

### 1. Economia del benessere e scienza delle finanze

#### 1.1. L'oggetto della scienza delle finanze

La scienza delle finanze ha come oggetto il finanziamento delle attività dello Stato nell'economia di mercato.

In un'economia di mercato la produzione di beni è effettuata da produttori, motivati dalla ricerca del profitto, che, con l'uso di fattori produttivi, di loro proprietà o procurati sul mercato, producono e offrono beni e servizi, a prezzi che dovrebbero riflettere i costi di produzione. Al mercato partecipano individui, consumatori e produttori, che prendono libere decisioni di acquistare o non acquistare i beni offerti, confrontando i prezzi di offerta con i propri bisogni e i mezzi a disposizione. I prezzi che emergono dalle attività di scambio sono il veicolo dell'informazione per produttori e consumatori del costo e della scarsità relativa dei beni e costituiscono la guida principale delle loro decisioni economiche. Questo modo di organizzare la produzione e lo scambio dei beni, di cui proprietà privata dei mezzi di produzione, libera iniziativa e concorrenza tra più produttori rappresentano i connotati fondamentali, si è rivelato, storicamente, più efficiente di altri modi, ad esempio di quelli sperimentati dalle economie pianificate dal centro.

Anche per il funzionamento dell'economia di mercato sono però necessari regole e vincoli. Un maestro del pensiero liberale come Luigi Einaudi (1874-1961), uno dei maggiori studiosi di scienza delle finanze della tradizione italiana, presidente della Repubblica dal 1948 al 1955, nelle splendide *Lezioni di politica sociale* del 1949, dopo avere descritto il ruolo e le funzioni positive del mercato prendendo lo spunto dalla vivida descrizione della fiera di un borgo di campagna, sente il bisogno di avvertire il lettore che:

tutti coloro i quali vanno alla fiera, sanno che questa non potrebbe avere luogo se, oltre ai banchi dei venditori i quali vantano a gran voce la bontà della loro merce, ed oltre la folla dei compratori che ammira la bella voce, ma prima vuole prendere in mano le scarpe per vedere se sono di cuoio o di cartone, non ci fosse qualcos'altro: il cappello a due punte della coppia di carabinieri che si vede passare sulla piazza,

la divisa della guardia municipale che fa tacere due che si sono presi a male parole, il palazzo del municipio, col segretario ed il sindaco, la pretura e la conciliatura, il notaio che redige i contratti, l'avvocato a cui si ricorre quando si crede di essere a torto imbrogliati in un contratto, il parroco, il quale ricorda i doveri del buon cristiano, doveri che non bisogna dimenticare nemmeno alla fiera. E ci sono le strade, le une dure e le altre fangose che conducono dai casolari della campagna al centro, ci sono le scuole dove i ragazzi vanno a studiare. E tante altre cose ci sono, che se non ci fossero, anche quella fiera non si potrebbe tenere o sarebbe tutta diversa da quel che effettivamente è<sup>1</sup>.

Il mercato non può esistere in assenza di altre istituzioni, senza le quali non potrebbe esprimere il meglio della sua funzione. Tra queste istituzioni, come mostrano bene gli esempi di Einaudi, un ruolo particolarmente importante, ma non unico, ha lo Stato, che manifesta la sua presenza in molteplici campi (sicurezza, giustizia, amministrazione, infrastrutture, scuole) e a molti livelli (nazionale e locale).

La presenza dello Stato è molto più pervasiva di quanto non si pensi. Lo studente che segue un corso universitario provi ad immaginare quante volte, nell'arco di una giornata, si imbatte nello Stato. Si sveglia, accende la luce: è fornita da un'impresa pubblica o regolata dallo Stato. Idem per l'acqua della doccia e il gas acceso per il primo caffè. Va all'università percorrendo strade la cui costruzione e manutenzione sono opera di un ente pubblico. Se qualcuno ha il vizio di fumarsi una sigaretta prima della lezione, sappia che oltre il 75% del suo prezzo è costituito da un'imposta indiretta messa dallo Stato. Idem per la benzina. La lezione è quasi sempre svolta in un edificio che è di proprietà di un ente pubblico, da cui dipende anche il docente che tiene la lezione.

Proviamo ora ad immaginare una normale esistenza. Nasciamo, molto spesso, in un ospedale pubblico o convenzionato con lo Stato. Passiamo la nostra infanzia e giovinezza in scuole, quasi sempre pubbliche. Entriamo nel mercato del lavoro: una parte non trascurabile di noi lavora alle dipendenze di un ente pubblico. Il salario o lo stipendio di tutti è in gran parte costituito da prelievi obbligatori, i contributi sociali, destinati a tutelarci da rischi sociali. Non c'è bene che acquistiamo che non sia gravato da una qualche imposta. Ci sposiamo, abbiamo figli, e spesso, soprattutto se non siamo molto ricchi, riceviamo aiuti economici dallo Stato o sgravi fiscali. Se ci ammaliamo siamo quasi sempre curati in ospedali pubblici. Vecchi, smettiamo di lavorare e riceviamo una pensione, pagata da un ente pubblico, la cui misura è stata definita da regole fissate da leggi dello Stato. Moriamo e siamo sepolti in un cimitero costruito dal Comune.

Studiare il ruolo dello Stato significa domandarsi:

Quali sono le ragioni della presenza pubblica in una economia moderna di mercato: perché c'è lo Stato? Perché ha queste dimensioni? Perché si preoccupa di certe cose e non di altre?

<sup>1</sup> L. Einaudi, *Lezioni di politica sociale*, con nota introduttiva di F. Caffè, Torino, Einaudi, 1964, p. 41.

Come è organizzato lo Stato, come si comporta, quali sono gli strumenti della sua azione? Viviamo in una democrazia rappresentativa in cui le decisioni politiche, fra cui anche quelle che riguardano le spese e le entrate pubbliche, devono seguire particolari procedure. L'attività amministrativa è organizzata nel territorio nazionale secondo modelli che devono essere spiegati e valutati.

Quali sono gli effetti economici delle azioni dello Stato: che accade se viene modificata, ad esempio, l'aliquota di un'imposta? Quali sono le conseguenze della decisione di aumentare o di diminuire le pensioni?

Le domande poste non sono oggetto dell'attenzione solo degli economisti. Anche altre discipline affrontano questi argomenti: il diritto, la sociologia, la scienza politica, la storia, la filosofia morale e molto spesso in questi capitoli avremo occasione di fare cenni al contributo che queste discipline possono fornire. Sono molti infatti i legami con la filosofia morale, che ci aiuta a capire quali regole sono alla base dei comportamenti etici individuali e sociali, con la scienza politica, che studia i processi e le istituzioni con cui gli uomini prendono decisioni collettive, con la sociologia, che affronta il comportamento degli uomini in quanto membri di una società, e questi legami saranno particolarmente evidenti in questo capitolo introduttivo in cui ci porremo le questioni più di fondo, sulla natura dello Stato e sul suo ruolo nella vita sociale. Ma anche nel capitolo secondo, quando sentiremo il bisogno di calare questi problemi nel contesto della società in cui viviamo oggi, faremo cenni alla storia, almeno quella più recente, del ruolo dello Stato nell'economia italiana. E anche quando concentreremo l'attenzione sui contenuti più propri di questa disciplina, soffermandoci a considerare l'attività dello Stato nelle sue funzioni di regolatore dell'attività economica e di utilizzatore di risorse finanziarie – spese ed entrate pubbliche – avremo occasione di utilizzare o misurarci con i risultati degli studi di diritto pubblico, amministrativo e, soprattutto, tributario. Il bisogno di uscire dai confini rigorosi, ma spesso un po' asfittici dell'economia sarà particolarmente forte quando affronteremo lo studio delle spese di protezione sociale (il welfare state), in cui sono particolarmente presenti aspetti di equità e motivazioni non strettamente economiche, che possono trovare risposta solo con l'aiuto della filosofia morale e della sociologia.

In misura preponderante il ruolo dello Stato si traduce in attività che hanno una rilevanza finanziaria. Lo studio delle ragioni e degli effetti del prelievo e della spesa pubblica è quindi l'oggetto principale della disciplina. Questo tipo di attività ha un riferimento formale ed emblematico nel Bilancio dello Stato: un documento poco attraente, ma molto importante, se Schumpeter ha potuto dire che esso è lo specchio di tutti i problemi politici e sociali di una nazione. L'attività dello Stato non si manifesta però solo attraverso l'impiego di strumenti finanziari. Non meno importante è l'attività con cui il potere pubblico definisce le regole del gioco che tutti i membri di una società devono rispettare nei loro comportamenti economici. Alcune attività, per essere esercitate da privati, richiedono autorizzazioni; la produzione deve essere effettuata rispettando certi requisiti che riflettono una volontà pubblica, e così via. Questo tipo di interventi, che possiamo in via generale indicare come «regolamentazioni», assume oggi un'importanza crescente ed



è alla base delle discussioni che si svolgono sul modo migliore di organizzare i grandi servizi sociali o le politiche di privatizzazione.

Affrontando lo studio della scienza delle finanze è naturale cercare di ordinare e classificare le funzioni che lo Stato deve assolvere attraverso lo strumento del bilancio e della regolamentazione. Quasi tutti i manuali fanno riferimento alla tripartizione proposta da R. Musgrave, uno dei più importanti studiosi di finanza pubblica del XX secolo, che, nell'opera *The Theory of Public Finance* del 1959<sup>2</sup>, propose di articolare l'attività finanziaria dello Stato e quindi la struttura del bilancio, che ne rappresenta la descrizione *ex ante* o *ex post*, in tre grandi Dipartimenti (la parola usata da Musgrave è *Branches*) che hanno come oggetto tre fondamentali funzioni: allocazione, redistribuzione e stabilizzazione.

L'*allocazione* cerca di capire in che modo lo Stato influenza l'efficienza economica, che sembra essere uno dei principali obiettivi che l'economia di mercato persegue nelle società capitalistiche contemporanee. Questa funzione è senza dubbio quella che riceve maggiore attenzione da parte degli studiosi di economia pubblica ed implicitamente ad essa si è fatto riferimento più sopra quando abbiamo fornito una prima definizione dell'oggetto dell'economia pubblica. Il paradigma dominante della teoria economica, il modello dell'equilibrio economico generale, costruito alla fine del secolo XIX da Walras e Pareto e affinato soprattutto negli anni '50 del secolo scorso da studiosi come Arrow e Debreu, tenta di fornire una spiegazione di come si formano i prezzi e le quantità prodotte in un'economia di mercato, sulla base del presupposto che gli attori economici agiscano secondo la razionalità economica: massimizzando i profitti, se produttori, e l'utilità, se consumatori. I risultati di questi comportamenti razionali godono di proprietà di ottimalità, cioè sono efficienti, in quanto esito di comportamenti razionali. Il contributo che l'economia pubblica ha dato a questo impianto teorico è iniziato con l'osservazione che esiste un particolare tipo di beni – i *beni pubblici* – che, come illustreremo più ampiamente in seguito, possiede caratteristiche intrinseche tali da mettere in crisi alcune proprietà dell'economia di mercato. Vi sono tipi di servizi (per esempio, la difesa, l'amministrazione della giustizia, la sicurezza pubblica, ecc.), che un'impresa privata non ha incentivo a produrre perché, se lo facesse, non potrebbe costringere gli acquirenti a pagare il prezzo del servizio da essi richiesto. I potenziali acquirenti, a loro volta, non hanno incentivo a rivelare quanto sarebbero disposti a pagare per tali beni. Ciò produce un fallimento del mercato e per questo tali servizi vengono solitamente forniti dallo Stato. La loro natura può darci una spiegazione dell'intervento pubblico. D'altro canto l'intervento dello

<sup>2</sup> L'impostazione di questo Corso risente in modo particolare dall'opera di R. Musgrave. Chi volesse approfondire il pensiero di questo studioso potrebbe utilmente consultare – oltre all'opera citata nel testo, che è uno dei trattati che hanno esercitato maggiore influenza sulla disciplina – una raccolta dei suoi scritti pubblicata dal Mulino, *Finanza pubblica, equità, democrazia* (1995). In particolare il saggio *Una breve storia della teoria fiscale* è molto formativo e la sua lettura è fortemente raccomandata dopo avere appreso gli elementi di base della scienza delle finanze. Il volume citato contiene anche un'ampia introduzione che inquadra il pensiero di questo autore nel contesto storico.

Stato, nella sua attività di regolazione delle attività private o di produzione dei beni e servizi che il mercato non ha interesse a produrre o non produce nella dimensione adeguata, finisce fatalmente per avere degli effetti allocativi. Imponendo vincoli, introducendo imposte per finanziare la propria attività, altera i prezzi, che come si è detto, sono i segnali che guidano gli operatori economici nella ricerca di soluzioni efficienti. Da ciò l'importanza di studiare quali siano le modalità di intervento che minimizzano le distorsioni prodotte dall'intervento pubblico.

