere l'algoritmo di riaggiustamento delle quote come segue.

Passo 0 All'istante iniziale 0, si riceve il compenso della vendita dell'opzione, quindi

$$V_0 = P_{\text{am}}(0, S_0).$$

Ora, occorre controllare se c'è esercizio.

- Se $V_0 = (K - S_0)_+$, l'opzione si conclude, e si esce dal programma con $\nu_0 = 0$.
- Se $V_0 > (K - S_0)_+$, si calcolano le quote di titolo rischioso e non rischioso da investire al passo successivo, date rispettivamente da

$$\phi_1 = \Delta_{\text{am}}(1, S_0) \quad \text{e} \quad \phi_1^0 = \Delta_{\text{am}}^0(1, S_0),$$

e si passa all'istante successivo.

Passo 1 All'istante 1, dapprima si osserva il prezzo del sottostante rischioso, e dunque si aggiorna il portafoglio usando le quote determinate al tempo 0:

$$V_1 = \phi_1 S_1 + \phi_1^0 (1+r).$$

A questo punto, si controlla se l'istante 1 è di esercizio o no.

- Se $V_1 = (K - S_1)_+$, l'opzione si conclude, e si esce dal programma con $\nu_0 = 1$.
- Se $V_1 > (K - S_1)_+$, si calcolano le quote di titolo rischioso e non rischioso da investire al passo successivo, date rispettivamente da

$$\phi_2 = \Delta_{\text{am}}(2, S_1) \quad \text{e} \quad \phi_2^0 = \Delta_{\text{am}}^0(2, S_1),$$

e si passa all'istante successivo.

⋮

Passo k Se si arriva all'istante k , vuol dire che non c'è stato esercizio e si conoscono le quote ϕ_k e ϕ_k^0 da investire rispettivamente nel titolo rischioso e nel titolo non rischioso. Quindi, si osserva il prezzo del titolo rischioso, ed il valore aggiornato del portafoglio è

$$V_k = \phi_k S_k + \phi_k^0 (1+r)^k.$$

Ancora una volta si controlla se c'è esercizio.

- Se $V_k = (K - S_k)_+$, l'opzione si conclude, e si esce dal programma con $\nu_0 = k$.
- Se $V_k > (K - S_k)_+$, si calcolano le quote di titolo rischioso e non rischioso da investire al passo successivo, date rispettivamente da

$$\phi_{k+1} = \Delta_{\text{am}}(k+1, S_k) \quad \text{e} \quad \phi_{k+1}^0 = \Delta_{\text{am}}^0(k+1, S_k),$$

e si passa all'istante successivo.

⋮

Passo N Se si arriva all'istante N , vuol dire che non c'è stato esercizio, il che ovviamente avviene senz'altro all'istante N (ricordiamo infatti che $V_N = P_{\text{am}}(N, S_N) = (K - S_N)_+$). Dunque, si osserva il sottostante S_N e, poiché si conoscono le quote ϕ_N e ϕ_N^0 da investire rispettivamente nel titolo rischioso e nel titolo non rischioso, il valore aggiornato del portafoglio è

$$V_N = \phi_N S_N + \phi_N^0 (1+r)^N.$$

In uscita dalla procedura sopra descritta, si ha in particolare l'istante di esercizio ν_0 ed il relativo valore del portafoglio V_{ν_0} . Ovviamente, dev'essere

$$V_{\nu_0} = (K - S_{\nu_0})_+.$$

Esercizio 6.7. Implementare al calcolatore l'algoritmo della copertura dinamica appena descritto, simulando una possibile traiettoria del processo di prezzo (sotto una misura diversa da quella neutrale al rischio!).

Si considerino gli stessi parametri suggeriti nell'Esercizio 6.1 e le formule (6.13) e (6.14) della copertura della put americana, a partire dalla formula all'indietro del prezzo come nella Proposizione 2.1.11.

Bibliografia

- [1] P. Baldi: *Equazioni differenziali stocastiche e applicazioni*. Pitagora Editrice, 2000.
- [2] F. Black, M. Scholes: The pricing of options and corporate liabilities. *Journal of Political Economy*, 81, 637-654, 1973.
- [3] P. Billingsley: *Probability and measure*. Wiley, 1995.
- [4] P. Glasserman: *Monte Carlo Methods in Financial Engineering*. Springer, 2003.
- [5] J.C. Hull: *Options, futures and other derivatives. 5th edition*. Prentice Hall, 2002.
- [6] D. Lamberton, B. Lapeyre: *Introduction to stochastic calculus in finance*. Chapman and Hall, 1996.
- [7] M. Musiela, M. Rutkowski: *Martingale methods in financial modelling*. Springer, 1998.
- [8] S.M. Ross: *An elementary introduction to Mathematical Finance. Options and other topics. Second edition*. Cambridge University Press, 2003.
- [9] D.M. Salopek: *American put options*. Addison Wesley Longman Inc, 1997.
- [10] D. Williams: *Probability with martingales*. Cambridge University Press, 1991.