La *redistribuzione* è la seconda funzione svolta dallo Stato e realizzata con il bilancio pubblico. L'introduzione di correttivi per migliorare le funzioni allocative fondamentalmente svolte dal mercato è infatti solo uno dei ruoli dello Stato. Spesso i suoi interventi riflettono una deliberata volontà di intervenire sulla distribuzione dei redditi e dei patrimoni che nella società si realizza attraverso l'operare del libero mercato. In assenza di tali politiche la distribuzione del reddito sarebbe determinata dalla distribuzione delle dotazioni iniziali dei soggetti che compongono una società, tutelata dal diritto di proprietà, e dai prezzi dei fattori (salari, interessi, rendite) e dei beni che si realizzano nel mercato. Anche se il mercato operasse nel modo più efficiente, la distribuzione del reddito e della ricchezza prodotta dal mercato non necessariamente corrisponderebbe a criteri di equità. La presenza di istituti quali ad esempio l'eredità, che costituiscono un cardine del sistema capitalistico, ha l'effetto di rendere estremamente vischiosa la modificazione della distribuzione delle risorse. Il mercato, per citare ancora le *Lezioni* di Einaudi, è adatto per

produrre beni e servizi, precisamente nella quantità e della qualità corrispondenti alla *domanda* degli uomini; non si afferma che il mercato indirizzi *altresì* la produzione a produrre beni e servizi nella quantità e nella qualità che sarebbe desiderata dagli stessi uomini. Sul mercato si soddisfano domande, non bisogni.

Nelle società moderne lo Stato ha il potere di modificare questa distribuzione e lo esercita in misura molto intensa. Le forme delle sue azioni redistributive sono molteplici. La più comune è rappresentata da trasferimenti monetari a favore di particolari gruppi di cittadini: gran parte delle spese di protezione sociale (si pensi ai sussidi di disoccupazione, alle pensioni sociali, ai trasferimenti a favore delle famiglie povere, agli aiuti economici alle famiglie con figli a carico) ha questa funzione. Della stessa natura sono le redistribuzioni che derivano dal modo in cui è costruito o modificato un sistema tributario. Se le imposte sul reddito sono fatte pagare in misura fissa per ogni cittadino o in misura proporzionale al reddito o in modo progressivo, si generano effetti distributivi molto diversi. La rilevanza di queste redistribuzioni è ben chiara nei momenti in cui avvengono riforme dei sistemi tributari, in cui si assiste ad accese discussioni e valutazioni su chi guadagna e chi perde dall'introduzione di una data riforma.

La redistribuzione può però essere attuata anche intervenendo sui prezzi dei beni. Si pensi ad esempio alle agevolazioni che sono spesso concesse ai soggetti economicamente più deboli quando si tratta di acquistare taluni servizi, come ad esempio i servizi di un asilo nido e un posto in una casa di riposo.

Infine le redistribuzioni possono essere realizzate attraverso la fornitura diretta di servizi ai cittadini, in funzione dei loro bisogni e non della capacità che essi hanno di pagare i servizi offerti. Un esempio molto importante di questo tipo è costituito da un servizio sanitario nazionale pubblico e universale, in cui le cure sono prestate gratuitamente in funzione del bisogno.

È in ogni caso molto difficile immaginare un intervento pubblico che non abbia effetti redistributivi, che non modifichi la distribuzione delle risorse, e quindi del benessere, da un soggetto ad un altro della società. Molte volte questi effetti non sono intenzionali, ma solo il risultato di azioni intraprese con altre finalità (allocative, ad esempio).

La *stabilizzazione* è la terza funzione dello Stato proposta da Musgrave, a cui è affidato il compito di garantire un livello di produzione il più vicino possibile a quello di pieno impiego. Il modello dell'equilibrio economico generale, nell'ambito delle astratte ipotesi che lo definiscono, ha come esito il pieno impiego dei fattori produttivi, fra cui riveste particolare interesse il lavoro. Ciò però non corrisponde a quanto osserviamo nella realtà storica. La disoccupazione esiste e in misura molto superiore a quella indotta dai normali flussi in entrata e in uscita dal mercato del lavoro. La *Teoria generale* di J.M. Keynes, uno dei libri di economia più importanti del secolo scorso, è stata elaborata nel 1936 proprio per cercare di spiegare questo fenomeno e da essa è nata una nuova disciplina economica, la macroeconomia. Le risposte fornite differiscono da scuola a scuola: rigidità dei prezzi, aspettative incoerenti, asimmetrie informative, ecc. In ogni caso l'esito è una domanda aggregata inferiore all'offerta potenziale. L'economia keynesiana ci insegna che la spesa pubblica e le imposte possono modificare il livello della domanda, operando attraverso canali differenziati. A loro volta spese e imposte sono influenzate dall'andamento del ciclo, e ciò complica il lavoro di chi deve tenere sotto controllo l'andamento dei saldi finanziari. L'efficacia delle politiche fiscali nazionali di stabilizzazione risulta poi indebolita dal grado sempre più elevato di integrazione reale. La prevedibilità degli effetti della politica fiscale è sempre più difficile in un mondo in cui i cambi sono flessibili ed è accentuata la libertà dei movimenti dei capitali. Per un paese europeo, in questa fase storica, l'analisi degli effetti macroeconomici dell'attività pubblica è fortemente condizionata dalla partecipazione all'Unione monetaria. Per il nostro paese, in particolare, il raggiungimento di questa meta ha imposto e pone una grande attenzione ai problemi derivanti dalla presenza di un debito pubblico elevato. La funzione di stabilizzazione – l'uso sistematico della spesa e della tassazione per controllare il ciclo – nel senso in cui l'intendeva Musgrave, in un periodo di grande entusiasmo per l'economia keynesiana, è stata posta in secondo piano negli ultimi decenni. Tuttavia la grande crisi che ha raggiunto il suo apice alla fine del 2008, innescata dagli squilibri del sistema finanziario, a partire dai fallimenti nel mercato dei mutui ipotecari negli USA, e che si è con rapidità straordinaria trasmessa a tutto il globo, ha riproposto con forza l'importanza delle politiche di sostegno della domanda attraverso il bilancio pubblico. Le politiche keynesiane stanno attraversando un momento di imprevista popolarità e sostegno anche da parte dei liberisti, che, fino a pochi anni fa, le consideravano sorpassate e irrilevanti. Nei paesi dell'Europa

l'intensità della crisi ha riproposto all'attenzione dei governi l'attualità delle regole che sono alla base del Patto di stabilità e crescita e ci si interroga su come principi di corretta gestione della finanza pubblica possano essere compatibili con interventi volti a promuovere la crescita delle economie.

Prima di abbandonare questo paragrafo introduttivo dedicato alla descrizione dei contenuti tradizionali della scienza delle finanze e quindi dei temi che saranno affrontati in questi capitoli, può essere utile richiamare un aspetto che sta a monte dello schema espositivo qui proposto, che ci consente di fare almeno un cenno ad un filone di studi che negli ultimi anni ha ricevuto molta attenzione e a cui non sarà qui possibile dedicare molto spazio. Lo studio delle funzioni dello Stato secondo la tripartizione musgraviana è tipico di una visione del funzionamento di una società in cui lo Stato è visto come uno, forse il più importante, degli attori della società, accanto agli attori economici, produttori e consumatori: un attore che ha proprie finalità e obiettivi. Lo Stato è visto come un benevolente *pater familias*, che ha come finalità il perseguimento del bene comune, in nome e per conto del popolo. In realtà l'economia pubblica, quando ha messo a fuoco il concetto di bene pubblico, di un bene, come si è detto, che il mercato non ha incentivo a produrre e del quale i cittadini non hanno incentivo a rivelare la loro disponibilità al pagamento, è stata stimolata ad interessarsi dei meccanismi di formazione della volontà politica (modelli di governo, sistemi elettorali, ecc.) attraverso i quali è possibile comprendere modi alternativi al mercato per spiegare la produzione, l'offerta e il finanziamento dei beni pubblici. Il campo della disciplina si è quindi allargato. Lo Stato non è più assunto come un dato della società, ma si è cercato di entrare dentro la sua «scatola nera», di spiegarne la genesi, di studiare le relazioni tra le preferenze dei cittadini e quelle non necessariamente coincidenti delle autorità politiche. Ci si è domandati se il politico non abbia preferenze e obiettivi distinti da quelli dei cittadini. Si è affrontato lo studio del ruolo della burocrazia e delle sue relazioni con il potere politico. Naturalmente questi problemi – la natura dello Stato, lo studio dei meccanismi politici – non sono stati creati *ex novo*. Esiste una tradizione di pensiero di scienza della politica molto rigogliosa, a cui gli economisti hanno attinto. L'aspetto che caratterizza questo campo di studi, indicato come *Public Choice* – studio delle scelte collettive – e sfociato nell'ultimo decennio in un campo ancora più ampio, noto come *Political Economy* – l'analisi economica dei comportamenti politici – è l'uso della teoria economica, fondata sull'ipotesi di comportamenti ottimizzanti di soggetti razionali, di norma motivati dall'interesse individuale, e dei suoi strumenti, per la costruzione di modelli formalizzati di analisi di comportamenti strategici<sup>3</sup>.

<sup>3</sup> Nella storia del pensiero economico e politico sono moltissimi coloro che hanno contribuito a questo filone di studi. È però largamente condiviso che J. Buchanan, americano, premio Nobel per l'economia, abbia avuto un ruolo di primo piano nello sviluppo della *Public Choice*. L'applicazione dell'analisi economica a campi diversi dalle azioni economiche in senso stretto è però anche legata all'opera di un altro americano, anch'esso premio Nobel per l'economia, Gary Becker, famoso per le sue provocatorie analisi economiche della famiglia.



Un esempio può essere utile per dare l'idea del modo di procedere di questo tipo di studi. Si cerchi di rispondere a questa domanda. Perché esistono strumenti di sostegno economico dei cittadini poveri, che non sono in grado di procurarsi mezzi di sostentamento? Secondo l'*impostazione tradizionale*, di cui la classificazione musgraviana è una derivazione, la ragione è semplicemente da ricercarsi nella volontà dello Stato di perseguire finalità di equità. Si presuppone che il governo preferisca una distribuzione del reddito non troppo squilibrata e quindi persegua tale obiettivo con lo strumento di un trasferimento monetario, prelevato a carico di tutta la collettività. Un cultore della *Public Choice*, ad esempio, uno ideologicamente orientato a favore della *New Right*, non si accontenterebbe di questa risposta e potrebbe obiettare che il governo adotta quelle politiche non perché ispirato da principi di uguaglianza. Di norma la distribuzione dei redditi non è simmetrica, ma presenta un *bias* verso sinistra: il numero dei soggetti che hanno un reddito inferiore alla media è maggiore di quello che ha un reddito superiore, ovvero il reddito medio è di solito superiore a quello mediano. In un sistema democratico in cui ogni testa ha un voto, è allora più facile che si formino maggioranze elettorali e quindi governi in cui sono rappresentati gli interessi dei più poveri. La spiegazione è quindi individuata nel meccanismo politico, all'interno del quale agirebbero le stesse motivazioni di *self-interest* proposte nelle analisi dei problemi economici. Potrebbe anche accadere che un altro cultore, questa volta con propensioni a favore di *impostazioni marxiane*, pure insoddisfatto della spiegazione tradizionale, argomenti che nella società il potere politico è sempre espressione della classe dominante, nel sistema capitalistico rappresentata da chi detiene la proprietà dei mezzi di produzione, che esercita, pur nel contesto di uno Stato legale, lo sfruttamento delle classi che da quella proprietà sono escluse. Le politiche di contrasto della povertà in questa prospettiva si spiegherebbero con l'esigenza delle classi dominanti, in una società formalmente democratica, di evitare tensioni sociali troppo forti. Si tratterebbe quindi di una «carità pelosa» volta al contenimento di potenziali tensioni sociali eversive dell'ordine costituito. Le due spiegazioni alternative, di proposito esposte in modo molto rozzo, sono forse sufficienti per dare l'idea delle interessanti direzioni, a cui, come si è visto, possono corrispondere propensioni ideologiche anche molto lontane, verso cui si può spingere l'analisi economica del comportamento politico.

## 1.2. Teoria positiva, teoria normativa ed Economia del benessere

La descrizione fatta del contenuto della scienza delle finanze ha implicitamente utilizzato proposizioni che possono essere ricondotte a categorie più generali delle scienze sociali note come *teoria positiva* e *teoria normativa*: due nozioni che è bene chiarire con cura prima di procedere nello studio dell'economia pubblica. Nello studio e nei dibattiti di economia i problemi a cui si cerca di dare risposta possono infatti essere posti in modi diversi.

Ci si può ad esempio domandare: perché un chilo di mele costa 3 euro mentre un'automobile di media cilindrata ne costa 20.000? Perché la produzione di gelati è di solito svolta da imprese private, mentre la difesa è

sempre compito dello Stato e i servizi sanitari sono spesso offerti da strutture pubbliche? Perché si paga un pedaggio nelle autostrade e non lo si paga nelle strade normali? Perché nei sistemi economici in cui viviamo esiste, in misura più o meno intensa, la disoccupazione?

Oppure ci si può chiedere: è bene o no introdurre i ticket sanitari? È meglio il sistema pensionistico pubblico a ripartizione o quello privato a capitalizzazione? Come fare per migliorare il servizio postale? Come si dovrebbe riformare l'imposta sul reddito? Cosa fare per ridurre la disoccupazione?

La differenza nel modo in cui sono formulate le domande dei due gruppi di esempi è molto importante e si deve essere ben sicuri di averla compresa. Nel primo gruppo il punto di interesse è capire le ragioni di un fenomeno sociale, individuare le regole che spiegano certi fenomeni e comportamenti sociali. Esattamente come fa uno studioso di fisica quando si domanda che cosa è la luce o perché si scivola sul ghiaccio. La finalità è puramente conoscitiva. Le risposte a queste domande costituiscono esempi di *teoria positiva*. Per poter formulare proposizioni di teoria positiva è necessario costruire una teoria. Il metodo con cui lo studio della realtà è affrontato consiste nella costruzione di modelli in cui la realtà viene semplificata, utilizzando un processo di astrazione che mette in rilievo solo gli aspetti di carattere economico ritenuti rilevanti. In tale processo è necessario mettere a fuoco le variabili oggetto dello studio, distinguendo le grandezze endogene, ovvero le incognite del problema (nel nostro esempio i prezzi e le quantità prodotte e scambiate dei beni), da quelle esogene o date (la dotazione iniziale dei fattori produttivi); si formulano ipotesi sui comportamenti dei soggetti interessati e sull'ambiente circostante (le caratteristiche delle funzioni di preferenza degli individui e delle funzioni di produzione, le ipotesi di comportamento ottimizzante); si impostano le relazioni che esistono fra tali variabili in sistemi di equazioni che rappresentano le regole del gioco (condizioni di equilibrio tra domande e offerte, vincoli di bilancio che devono essere rispettati) e si derivano infine i risultati, vale a dire i valori di equilibrio delle variabili endogene.

Nelle domande del secondo gruppo invece l'interesse è la ricerca degli strumenti che consentano di raggiungere una situazione desiderabile (un buon servizio sanitario o postale, una buona riforma fiscale, il pieno impiego). Poiché tali situazioni interessano tutti gli individui di cui è composta una collettività, si tratta di ricercare una situazione di ottimo sociale. Le risposte a queste domande costituiscono esempi di *teoria normativa*. Le risposte a queste domande costituiscono esempi di *teoria normativa*.

Per poter avanzare una proposizione normativa sono necessari due presupposti fondamentali.

In primo luogo è necessario possedere una corretta teoria positiva del fenomeno che si studia. Non posso sperare di trovare gli strumenti per eliminare la disoccupazione se, prima, non mi sono chiarito le idee sulle sue cause. In altre parole *la teoria normativa presuppone la teoria positiva*.

In secondo luogo è necessario avere un'idea di che cosa sia il «bene della collettività», sapere che cosa sia ciò che gli economisti chiamano un «ottimo sociale». La risposta sarebbe possibile e facile se tutti avessimo le stesse idee: in tal caso sarebbe giusto e bene ciò che tutti pensano sia tale. Una società in cui vi fosse sempre unanimità di consensi sarebbe forse un po' noiosa, ma sicuramente in grado di prendere con facilità decisioni sociali. Nella realtà

l'unanimità si manifesta solo raramente: gli uomini hanno idee diverse su ciò che è bene, hanno cioè diversi *giudizi di valore*, di cui non è facile dire con certezza, sulla base di ragionamenti puramente razionali, se uno sia in via assoluta preferibile all'altro. Intervengono fattori ideologici, motivazioni pre-scientifiche, interessi, su cui la teoria economica e in generale la scienza non può sempre dare un contributo.

Nello studio e nelle discussioni di economia pubblica proposizioni di tipo positivo e di tipo normativo si presentano spesso insieme e possono essere molto intrecciate le une alle altre. È però essenziale sapere distinguerle correttamente: in tale capacità di distinzione si ha spesso il contributo più significativo che l'economista può dare alla spiegazione dei fenomeni sociali, aiutando a distinguere ciò che è risultato della ricerca della conoscenza e ciò che dipende da giudizi di valore. Quando si è di fronte a due proposizioni positive contrastanti si può sempre sperare che una teoria più elaborata o una più adeguata verifica empirica possano aiutarci nella ricerca della verità. Di fronte ad una proposizione normativa, ci si deve di solito fermare a giudizi del tipo: «se si ritiene che il bene sia  $x$ , allora è opportuno seguire queste politiche; se si ritiene che il bene sia  $y$ , allora sono appropriate quest'altre politiche».

L'importanza della distinzione tra teoria positiva e normativa è così grande in economia che la teoria normativa è diventata una disciplina a sé stante e parallela rispetto alla teoria positiva e la si è chiamata Economia del benessere (*Welfare Economics*), che si è così affiancata alla teoria dell'Equilibrio economico generale, il paradigma scientifico dominante. La microeconomia ci ha abituato a presupporre che il comportamento degli individui sia razionale: i consumatori massimizzano le loro preferenze, dati i vincoli di bilancio; i produttori massimizzano i profitti, in un contesto istituzionale in cui sono definiti i diritti di proprietà (le dotazioni iniziali di risorse), sono note le tecniche produttive, i costi dei fattori e le regole del gioco, il mercato concorrenziale. Questa teoria spiega come si determinano i prezzi dei beni e le quantità prodotte in un'economia decentrata. Se le ipotesi dette corrispondono alla realtà, la produzione dei beni e il loro scambio avvengono nel rispetto di alcune condizioni: per ciascun individuo il tasso marginale di sostituzione tra due beni, vale a dire il rapporto tra le loro utilità marginali, eguaglia il rapporto tra i loro prezzi; per le imprese, il tasso marginale di sostituzione tra due fattori, il rapporto tra le loro produttività marginali, eguaglia il rapporto dei loro prezzi; per l'economia nel suo complesso il saggio marginale di trasformazione, che definisce la quantità di un bene a cui si deve rinunciare se si desidera produrre un'unità in più di un altro bene, eguaglia i tassi marginali di sostituzione. Il modello così costruito ci fornisce, quindi, un'interpretazione della realtà studiata, dei prezzi e delle quantità di beni prodotti, che risultano spiegati dalle preferenze, dalle tecniche produttive e dall'ambiente istituzionale di un'economia decentrata.

L'Economia del benessere, la teoria economica normativa, le cui basi sono state costruite sostanzialmente nel ventennio tra il 1930 e il 1950, studia il funzionamento di un'economia di produzione e di scambio da un altro punto di vista. Essa si domanda quale debba essere la configurazione ottimale di un sistema economico, in cui siano presenti più individui, con diversi sistemi di

preferenze e diverse dotazioni iniziali di fattori (capacità lavorativa e capitali) e di beni. Il suo fine è definire un ottimo sociale, vale a dire la quantità di beni da produrre e la distribuzione degli stessi che consentono di realizzare una situazione di massimo benessere collettivo. È una teoria normativa perché vuole definire ciò che è bene o è male, fornire prescrizioni per il raggiungimento di determinati fini, individuare *norme* da seguire, fornire regole che consentono di ordinare diversi stati sociali sotto il profilo della loro preferibilità per la collettività.

L'Economia del benessere si fonda su presupposti molto generali. La ricerca è svolta in un contesto istituzionale vuoto; non si fa cioè alcuna ipotesi sulle caratteristiche del modo di produzione della società considerata, che può essere sia un'economia decentrata sia un'economia pianificata.

Con l'Equilibrio economico generale ha in comune una visione di tipo *individualistico*, dato che pone al centro dell'attenzione l'individuo razionale, che è quindi considerato il migliore giudice di se stesso, e di tipo *utilitaristico*, dato che il benessere dell'individuo dipende esclusivamente dalla quantità dei beni da lui stesso consumati.

L'Equilibrio economico generale solitamente non pone attenzione alla natura e al ruolo dello Stato nella società. L'Economia del benessere invece formula ipotesi a questo riguardo, sposando un'impostazione *non organica della società*. Lo Stato esiste, ma non è un'autonoma fonte di valori. La sua volontà è nulla più di quella che risulta dall'aggregazione delle volontà degli individui che ne fanno parte.

Nella definizione dei fini a cui la società deve tendere – l'ottimo sociale – con riferimento alla produzione e allo scambio, l'Economia del benessere attribuisce un particolare valore al *principio di efficienza*, noto anche come *principio di Pareto*.

Qual è il contenuto del principio di Pareto? Con riguardo alla *produzione*, il principio di efficienza è strettamente connesso all'idea di razionalità. Se gli uomini sono razionali, dovrebbe esservi unanime consenso che se, a parità di impiego di fattori, con la tecnica A posso produrre  $x$  beni e con la tecnica B posso produrre  $y = x + 1$  beni, la tecnica B sia preferibile alla tecnica A.

Con riguardo allo *scambio*, il principio di Pareto afferma che, dovendo distribuire tra gli individui della società una certa quantità di beni, una riallocazione delle risorse che migliori il benessere di un individuo senza arrecare danno agli altri rappresenta un miglioramento del benessere per la società. La formulazione dell'efficienza nello scambio, logicamente simmetrica a quella dell'efficienza nella produzione, pur apparendo largamente condivisibile, comporta tuttavia l'introduzione di un giudizio di valore, in cui sono implicite valutazioni distributive. La sua accettazione, ad esempio, cessa di essere unanime se gli individui non sono motivati solo da comportamenti egoistici, ma, come accade in talune circostanze, anche dall'altruismo o, al contrario, dall'invidia.

La razionalità dell'efficienza produttiva e la plausibilità dell'efficienza nello scambio sono sembrate agli studiosi di Economia del benessere punti di riferimento utili per porre le basi della ricerca dell'ottimo sociale, tali da motivare gli sforzi per individuare le situazioni in cui tali principi possono essere realizzati (i punti di ottimo paretiano). Deve però essere ben chiaro che



concentrare l'attenzione solo sui punti di ottimo paretiano non ci consente di fare molti passi avanti nell'indagine normativa con riferimento ai problemi che pone la vita associata, in cui quasi sempre la scelta tra diverse possibili linee di azioni o riforme comporta vantaggi per alcuni individui e svantaggi per altri. In questi casi il principio di Pareto non è in grado di dirci se il risultato dell'azione è preferibile, dal punto di vista sociale, rispetto allo *status quo*. L'economista a questo punto di solito si arresta. La teoria economica ha infatti dimostrato che, se si vuole restare rigorosamente fedeli ai presupposti descritti dell'Economia del benessere (razionalità, individualismo, ecc.) e al principio di efficienza paretiana, è *impossibile* arrivare a definire una regola di decisione collettiva e cioè una regola di scelta fra diversi stati sociali che conduca al massimo benessere della società (teorema dell'impossibilità di Arrow).

Accanto al principio di efficienza paretiana si devono quindi prendere in considerazione altri principi, che riflettono giudizi di valore, nozioni di *equità* e di *giustizia*. Si tratta però di un terreno molto difficile da esplorare per l'economista, che lo porta a sconfinare nel campo della filosofia morale. Esistono infatti diverse nozioni di equità, tutte potenzialmente rilevanti ma non sempre tra loro compatibili. A questo punto gli economisti generalmente si limitano a descrivere e ad assumere possibili alternative formulazioni della giustizia e a trarne le conseguenze, in termini di azioni di riforma. In questo modo le conclusioni dell'Economia del benessere risultano condizionali alla versione del principio di equità che si è di volta in volta assunta. Gli esiti della ricerca dipendono da giudizi di valore.

La pura razionalità economica non è sufficiente per condurci alla determinazione dell'ottimo sociale. Si deve concludere che questo percorso di ricerca è approdato ad un vicolo cieco e si è rivelato infruttuoso? Vi è chi ha espresso posizioni critiche di questo tipo, ma è comunque giusto riconoscere che l'impianto teorico descritto ha il pregio di indurre chi si avventuri nella valutazione sociale delle politiche economiche a distinguere accuratamente le proposizioni che possono essere compatibili con i fondamenti dell'Economia del benessere e il principio di efficienza paretiana, su cui molti potrebbero trovarsi d'accordo, e le proposizioni che dipendono in modo più pesante da giudizi di valore relativi alla distribuzione delle risorse tra i membri della società. Uno degli aspetti più interessanti di questo modo di ragionare è mostrare se e quando vi sia compatibilità tra efficienza ed equità, o, per usare l'espressione inglese, se esista o meno un *trade-off* tra efficienza ed equità.

L'economista sembra quindi avere circoscritto il proprio campo di indagine, rivolto allo studio dell'efficienza, consapevole che questa è solo una parte della storia. Un atteggiamento prudente che va salutato con favore. Ma non si deve dimenticare che anche questo campo base, su cui ci si attesta nell'impossibilità di raggiungere la meta con gli strumenti a disposizione, non è affatto un terreno sicuro, in cui tutti possano ritrovarsi. Il concetto di efficienza paretiana appena descritto si inserisce pur sempre in un contesto in cui prevale una visione individualistica e utilitaristica che non consente di cogliere importanti dimensioni del comportamento umano.

In secondo luogo, e in termini più generali, soprattutto in riferimento alla teoria economica positiva, in un terreno in cui sono in gioco aspetti

riguardanti diverse scienze umane, bisogna imparare a riconoscere le deformazioni professionali. In genere un economista è molto orgoglioso di poter fornire spiegazioni puramente economiche dei problemi che studia. È di solito un titolo di merito mostrare, ad esempio, che un fenomeno, spiegabile a prima vista solo sulla base di ragioni di equità (si pensi all'esempio fatto sopra delle politiche contro la povertà), non è l'esito di scelte motivate da valori estranei alla teoria neoclassica (ad es. l'altruismo), ma dipende dalla presenza di più complesse logiche compatibili con i presupposti della teoria positiva dominante. Molto spesso questa ricerca può però portare ad esasperazioni nella formulazione delle ipotesi e nella costruzione dei modelli, di cui si finisce per privilegiare più l'aspetto formale che la capacità interpretativa.

### 1.3. I due teoremi dell'Economia del benessere e l'ottimo sociale

È utile ora richiamare in modo più puntuale i principali risultati dell'Economia del benessere e del processo logico che porta all'individuazione dell'ottimo sociale. L'esposizione è finalizzata solo al richiamo dei risultati della teoria: il lettore interessato può fare riferimento ad un buon testo di microeconomia per trovare l'illustrazione del percorso analitico necessario per ottenere questi risultati o consultare gli aggiornamenti disponibili sul sito [www.mulino.it/aulaweb](http://www.mulino.it/aulaweb).

*La determinazione dell'ottimo sociale.* Come possiamo definire una situazione di ottimo sociale? La risposta non è univoca, perché l'eterogeneità degli individui pone il problema di definire: ottimo per chi? ottimo rispetto a che cosa? Per affrontare questo problema gli economisti del benessere (sulla scia di Bergson e Samuelson) hanno introdotto il concetto di *funzione del benessere sociale*.

Con riferimento a una società composta da due soli individui, essa può essere così scritta:

$$W = W(U_1, U_2)$$

Si tratta di una funzione che ha come argomenti le utilità degli individui (1 e 2) e che misura il benessere collettivo. La formulazione descritta dalla funzione  $W$  è del tutto generale, ma, come si è spiegato sopra, la sua forma dipende dalle nozioni di *equità* che si è deciso di assumere nell'analisi. Rappresenta la valutazione che la società, ad esempio attraverso la classe governante, sia essa eletta democraticamente o meno, ha dell'equità, che nella formulazione qui seguita si manifesta sotto forma di valutazioni sulle modalità con cui i beni vengono distribuiti ai due individui, determinandone i rispettivi livelli di benessere.

Una volta assunta una data funzione  $W$ , che l'Economia del benessere si limita a presupporre nella propria ricerca dell'ottimo sociale, si possiede la chiave per valutare in termini normativi qualsiasi problema di politica economica, e per affermare se un certo stato sociale, ad esempio quello che risulta descritto dalla distribuzione delle risorse in seguito all'introduzione

di una misura di politica economica, sia o meno preferibile ad un altro, ad esempio lo *status quo*, oppure ad uno stato del mondo in cui si prenda una misura alternativa.

Abbiamo già messo in guardia sui presupposti filosofici che sono alla base di tale funzione. Essa implica che il benessere sociale dipenda esclusivamente dall'utilità dei soggetti che appartengono alla società. Si tratta quindi di un concetto del tutto coerente con i primi due presupposti ricordati nel precedente paragrafo. Si assume, cioè, un punto di vista individualista e non si ammette la possibilità di definire il bene della società se non in funzione del benessere degli individui che la compongono. Non esiste, come vorrebbero le teorie organicistiche della società, il «bene dello Stato» in sé. Anche se la classe di governo non traesse la propria legittimazione da un processo di voto democratico, come ad esempio nel caso di una dittatura, si suppone che il governante valuti il benessere della società solo in funzione del benessere dei cittadini (e non in base ad altri valori, ad esempio il volere di una divinità o il patrimonio del dittatore).

In secondo luogo, nella quasi totalità delle sue applicazioni, il benessere degli individui, vale a dire gli argomenti della funzione  $W$ , è rappresentato dai beni e servizi (qui si suppone siano due) di cui gli individui dispongono e possono godere:

$$U_i = U_i(X_{i1}, X_{i2}) \quad i = 1, 2$$

ove  $X_{ij}$  rappresenta la quantità del bene  $j$  posseduta e consumata dall'individuo  $i$ . Si accetta cioè quasi sempre un punto di vista egoistico e non si ammette che il benessere di un individuo possa dipendere dal benessere dell'altro.

Immaginiamo allora l'esistenza di una funzione  $W$ . Possiamo porci un'ulteriore domanda. Date certe dotazioni iniziali di fattori e di beni, date certe tecniche produttive, come dovrebbero essere organizzate la produzione e la distribuzione dei beni tra i vari individui della società al fine di realizzare il massimo benessere sociale?

La teoria ha raggiunto una prima conclusione importante. *Pur senza alcun bisogno di caratterizzare il contesto istituzionale*, vale a dire senza dover specificare se ci troviamo in un'economia decentrata in cui esistono imprese private o in un'economia a pianificazione centralizzata, è possibile definire alcune condizioni che devono essere rispettate affinché si possa raggiungere il massimo benessere collettivo.

Consideriamo due individui (1 e 2), due beni ( $X_1$  e  $X_2$ ) e due fattori ( $K$  e  $L$ ). In termini formali il problema è quello di massimizzare il benessere sociale

$$\max W = W(U_1, U_2)$$

tenendo conto dei vincoli costituiti dalle preferenze individuali, dalle tecniche di produzione e dalle dotazioni e diritti sui fattori produttivi.

$$\begin{aligned} U_1 &= U_1(X_{11}, X_{12}) \\ U_2 &= U_2(X_{21}, X_{22}) \end{aligned} \quad \text{Preferenze}$$

$$\begin{aligned} X_1 &= X_1(K_1, L_1) \\ X_2 &= X_2(K_2, L_2) & \text{Tecnologia} \\ K_1 + K_2 &= K \\ L_1 + L_2 &= L & \text{Risorse} \end{aligned}$$

e ove

$$\begin{aligned} X_{11} + X_{21} &= X_1 \\ X_{12} + X_{22} &= X_2 \end{aligned}$$

Tale problema è stato affrontato e risolto per la prima volta da un economista italiano all'inizio del '900, Enrico Barone, e poi anche da Oskar Lange, un famoso studioso di sistemi pianificati, e in seguito perfezionato da altri economisti.

Essi hanno derivato alcune condizioni che possiamo tenere *distinte in due gruppi*.

Le condizioni del primo gruppo sono note come *condizioni di efficienza paretiana*. Nella situazione di ottimo:

- il saggio marginale di sostituzione tra i beni  $X_1$  e  $X_2$  dell'individuo 1 è identico all'analogo saggio marginale di sostituzione dell'individuo 2 (efficienza nello scambio);
- i saggi marginali di sostituzione tecnica dei fattori produttivi  $K$  e  $L$  sono uguali nella produzione di entrambi i beni (efficienza nella produzione);
- per il sistema nel suo complesso, il saggio marginale di trasformazione tra due beni calcolato sulla frontiera della produzione eguaglia il tasso marginale di sostituzione dei due individui (efficienza nella composizione del prodotto).

Se queste tre condizioni sono rispettate, siamo in presenza di una *situazione Pareto efficiente*, e cioè di una situazione di ottimo paretiano in cui non è possibile riorganizzare la produzione e la distribuzione dei beni in modo da migliorare la posizione di benessere di un individuo senza danneggiare quella di un altro individuo.

L'insieme delle possibili posizioni Pareto efficienti è *infinito* e viene solitamente rappresentato in un grafico molto importante, in cui nel piano ( $U_1, U_2$ ) sono rappresentati i livelli di benessere dei due individui della società, ed è tracciata la grande frontiera dell'utilità (la curva  $BCPDE$  della figura 1.1). Tutti i punti che giacciono sulla curva rappresentano punti di efficienza paretiana e vengono anche indicati come posizioni di First Best. I punti situati all'interno della frontiera, ad esempio il punto  $A$ , non rispettano invece una o più delle condizioni di efficienza paretiana. Ciò significa che sarebbe possibile, mediante appropriate modificazioni delle combinazioni produttive disponibili, della composizione della produzione e della sua distribuzione, migliorare il benessere sociale, fino a raggiungere posizioni Pareto efficienti, situate sulla grande frontiera. Prendendo come punto di partenza il punto  $A$ , che non è un punto di First Best, tutte le posizioni che si situano nell'area  $ACD$  sono indicate come miglioramenti paretiani. Di queste, solo quelle che giacciono sul tratto della frontiera  $CD$  rappresentano situazioni di First Best,



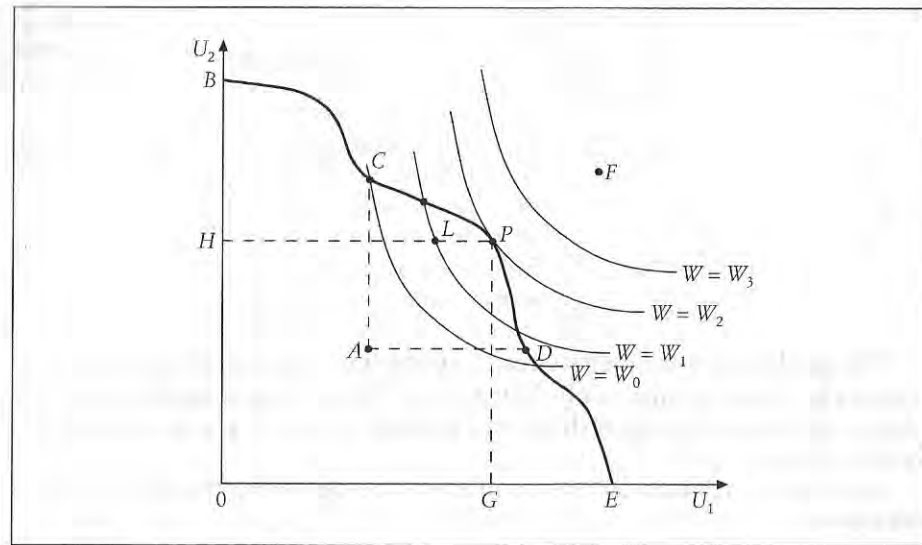


FIGURA 1.1. Determinazione dell'ottimo sociale.

in cui non è possibile aumentare il benessere di un individuo senza ridurre quello dell'altro. I punti al di fuori della grande frontiera, ad esempio il punto  $F$ , descrivono invece livelli di benessere degli individui che non sono raggiungibili dall'economia considerata con le tecniche e le risorse disponibili.

Le condizioni del secondo gruppo riguardano la funzione del benessere sociale, la quale – fatto molto importante alla luce di quanto diremo – non compare nella definizione delle condizioni del primo gruppo. Esse affermano che nella situazione di ottimo sociale l'utilità marginale sociale prodotta dall'incremento di un'unità addizionale di bene,  $X_1$  o  $X_2$ , messa a disposizione di un individuo deve essere uguale per ogni individuo. Partendo da una data grande frontiera dell'utilità e da una funzione del benessere sociale, questo secondo gruppo di condizioni consente di individuare il punto della grande frontiera dell'utilità a cui corrisponde il massimo benessere sociale.

Nella figura 1.1 sono rappresentate infatti alcune curve di livello (del tutto analoghe alle curve di indifferenza studiate nella teoria della domanda in microeconomia, con la differenza che in questo caso gli argomenti della funzione  $W$  sono le utilità degli individui della società e non le quantità dei beni consumati dal soggetto; il soggetto titolare della funzione del benessere è il governo, mentre nella teoria della domanda è il consumatore). Tutti i punti che giacciono su una curva (ad esempio quella a cui corrisponde un livello di benessere sociale pari a  $W = W_0$ ) definiscono un identico livello di benessere sociale. Esistono infinite curve di livello non intersecantesi. Il benessere sociale aumenta via via che ci si sposta su curve collocate sempre più a nord-est ( $W_1$ ,  $W_2$ , ecc.).

L'ottimo sociale è definito dal punto  $P$ , ove la curva di benessere sociale con indice più elevato risulta tangente alla grande frontiera dell'utilità. Tra tutti gli infiniti punti Pareto efficienti, l'ottimo sociale risulta essere un solo punto. In  $P$  si ha un livello di benessere sociale pari a  $W = W_2$ , a cui corri-

sponde una distribuzione delle risorse tra i due individui tale da generare un livello di benessere individuale rispettivamente pari a  $OG$  e  $OH$ .

In questo processo di individuazione dell'ottimo sociale è importante osservare che i problemi di efficienza paretiana (il primo gruppo di condizioni sopra ricordate che concorrono alla determinazione della grande frontiera dell'utilità) risultano separati da quelli relativi alla distribuzione (la scelta del punto sulla grande frontiera dell'utilità, utilizzando la funzione del benessere sociale). La separabilità dei due gruppi di condizioni per l'esistenza di un ottimo sociale è di estrema importanza nello studio dell'economia pubblica, in quanto costituisce il fondamento della distinzione, che abbiamo ricordato nel precedente paragrafo, tra aspetti di efficienza (secondo l'accezione paretiana) ed equità. È tale distinzione che suggerisce all'economista il ruolo di mostrare quali siano i punti sulla grande frontiera dell'utilità, caratterizzati dall'efficienza (nel senso paretiano), lasciando al filosofo morale il compito di definire la forma della funzione di benessere sociale di una data classe di governo, e a quest'ultimo l'individuazione della soluzione, rappresentata dal punto  $P$ .

*I due teoremi dell'Economia del benessere.* Partendo da questo impianto analitico è poi stato possibile dimostrare due importanti teoremi, noti appunto come i due teoremi dell'Economia del benessere.

Il primo teorema dell'Economia del benessere afferma che le condizioni di efficienza paretiana (nella produzione, nello scambio e nella composizione del prodotto) sono realizzate in una particolare configurazione istituzionale: un'economia decentrata di concorrenza perfetta. In altre parole, se lasciamo operare un mercato in condizioni di concorrenza perfetta, come quella descritta dalla teoria positiva dell'Equilibrio economico generale, partendo da date dotazioni iniziali, l'allocatione dei fattori produttivi e la produzione dei beni che ne emergono sono Pareto efficienti. L'importanza di questo teorema non va sottovalutata: esso sancisce che un'economia di mercato di concorrenza perfetta possiede caratteristiche di ottimalità paretiana.

Da ciò discendono conseguenze interessanti in termini normativi. Se un sistema economico si fonda su un'economia di mercato decentrata appare conveniente favorire tutte le iniziative che consentono di realizzare condizioni di concorrenza perfetta, perché ad esse corrisponde una situazione di ottimo paretiano. In questo senso la dimostrazione del primo teorema dell'Economia del benessere ci indica già una motivazione dell'intervento pubblico (come ad esempio le politiche di liberalizzazione dei mercati, abolizioni di posizioni di rendita e di monopolio legale, ecc.).

La soluzione Pareto efficiente definita dal mercato concorrenziale dipende però dalla distribuzione iniziale delle risorse. Partendo da una diversa distribuzione, si arriva ad una diversa soluzione sulla grande frontiera dell'utilità, anch'essa Pareto efficiente, ma caratterizzata da una diversa distribuzione del benessere tra gli individui e quindi da un diverso valore della funzione del benessere sociale. Dal punto di vista sociale il mercato di concorrenza perfetta permette di realizzare l'efficienza paretiana, ma non può risolvere il problema della distribuzione ottimale del benessere tra gli individui. In altre parole, dal punto di vista distributivo la soluzione individuata da un'economia di mercato solo per caso può coincidere con quella di ottimo sociale.

Nell'esempio descritto dalla figura 1.1, si può immaginare che, data una certa distribuzione iniziale delle risorse, il funzionamento del mercato concorrenziale consenta di realizzare la situazione descritta dal punto  $C$ , sulla grande frontiera delle utilità. A tale punto corrisponde un livello di benessere sociale pari a  $W = W_0$  che non è però quello corrispondente all'ottimo sociale (punto  $P$ ). Il punto  $C$  (di First Best) appare preferibile al punto  $A$ , che non è Pareto efficiente, ma non corrisponde all'ottimo sociale.

Alla difficoltà appena accennata sembrerebbe venire in soccorso il secondo teorema dell'Economia del benessere.

Il secondo teorema dell'Economia del benessere afferma che, modificando opportunamente le dotazioni iniziali con particolari strumenti di redistribuzione, imposte o sussidi in somma fissa (*lump sum taxes*), un'economia concorrenziale consente di raggiungere qualsivoglia stato sociale Pareto efficiente sulla grande frontiera dell'utilità. Grazie a questo teorema si può allora fornire un ulteriore motivo per sostenere la validità normativa di un'economia decentrata di tipo concorrenziale. Lasciandola operare si risolve infatti il problema dell'impiego efficiente delle risorse (problema in cui hanno fallito le economie pianificate dal centro) e si è quindi in grado di raggiungere un punto sulla grande frontiera dell'utilità. La soluzione offerta dal mercato può però non essere gradita dal punto di vista distributivo: può darsi infatti che, sulla base della funzione del benessere sociale, in quella società si preferisca un altro punto della grande frontiera dell'utilità. Il problema può allora essere risolto utilizzando *lump sum taxes*, che consentono di passare alla soluzione ottimale individuata dalla funzione del benessere sociale, come punto di tangenza fra tale funzione e la grande frontiera dell'utilità.

Ancora nella figura 1.1 l'uso di *lump sum taxes* potrebbe consentire di spostare l'equilibrio dal punto  $C$  al punto  $P$ .

Cos'è una *lump sum tax*? È una forma di imposta/sussidio che redistribuisce risorse senza influenzare i segnali (i prezzi relativi) che i consumatori e i produttori hanno come punto di riferimento nel compiere le proprie scelte in un mercato concorrenziale. Essa consente di evitare distorsioni nei comportamenti dei soggetti e non violare le tre condizioni di efficienza paretiana. Una *lump sum tax* è dunque caratterizzata dal fatto di essere commisurata a fattori che sono esogeni, cioè fuori dal controllo dell'individuo a cui viene applicata. Si pensi, ad esempio, ad imposte commisurate ad abilità naturali degli individui. Se ciò non accadesse, le imposte potrebbero generare reazioni di comportamento che allontanano il sistema dalle condizioni di efficienza paretiana. Se, ad esempio, si introduce un'imposta commisurata al consumo di un bene, il prezzo con cui i consumatori si confrontano non rifletterà solo le condizioni di produzione (costi marginali) e le preferenze degli individui (utilità marginali), ma anche il livello dell'aliquota dell'imposta. La condizione di impiego efficiente del reddito sarebbe violata. Le uniche forme di imposte o sussidi ammissibili sono dunque strumenti non correlati a variabili rilevanti per le scelte economiche degli operatori, e che dovranno essere definiti in misura diversa per ciascun individuo della società al fine di definire la dotazione iniziale compatibile con l'ottimo sociale che si desidera raggiungere.

La realizzazione di un sistema di imposte/sussidi di questo tipo è chiaramente impossibile. Il governante, cioè il soggetto politico che ha il controllo della funzione del benessere sociale e il compito di massimizzarla, dovrebbe disporre di numerose informazioni sugli individui della società, ad esempio, conoscere le funzioni di preferenza di tutti, ed essere in grado di individuare strumenti fiscali che ben difficilmente potrebbero funzionare senza costi di gestione enormi.

Se, per correggere la distribuzione realizzata dal mercato concorrenziale, il governo non è in grado di usare strumenti neutrali come le *lump sum taxes*, ma utilizza imposte che generano distorsioni, l'intervento correttivo comporta l'abbandono della grande frontiera e quindi di soluzioni di First Best. Nella figura 1.1, posto che  $C$  sia il punto di First Best realizzato dal mercato concorrenziale, l'intervento correttivo da parte dello Stato mediante strumenti distorsivi potrebbe portare alla soluzione descritta dal punto  $L$ , che giace sulla curva di benessere sociale con indice  $W = W_1$ . Si noti che, nell'esempio fatto, al punto  $L$  corrisponde un livello di benessere sociale che è superiore a quello del punto  $C$ , anche se  $L$  non rappresenta un punto di First Best, ma, come viene solitamente indicato, un punto di Second Best. Per realizzare una distribuzione del reddito più coerente con la funzione del benessere sociale, si è sacrificata l'efficienza. Per usare l'efficace immagine proposta negli anni '60 dall'economista americano A. Okun, lo Stato, nella ricerca di una situazione più equa, trasferisce risorse usando un secchio con il fondo bucato: la redistribuzione necessariamente causa perdite per la società vista nel suo complesso.

In realtà, il valore del secondo teorema è prevalentemente di tipo negativo. Ci insegna cioè che se ci poniamo nella prospettiva di privilegiare l'economia di mercato come strumento per realizzare l'ottimo sociale, anche se l'economia reale corrispondesse esattamente ad un sistema di concorrenza perfetta (e già questo è un bel problema!), il problema distributivo non sarebbe ancora risolto. La conclusione da trarre è che, sebbene l'economia decentrata e concorrenziale assolva un ruolo importante che può essere valutato con favore dal punto di vista normativo, essa non è in grado di garantirci il raggiungimento della soluzione desiderata di First Best sulla grande frontiera dell'utilità.

Nella realtà in cui viviamo ci troveremo quasi sempre in situazioni di Second Best, dovute o all'assenza delle condizioni di concorrenza perfetta o all'uso di strumenti non neutrali messi in atto per realizzare equilibri distributivi coerenti con la funzione del benessere sociale.

Queste conclusioni hanno senza dubbio l'effetto di deludere la grande aspirazione dei pionieri dell'Economia del benessere di poter separare i problemi dell'efficienza da quelli della distribuzione e di consacrare il valore normativo del mercato di concorrenza perfetta. Dal punto di vista del metodo la rilevanza della distinzione tra efficienza ed equità conserva tuttavia un notevole valore euristico e può utilmente servire da guida nell'analisi dei problemi di finanza pubblica. La realtà non si conformerà mai al modello della concorrenza perfetta, ma nello studiarla può essere utile avere come punto di riferimento, come *benchmark*, i risultati di queste elaborazioni teoriche, oltre che i limiti della loro validità.



#### 1.4. L'interpretazione del primo teorema dell'Economia del benessere in equilibrio parziale

Le condizioni che devono essere soddisfatte perché un'economia concorrenziale goda della proprietà di ottimo paretiano non sono di troppo ardua comprensione; è però complesso l'uso, per lo studio di effetti di politiche fiscali, dello schema analitico da cui esse sono dedotte. Anche quando ci si limiti al caso più semplice di due soli individui, due beni, due fattori, si tratta pur sempre di governare un modello di equilibrio economico generale, in cui sono presenti interdipendenze tra variabili e vincoli che le legano che non sono facili da rappresentare in modo semplice. È forse per questa ragione che nella didattica vengono spesso utilizzati strumenti meno complessi, ma anche meno rigorosi. Tali sono i concetti di surplus del consumatore e del produttore, elaborati nell'ambito di schemi di equilibrio parziale, vale a dire di un'economia costituita dal mercato di un solo bene. Può essere utile richiamare questa nozione, già studiata nei corsi introduttivi di microeconomia, per illustrare successivamente alcune corrispondenze concettuali tra un'analisi di equilibrio generale e una di equilibrio parziale.

In un'analisi di equilibrio parziale, indicata anche come analisi marshalliana, in onore dell'economista inglese che, alla fine del XIX secolo, l'ha introdotta nella teoria economica, l'attenzione viene concentrata su una singola industria, che si suppone produca un bene omogeneo. Gli attori dell'economia sono costituiti dai consumatori e dai produttori del bene  $X$ , che esprimono domanda e offerta del bene in un mercato, che considereremo di concorrenza perfetta, in cui cioè esiste una molteplicità di consumatori e produttori, che massimizzano rispettivamente utilità e profitti, e che assumono il prezzo come un dato (*price-taker*) (si veda la figura 1.2a).

La domanda  $D$  è il risultato dell'aggregazione delle quantità domandate da tutti i consumatori. Per ogni livello di prezzo, il consumatore domanderà il bene sino al punto in cui il prezzo è uguale alla valutazione-utilità *marginale* data al bene. Il trapezio  $0Q'\beta\gamma$  è il *surplus lordo dei consumatori*, che può essere assunto come una rappresentazione del benessere dei consumatori. Quando il prezzo di  $X$  è  $P'$ ,  $Q'$  è infatti la domanda e il consumo; nella misura in cui, con una certa perdita di rigore, la domanda possa essere assunta come una rappresentazione dell'utilità marginale del bene in questione, l'utilità complessiva derivabile dal consumo è pari all'area indicata. Per godere di tale surplus lordo i consumatori pagano una somma pari all'area  $0Q'\beta P'$ . Si definisce quindi *surplus netto dei consumatori* la differenza tra il surplus lordo e tale costo, pari al triangolo  $P'\beta\gamma$ . Il surplus netto è positivo perché la valutazione che i consumatori danno alle unità *inframarginali* del bene in questione, per la nota decrescenza dell'utilità marginale di un bene, è maggiore della valutazione data all'unità di consumo marginale, pari al prezzo, unico per tutti i consumatori.

L'offerta  $S$ , con ragionamento del tutto simmetrico, è il risultato dell'aggregazione delle curve dei costi marginali dei produttori. In concorrenza, essendo *price-takers*, essi fissano l'offerta al punto in cui il costo marginale è uguale al prezzo. Per la combinazione prezzo-quantità  $(P', Q')$ , l'area  $0Q'\beta P'$  rappresenta il ricavo della produzione  $Q'$ , che può essere interpretata come

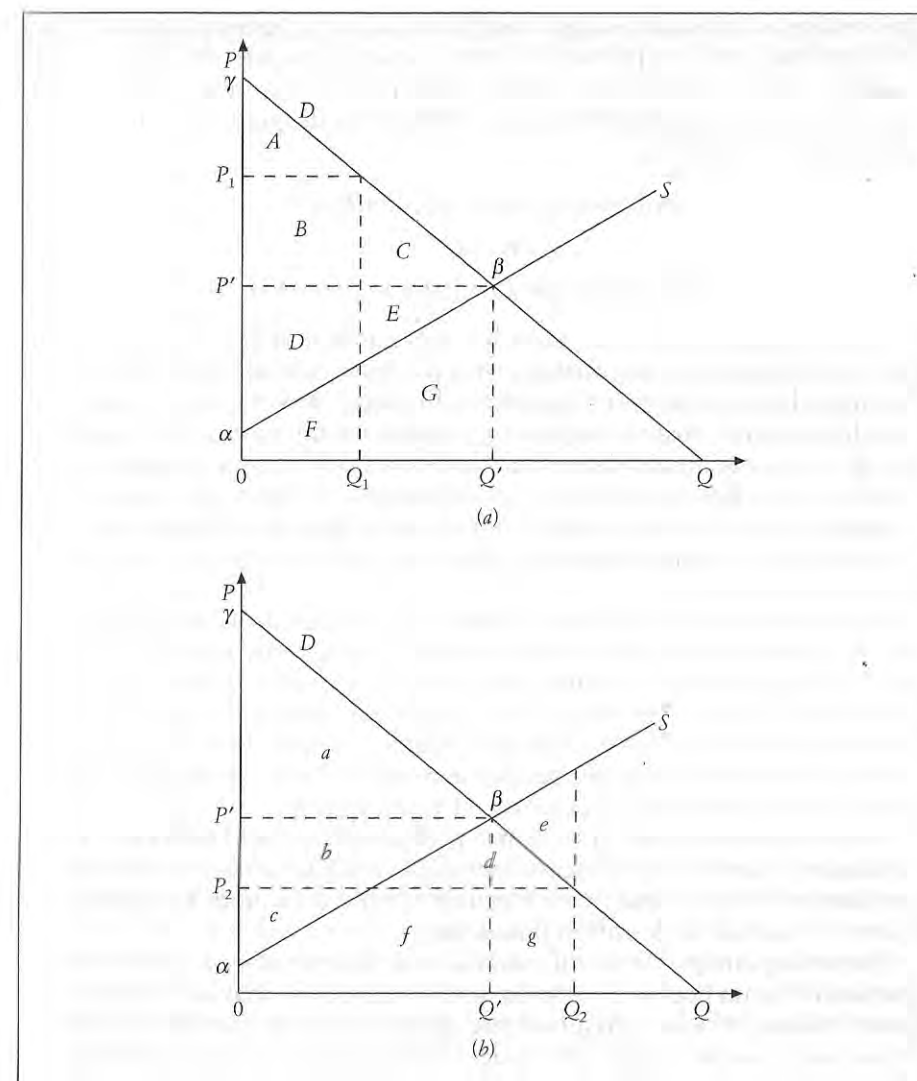


FIGURA 1.2. Surplus del consumatore e del produttore e ottimo paretiano.

misura del *surplus lordo dei produttori*. L'area  $0Q'\beta\alpha$ , sottostante alla curva di offerta, è invece rappresentativa del costo complessivo che i produttori devono sopportare per realizzare l'offerta  $Q'$ . La differenza pari al triangolo  $\alpha\beta P'$  è definita *surplus netto dei produttori*. Tale surplus è positivo perché esistono produttori (tutti ad esclusione dei produttori marginali per i quali i costi marginali sono uguali a quelli medi) per i quali il costo marginale, uguale al prezzo, risulta essere superiore al costo medio in corrispondenza alla quantità da essi offerta.

Questo schema familiare presenta alcune fondamentali analogie con lo schema concorrenziale di equilibrio economico generale e consente di

stabilire alcune utili connessioni con l'Economia del benessere studiata nei precedenti paragrafi. In particolare le condizioni ricordate per definire l'ottimalità paretiana della concorrenza (efficienza nello scambio, nella produzione e complessiva) sono traducibili all'interno di questo schema con la fondamentale condizione:

$$\text{Costo marginale dei produttori} =$$

$$\text{Prezzo} =$$

$$\text{Valutazione marginale dei consumatori}$$

A ciò infatti si riduce, in un mondo in cui esista un solo bene, l'uguaglianza dei tassi marginali di sostituzione con i prezzi relativi dei beni (valutazione marginale dei consumatori = prezzo) e l'uguaglianza tra i saggi marginali di trasformazione (costo marginale = prezzo). In un contesto di equilibrio parziale, in un'economia decentrata, la ricerca dell'ottimo paretiano equivale alla ricerca delle condizioni che consentano di realizzare l'eguaglianza tra prezzo, che è la misura della valutazione marginale del bene attribuita dai consumatori, e costo marginale, che misura il sacrificio di risorse per la produzione del bene.

Poiché in questo universo esistono solo consumatori e produttori, è naturale assumere come misura del benessere complessivo dell'economia la somma del surplus netto dei consumatori e dei produttori. Esso è quindi pari all'area  $\alpha\beta\gamma$ . È bene avere sempre ben presenti le ragioni che lo producono: le rendite di alcuni consumatori, per i quali la valutazione marginale dei beni acquistati è superiore al prezzo e le rendite di produttori, per i quali il prezzo = costo marginale è superiore al costo medio.

Produttori interessati al massimo profitto decidono l'offerta in corrispondenza al punto in cui il ricavo marginale è uguale al costo marginale. In concorrenza, il ricavo marginale è uguale al prezzo e quindi la concorrenza rispetta le condizioni di ottimo paretiano.

Questo aspetto può essere illustrato con lo schema di equilibrio parziale, mostrando che un prezzo diverso da quello realizzabile in concorrenza (nella figura il prezzo  $P'$  e la corrispondente quantità  $Q'$ ) produrrebbe un benessere sociale (misurato dalla somma dei surplus di produttori e consumatori) inferiore a quello realizzabile in concorrenza, oltre a causare redistribuzioni di surplus tra categorie di soggetti (nel nostro caso consumatori e produttori).

Immaginiamo che, a causa di qualche imperfezione del mercato (ad esempio il potere monopolistico dei produttori) venga fissato un prezzo del bene pari a  $P_1 > P'$ . Sempre con riferimento alla figura 1.2a, la quantità domandata si ridurrà a  $Q_1$ . Il surplus netto dei consumatori – che ora è più conveniente rappresentare con le aree indicate nella figura con le lettere maiuscole –, prima pari a  $A + B + C$ , si ridurrà al triangolo  $A$ , con una variazione pari a  $-(B + C)$  (si veda anche la tabella 1.1). I produttori, il cui surplus netto era  $D + E$ , godranno ora di un surplus netto pari a  $B + D$ , con una variazione pari a  $+B - E$ . Non è possibile dire se esso sia maggiore o minore di quello precedente: dipende dalla dimensione relativa delle aree  $B$  ed  $E$ . Si noti tuttavia che l'area  $B$  è una parte del surplus dei consumatori che i produttori, con l'aumento dei prezzi, hanno sottratto ai consumatori. Per l'economia nel suo

TAB. 1.1. Effetti di variazioni dei prezzi del surplus dei consumatori e produttori

|                       | Surplus netto |              |                     |
|-----------------------|---------------|--------------|---------------------|
|                       | Consumatori   | Produttori   | Totale              |
| 1 Iniziale            | $A + B + C$   | $D + E$      | $A + B + C + D + E$ |
| 2 Aumento a $P_1$     | $A$           | $B + D$      | $A + B + D$         |
| 3 = 2 - 1 Variazione  | $-B - C$      | $+B - E$     | $-C - E$            |
| 4 Iniziale            | $+a$          | $+b + c$     | $+a + b + c$        |
| 5 Diminuzione a $P_2$ | $+a + b + d$  | $c - d - e$  | $+a + b + c - e$    |
| 6 = 5 - 4 Variazione  | $+b + d$      | $-b - d - e$ | $-e$                |

complesso vi è stata una perdita di benessere misurabile dall'area  $(C + E)$ . Essa può essere scomposta in due parti: una parte riguardante i consumatori, pari all'area  $C$ , dovuta al fatto che al prezzo più elevato alcuni consumatori decideranno di non domandare e quindi consumare il bene; una parte riguardante i produttori, per i quali la riduzione della quantità domandata da  $Q'$  a  $Q_1$  comporta una riduzione delle rendite. Il surplus netto complessivo restante è ora ripartito in modo più sfavorevole ai consumatori: l'area  $A$  va ai consumatori e l'area  $B + D$  ai produttori.

Il benessere complessivo dell'economia si ridurrebbe anche nel caso in cui nel mercato si fissasse un prezzo inferiore a quello di concorrenza. Si consideri ora la figura 1.2b e si immagini che venga fissato il prezzo  $P_2 < P'$ . La domanda sarà pari alla quantità  $Q_2$ , maggiore di  $Q'$ . Il surplus netto dei consumatori diviene ora pari alle aree, descritte con lettere minuscole:  $a + b + d$ . Ora i produttori sono costretti a produrre la quantità  $Q_2$  con costi pari a  $d + e + f + g$ , ottenendo ricavi pari a  $c + f + g$ . Il loro surplus netto è quindi pari a  $c - d - e$ . La riduzione del surplus netto complessivo (come si può vedere anche nella tabella) è pari ad  $e$ . Il surplus dei consumatori è aumentato, a scapito dei produttori, in misura pari a  $b + d$ , ma tale vantaggio non è compensato dalla perdita subita da questi ultimi.

Abbiamo così dimostrato che il prezzo di equilibrio di un mercato di concorrenza perfetta (prezzo = costo marginale) è coerente con la massimizzazione del benessere dell'economia, misurato dalla somma dei surplus dei consumatori e dei produttori.

Nei prossimi capitoli questo tipo di strumentazione sarà spesso utilizzato per valutare gli effetti di politiche fiscali. Ciò ci consentirà di trarre conclusioni rilevanti per la politica economica, anche senza ricorrere alla complessa, ma certo più rigorosa, analisi di equilibrio generale.

### 1.5. La funzione del benessere sociale

Sviluppiamo ora qualche ulteriore osservazione sulle caratteristiche della funzione del benessere sociale, al fine di mostrare in che senso essa dipenda da giudizi di valore relativi alla distribuzione del benessere tra i soggetti di una società. Si è detto che la funzione di benessere sociale è uno strumento che consente di ordinare in termini di benessere diversi possibili stati sociali, ovvero diverse configurazioni di produzione, scambio e distribuzione dei beni



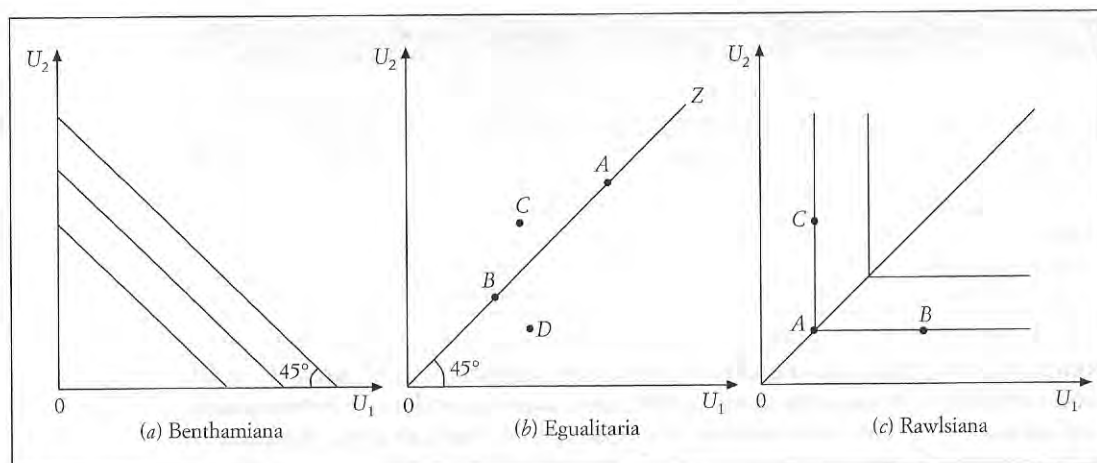


FIGURA 1.3. Funzioni del benessere sociale.

tra gli individui. Esistono diverse possibilità di esprimere questa funzione. Qui ne considereremo solo alcune, a titolo esemplificativo: la benthamiana, la egualitaria e la rawlsiana.

La *funzione benthamiana* si richiama al filosofo inglese J. Bentham, vissuto a cavallo tra il '700 e l'800, fondatore dell'utilitarismo, una concezione filosofica che, seppure aggiornata, continua ad essere il fondamento dell'Economia del benessere moderna. Afferma Bentham, nell'opera *Fragment on Government* del 1776: «L'assioma fondamentale è che [...] la misura di ciò che è giusto o sbagliato è la massima felicità del più grande numero di persone». È quindi naturale che una funzione del benessere ispirata a tale studioso venga espressa analiticamente nel seguente modo:

$$W = U_1 + U_2$$

Il benessere è appunto dato dalla *somma delle utilità* degli individui. Questa concezione del benessere mira a massimizzare il benessere complessivo, trascurando completamente il modo in cui esso è distribuito tra gli individui. In quest'ottica, il benessere della società aumenta se tutti stanno meglio, ma anche se aumenta solo il benessere di chi è già ricco, purché il suo miglioramento sia maggiore del danno inflitto agli altri membri della società.

Graficamente, nel piano  $(U_1, U_2)$  le curve di indifferenza corrispondenti alla funzione del benessere benthamiana sono rette inclinate negativamente e con una pendenza di  $45^\circ$  (cfr. figura 1.3a): tutte le combinazioni di  $U_1$  e  $U_2$  che stanno su una stessa linea di questo tipo hanno infatti uguale somma.

La *funzione egualitaria* è invece caratterizzata da questa condizione:

$$U_1 = U_2$$

che valorizza l'obiettivo secondo cui ogni individuo della società deve raggiungere lo stesso *livello* di benessere.

Graficamente le curve di indifferenza corrispondenti a questa funzione del benessere possono essere rappresentate da curve di indifferenza pun-

tiformi, situate sulla semiretta che interseca il piano  $(U_1, U_2)$  a  $45^\circ$ . Per un egualitario (cfr. figura 1.3b) sono interessanti solo i punti sulla semiretta  $OZ$ . In particolare il punto  $A$  è preferibile a  $B$ , che è preferito ai punti  $C$  e  $D$ . Si noti che secondo questa funzione del benessere sociale risultano preferiti punti, come  $B$ , rispetto ad altri, il punto  $C$ , in cui ciascuno degli individui sta meglio. Tale funzione poi non discrimina tra i punti, ad esempio  $C$  e  $D$ , che non appartengono alla semiretta  $OZ$ .

La *funzione rawlsiana* può essere interpretata come una specificazione particolare di una funzione egualitaria ed è altresì coerente con il cd. principio del *max-min*. Secondo questa funzione, il benessere aumenta se viene migliorata (*max*) la posizione del soggetto che sta peggio (*min*) nella società. Analiticamente:

$$W = \min_b (U_b) \quad \text{con } b = 1, 2$$

Graficamente essa può essere rappresentata da curve di indifferenza ad angolo retto, il cui angolo si situa sulla semiretta che interseca il piano a  $45^\circ$  (cfr. figura 1.3c).

Perché la mappa di curve di indifferenza che rappresenta la funzione del benessere sociale rawlsiana assume questa forma ad  $L$ ? Si consideri una situazione iniziale in cui le risorse dell'economia siano allocate tra gli individui 1 e 2 in modo perfettamente egualitario con il risultato che 1 e 2 godono del medesimo livello di utilità. Tale situazione sarà per esempio rappresentata da un punto  $A$  sulla bisettrice a  $45^\circ$ , un punto della funzione del benessere egualitaria. Supponiamo ora che aumenti esogenamente la quantità di risorse a disposizione dell'economia (cade la manna dal cielo) e queste risorse aggiuntive vengano attribuite interamente all'individuo 1 (la manna cade tutta nelle mani di 1). La nuova allocazione sarà allora rappresentata da un punto  $B$  con identica ordinata rispetto ad  $A$  ma con ascissa maggiore. Dato che in corrispondenza al punto  $B$  l'individuo 2 gode di un'utilità minore dell'individuo 1, il benessere sociale, secondo Rawls, sarà ora uguale all'utilità dell'individuo 2, che è rimasta immutata rispetto alla situazione iniziale  $A$ . Di conseguenza le allocazioni  $A$  e  $B$  corrisponderanno ad un identico livello di benessere sociale e quindi saranno punti della medesima curva di indifferenza sociale. Analogo discorso può essere ripetuto nel caso in cui l'unico beneficiario dell'aumento di utilità sia l'individuo 2 (nuova allocazione rappresentata dal punto  $C$ ).

John Rawls, da cui questa funzione trae il suo nome, è stato uno dei più grandi filosofi morali e politici del XX secolo, la cui opera, *Una teoria della giustizia* del 1971, ha suscitato un grande interesse nel campo della filosofia morale e anche fra gli economisti. La sua costruzione teorica non è tuttavia bene espressa dal solo principio del *max-min* citato e merita quindi qualche ulteriore considerazione. Essa infatti fornisce il fondamento filosofico di una teoria più generale della giustizia sociale, all'interno di una teoria del contratto sociale che solo parzialmente può essere considerata compatibile con l'impostazione welfaristica di cui ci stiamo occupando. La sua posizione filosofica si contrappone in modo molto netto al pensiero liberale classico, pur accettando in modo pieno punti di vista democratici e pluralisti. Per avere

un'idea della distanza tra queste concezioni, basti ricordare che secondo uno dei più grandi rappresentanti del pensiero liberale, F. von Hayek, il termine di giustizia sociale dovrebbe essere eliminato dal lessico della teoria sociale, trattandosi semplicemente di una superstizione di tipo religioso. Ben diverso è il punto di vista di Rawls: per questi giustizia significa equità. E per fare comprendere come essa possa essere il fondamento del vivere sociale, Rawls, così come avevano fatto i grandi filosofi contrattualisti del '600 e del '700, Hobbes, Locke e Rousseau, utilizza la metafora del *contratto sociale*. Egli si domanda quali sarebbero i principi che ipotetici individui che intendano costruire una società nella «posizione originale» sarebbero disposti ad accettare come base del loro vivere civile. Perché l'espressione dei valori avvenga in modo corretto, Rawls osserva che dovremmo immaginare una situazione in cui le opinioni degli individui sulla giustizia non siano influenzate dalla loro posizione concreta nella scala sociale. Per quanto mi sforzi di essere oggettivo, le mie opinioni sul diritto di proprietà possono infatti essere ben diverse a seconda che io non possieda nulla o disponga di un patrimonio ingente. Il contratto sociale può quindi essere correttamente definito solo se immaginiamo che gli individui nella posizione originaria non sappiano quale sarà il ruolo che a loro sarà riservato dalla sorte nella società che deve essere costruita: essi devono, cioè, ragionare «dietro il velo dell'ignoranza». Utilizzando questo esperimento concettuale, e supponendo inoltre che gli individui abbiano fondamentalmente un atteggiamento di avversione per il rischio, egli argomenta e deduce i seguenti principi:

1. *principio di libertà*: ogni persona ha un identico diritto alla massima estensione della sua libertà compatibilmente con l'analoga libertà degli altri individui;

2. *principio di differenza*: le disuguaglianze sociali devono essere risolte in modo che si realizzi il massimo beneficio per il meno avvantaggiato e si applichino a situazioni aperte a tutti gli individui in condizioni di eguali opportunità.

Si noterà che il primo principio è coerente con principi etici come quello cristiano (Ama il prossimo tuo come te stesso) o kantiano (Non fare agli altri quello che non vorresti fosse fatto a te). È in questo principio che Rawls trova il fondamento alla garanzia per ogni membro della società di un insieme fondamentale di *diritti primari* (alla vita, alla libertà, al rispetto da parte degli altri, ad un minimo di risorse, ecc.). Un principio il cui soddisfacimento Rawls considera prioritario in una società che disponga di sufficienti risorse e che, come vedremo nel capitolo ottavo, presenta forti affinità con il cd. egualitarismo specifico.

Il secondo, che dipende crucialmente dall'idea del velo dell'ignoranza e dell'avversione per il rischio, dà invece fondamento ad un principio distributivo. Questo secondo principio è passibile, se pure con qualche forzatura e semplificazione, di un'interpretazione nell'ambito delle teorie welfaristiche ed ha sollecitato in modo particolare l'interesse degli economisti.

La figura 1.4 mostra le relazioni tra le due diverse definizioni di benessere e il criterio paretiano. Supponiamo di partire dal punto A, che non si trova sulla grande frontiera dell'utilità e che caratterizza una situazione in cui  $U_1$  si trova in una posizione di maggiore vantaggio rispetto a  $U_2$ . Si vogliano defini-

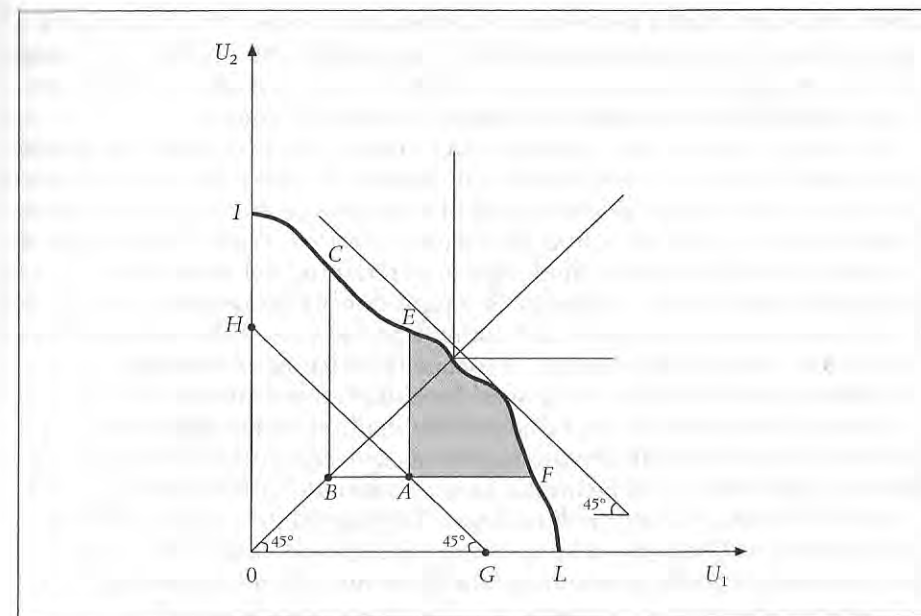


FIGURA 1.4. La grande frontiera dell'utilità.

re gli spazi possibili di miglioramento del benessere alla luce del principio di Pareto, della funzione benthamiana e della funzione Rawls-egualitaria. AEF definisce l'insieme dei punti di miglioramento paretiano. L'area BCF, con esclusione dei punti sui segmenti CB e BF, indica i punti di possibile miglioramento *à la* Rawls e l'area HGLI, con esclusione dei punti sul segmento HG, indica i punti di miglioramento benthamiano. Come si nota, sia la funzione benthamiana sia la rawlsiana considerano come miglioramenti del benessere sociale punti che non sono anche Pareto superiori (le due funzioni infatti ammettono confronti di utilità che Pareto esclude). Non tutti i punti Pareto superiori sono preferibili secondo Rawls, che non assegna valore positivo al miglioramento del ricco (si considerino ad esempio i punti sul segmento AF: non rappresentano un miglioramento per la rawlsiana, pur costituendo miglioramenti paretiani). Tutti i punti di miglioramento paretiano sono invece anche punti di miglioramento benthamiano.

#### 1.6. La valutazione delle riforme in Second Best e il trade-off Efficienza-Equità

Gli strumenti analitici sin qui illustrati potranno apparire molto astratti. Essi sono però utili per identificare i conflitti che nascono quando si discute di riforme economiche e di obiettivi degli interventi pubblici.

Nei dibattiti di politica economica le alternative vengono talora presentate come *trade-off* tra efficienza ed equità, proposti, ad esempio, in questa forma: secondo alcuni, i «riformisti», la proposta  $x$  avrebbe l'effetto di aumentare l'efficienza del sistema economico; la sua realizzazione ha però anche



l'effetto di peggiorare il benessere di un qualche gruppo della collettività, i «conservatori», che tenderanno quindi ad opporsi al mutamento. La ricerca dell'efficienza parrebbe entrare in conflitto con il mantenimento della preesistente distribuzione, creando un costo in termini di equità.

In realtà, i concetti qui studiati ci impongono un'analisi più sottile e ci consentono di comprendere come la valutazione dei pro e dei contro di possibili riforme sia un'operazione molto più complessa. Al generico concetto di miglioramento dell'efficienza dobbiamo infatti sforzarci di sostituire il concetto di «miglioramento di efficienza paretiano», nel senso definito nel precedente paragrafo; al concetto di equità dobbiamo sostituire quello di «livello del benessere sociale», che, come abbiamo visto, è diverso a seconda delle diverse funzioni del benessere sociale (benthamiana, rawlsiana, ecc.).

L'interpretazione della realtà sarebbe abbastanza semplice se il mondo funzionasse come prescrivono i due teoremi dell'Economia del benessere e il libero funzionamento del mercato concorrenziale consentisse di raggiungere un punto sulla frontiera delle utilità. In questo mondo di First Best, al mercato sarebbe affidato il compito di realizzare l'efficienza paretiana, e allo Stato, eventualmente utilizzando *lump sum taxes*, il compito di definire il massimo benessere sociale coerente con una data funzione di benessere sociale. Ma la realtà in cui viviamo non è questa, per due ragioni fondamentali.

La prima è connessa all'impossibilità, già accennata nel paragrafo 1.3, di utilizzare il secondo teorema dell'Economia del benessere, a causa dell'impraticabilità delle *lump sum taxes*. In questo contesto, l'intervento pubblico, nel perseguire obiettivi di redistribuzione e/o stabilizzazione, non può evitare di fare uso di strumenti distorsivi (il «secchio bucato» di Okun). La ricerca di una distribuzione migliore genera una perdita di efficienza paretiana.

La seconda circostanza è connessa all'impossibilità di conformare la realtà alla concorrenza perfetta in condizioni di pieno utilizzo delle risorse produttive disponibili. Una delle più patenti violazioni delle condizioni di First Best, in un'economia di mercato, è costituita dalla presenza della non piena occupazione. Il fatto che non tutti coloro che desiderano lavorare abbiano l'opportunità di farlo è il più evidente segno che il sistema economico non sta sfruttando al massimo le proprie possibilità di generare benessere. Per questa e per molte altre imperfezioni dei mercati, lo *status quo* del mondo in cui viviamo è necessariamente costituito da una posizione di Second Best. Anche in questo contesto, la ricerca di una maggiore efficienza paretiana può entrare in conflitto con l'equità, così come questa risulta descritta dalla *funzione del benessere sociale*.

### 1.7. Limiti dell'Economia del benessere

Il giudizio positivo sull'Economia del benessere come utile punto di riferimento per inquadrare la valutazione delle politiche economiche, al cui interno un ruolo di primo piano hanno le politiche delle entrate e delle spese pubbliche, va temperato da molte cautele. Alcune sono già state ricordate. In questo paragrafo si fa cenno ad altre due, di carattere più generale.

Le critiche più importanti mettono in discussione le due ipotesi cardine dell'Economia del benessere: l'individualismo, in base al quale il bene è sempre riconducibile solo alla valutazione dei singoli, e l'utilitarismo, secondo cui il benessere è definito solo in termini di soddisfacimento di bisogni.

#### ► I beni meritori

In talune circostanze l'adozione di un punto di vista individualistico non consente di fornire una spiegazione di forme di intervento pubblico che pure vediamo correntemente realizzate nelle moderne società democratiche, e che sembrano negare il presupposto secondo cui l'individuo è il migliore giudice del proprio benessere. Come si spiegano le norme in materia di uso di droghe, di consumo degli alcolici da parte dei minori, di adozione delle cinture di sicurezza? O i casi in cui lo Stato impone obbligatoriamente il consumo di certi servizi, come ad esempio nella sanità e nell'istruzione primaria? Devono essere viste come deviazioni dai dettami della teoria, che sarebbe auspicabile fossero rimosse, o tali comportamenti possono trovare giustificazione all'interno di una teoria comprensiva del ruolo dello Stato?

Per tenere conto di questi aspetti, Musgrave ha introdotto nella letteratura finanziaria la categoria dei *merit goods*, tradotta in italiano con beni di merito o beni meritori.

I *merit goods* sono beni che non presentano problemi di «non rivalità» o di «non esclusione», nozioni che, come vedremo tra poco, sono alla base della definizione di un bene pubblico; sono beni destinati a soddisfare bisogni trasversali alla tradizionale distinzione tra beni privati e beni pubblici e caratterizzati da un problema di interferenza nelle preferenze dei consumatori. In questi casi lo Stato sovrappone il proprio punto di vista a quello individuale imponendo certi comportamenti che possono essere interpretati come di tipo paternalistico, ispirati cioè dal desiderio di fare il bene del destinatario anche se questo non ne è persuaso.

L'introduzione di tale principio può creare un'incrinatura nell'impianto teorico welfaristico con conseguenze molto rilevanti. Una volta che si ammetta questo principio si può essere indotti a ritenere che lo Stato, come entità distinta dai cittadini, possa avere proprie finalità, propri bisogni, che non necessariamente coincidono con la volontà dei cittadini. Il fatto che gli esempi citati si riferiscano a interventi che molti riterrebbero positivi, rivolti cioè a realizzare il bene, non deve tranquillizzare. Su questa via si potrebbe teorizzare lo stato etico, a cui gli individui devono obbedienza e che storicamente è stato supporto ideologico a dittature.

I pericoli accennati non devono tuttavia scoraggiare troppo. La presenza di interferenze nelle preferenze individuali è un fenomeno fatalmente diffuso nella vita sociale. Va semmai ricordato che queste interferenze si manifestano non solo quando si discute di problemi di intervento pubblico, ma sono anche molto diffuse nei rapporti di mercato. Si pensi, ad esempio, al ruolo che ha la pubblicità nel «forzare» le preferenze individuali.

La giustificazione delle interferenze è quindi un punto delicato a cui molti studiosi, filosofi in particolare, hanno cercato di dare una risposta. Alcuni dei tentativi più interessanti cercano di teorizzare la presenza di sistemi di preferenza multipli. Gli individui, allorché compiono scelte individuali ordinarie (ad esempio che beni acquistare, che lavoro fare, ecc.), agirebbero sulla base di un sistema di preferenze diverso da quello che informa il loro comportamento in sfere di scelta più elevate come quella politica o etica. Questa molteplicità di sistemi di preferenze potrebbe giustificare scelte contraddittorie alla luce di un solo sistema di preferenze.

Il problema si pone non solo per il rapporto tra individuo e Stato, ma anche per altri aspetti della teoria economica di impostazione utilitarista. Chi conosce l'economia politica e sa apprezzare il suo modo di procedere per ragionamenti astratti non di rado resta perplesso di fronte a tentativi di fornire una spiegazione della realtà affidandosi esclusivamente al principio individualistico, che quasi sempre viene descritto in termini egoistici.

#### ► *Equità consequenziale ed equità procedurale*

Il presupposto utilitaristico dell'Economia del benessere è stato oggetto di critiche, anche per le implicazioni che esso ha nella definizione di equità. Quella implicita negli schemi dell'Economia del benessere è del tipo noto come *equità consequenziale*, così chiamata perché ciò che conta, ai fini delle valutazioni del benessere, sono le conseguenze dei risultati prodotti sui soggetti interessati dalle misure introdotte. Non è difficile fare esempi che mostrano come questa impostazione possa dare luogo a paradossi. Dare del denaro ad un tossicodipendente può consentire la soddisfazione di un bisogno individuale urgentissimo, soggettivamente molto più intenso di quello che con le stesse risorse si potrebbe arrecare ad un lavoratore povero che ha molti figli da mantenere. Sareste disposti a ritenere preferibile il trasferimento al primo piuttosto che al secondo?

I dubbi sulla validità dell'impianto dell'Economia del benessere, che l'esempio fatto istilla, dipendono dal fatto che la nozione di equità consequenziale che ne sta alla base non appare pienamente soddisfacente in ogni circostanza. Accanto ad essa, si può concepire anche una nozione di equità che non è tanto interessata ai *risultati*, ma al fatto che la vita sociale sia definita da *regole* eque. Si usa in questo caso il termine di *equità procedurale*. Secondo questa nozione è giusta la società in cui tutti i membri hanno pari opportunità di realizzare un progetto di vita soddisfacente. L'equità si realizza non tanto cercando di garantire a tutti una vita felice, ma mettendo tutti nelle condizioni di disporre degli strumenti per poter realizzare il proprio progetto. È equo, ad esempio, offrire a tutti l'opportunità di andare a scuola, non di garantire lo stesso risultato (il diploma). L'enfasi è sul *processo* che porta ad un determinato risultato e all'interno di questo processo si attribuisce solitamente molto valore ai diritti fondamentali (alla vita, alla libertà di espressione, al diritto al lavoro, ecc.), alle regole del processo piuttosto che all'esito. Secondo questa nozione se un soggetto, per ragioni che possono dipendere dal suo impegno, ma anche

dal caso, non riesce nel suo intento, non dovrebbe essere meritevole di aiuto da parte della società.

L'equità procedurale sottolinea molto la responsabilità dell'individuo nella società. Esistono diritti fondamentali, di cui tutti gli esseri umani come membri di una società devono godere, ma essi sono fondati sull'esistenza di una relazione di reciprocità: così, ad esempio, il dovere al lavoro è sottolineato come elemento di scambio rispetto al diritto di essere tutelato in caso di bisogno.

Della nozione di equità procedurale esistono molteplici varianti, che hanno implicazioni abbastanza differenti in termini di suggerimenti per l'azione politica. Una versione estrema, suggerita dal filosofo Nozick, sostiene che lo scopo essenziale di un'attività pubblica è solo quello di garantire l'invulnerabilità dei diritti individuali, fra cui vi è *in primis* quello della proprietà. Ogni altro intervento sarebbe da considerarsi proceduralmente scorretto. Da tale visione, come è facile intuire, discende un ruolo dello Stato minimale, che deve concentrarsi solo su limitatissimi compiti fra cui *non* vi è quello di redistribuire le risorse.

Una versione più tradizionale, che risale al pensiero liberale, è quella che sottolinea l'esigenza di garantire eguaglianza dei punti di partenza. Einaudi, che abbiamo citato all'inizio di questo capitolo, era un sostenitore di questo principio. In questa versione il compito dello Stato è di garantire condizioni di parità sufficienti (sotto forma di possibilità di istruirsi, non discriminazione rispetto al censo o allo stato sociale nelle professioni e nella vita politica, ecc.), lasciando però che ognuno faccia la propria corsa nella vita. Una visione di questo tipo, in cui è sottolineata l'esigenza di mantenere vivo il senso di responsabilità individuale, può suggerire un intervento ridotto, ma non minimo dello Stato.

La valorizzazione dell'equità in senso procedurale ha influenzato molto il dibattito sui sistemi di welfare in Europa ed ha un rilievo notevole nell'ambito del modello di welfare noto come *Flexi-security* che riprenderemo nel capitolo settimo.

#### ► *L'approccio dello sviluppo umano*

Esistono poi versioni più sofisticate e articolate di equità non welfare, che, partendo dalla nozione di diritti fondamentali di cittadinanza che devono essere garantiti, hanno messo a fuoco concetti più appropriati dell'utilità individuale, funzione della quantità di beni di cui un soggetto può disporre, per descrivere il benessere, o meglio il *well-being* individuale, introducendo elementi non solo di tipo quantitativo, ma anche e soprattutto di tipo qualitativo. Secondo il filosofo ed economista A. Sen, il benessere (che, per distinguerlo, anche lessicalmente, dalla nozione tradizionale è indicato con il sostantivo *well-being*, anziché *welfare*) può essere definito solo in termini di allargamento degli spazi della libertà individuale, nel senso della possibilità di autodeterminazione e possibilità di scelta di un proprio progetto di vita. Danno contenuto a questa impostazione due concetti fondamentali: i *funzionamenti* e le *capacità*. I primi descrivono il



modo in cui i bisogni individuali, materiali e non, vengono soddisfatti. Ad essi appartengono quelli più elementari, come essere nutriti, vestiti, avere la possibilità di muoversi e comunicare, ma possono riguardare anche aspetti più complessi, come poter partecipare attivamente alla vita pubblica, sentirsi integrati in un dato ambiente sociale. Questi elementi costitutivi del benessere, tipicamente multidimensionali, potranno diventare effettivi se si traducono in insiemi di capacità, che rappresentano l'insieme dei funzionamenti potenzialmente attivabili da parte di un individuo e forniscono una misura, non solo quantitativa, ma soprattutto qualitativa, del grado della realizzazione della libertà. Lo «sviluppo umano» – un modo in cui questa impostazione viene di frequente indicata – è quindi rappresentato dall'estensione delle capacità di funzionamento di un soggetto. È dallo sviluppo di idee di questo tipo, oggetto di elaborazione da parte della teoria dello sviluppo umano, che ci si attendono contributi in grado di superare la visione alquanto asfittica, sotto il profilo dei valori, dell'Economia del benessere tradizionale.

Entrambe le nozioni di equità presentate (conseguenziale e procedurale), nelle molteplici varianti a cui abbiamo fatto cenno, presentano aspetti positivi e negativi, e nella discussione dei problemi economici sociali si fa spesso uso dell'una o dell'altra nozione di equità. È molto importante essere consapevoli che esse possono portare a prescrizioni politiche diverse e sapere riconoscere nelle discussioni di politica la presenza di elementi dell'una o dell'altra.

### ► L'economia della felicità

Economisti e filosofi del '700 come A. Smith e J. Bentham consideravano il raggiungimento della felicità come uno degli obiettivi fondamentali degli uomini e delle società. Per lungo tempo, nel corso dell'800, gli economisti erano convinti che fosse possibile misurare e confrontare il benessere delle persone. In seguito, la teoria economica neoclassica, a partire dal contributo di Pareto, ha respinto l'idea che i livelli di benessere individuali siano misurabili in senso cardinale e comparabili. In realtà, confronti tra livelli di benessere sono sempre stati fatti, utilizzando il reddito o il consumo come approssimazioni del tenore di vita delle persone.

Un recente filone di ricerca, che ha preso il nome di *economia della felicità*, non solo ammette esplicitamente la confrontabilità dei livelli di benessere, ma quantifica questi ultimi non utilizzando il reddito, bensì indicatori soggettivi che misurano quanto le persone si sentono soddisfatte della propria vita. Il modo per ottenere informazioni di questo tipo è in genere molto semplice: nel caso più tipico, si tratta delle risposte che, in indagini campionarie, gli intervistati danno a domande del tipo «in una scala da 0 a 10, quanto sei soddisfatto della tua vita?». Questo approccio si basa quindi sulla auto-valutazione che le persone forniscono del proprio livello di benessere. La teoria neoclassica tradizionale, invece, ammette la possibilità di osservare solo il comportamento degli agenti, non le loro dichiarazioni dirette circa l'utilità goduta.

Rispetto al concetto rawlsiano di bene sociale primario (che comprende il reddito e la ricchezza), l'economia della felicità non si ferma al reddito,

perché lo considera, come nell'approccio delle capacità, solo strumentale al raggiungimento di obiettivi più importanti, cioè una delle tante possibili determinanti del livello di (in)felicità. Anche se l'economia della felicità e l'approccio delle capacità condividono la valutazione del reddito come strumento e non come obiettivo, Sen ha criticato l'approccio dell'economia della felicità, sostenendo che se una persona vive da lungo tempo in condizioni degradate, può assuefarsi a questo stato e considerarsi soddisfatta della propria vita, anche se la sua sfera di capacità è estremamente limitata e insufficiente per vivere una vita degna. A sua volta, l'economia della felicità rimprovera alla teoria delle capacità un eccesso di paternalismo, che può sfociare nella formulazione di una «lista» di funzionamenti che tutti dovrebbero possedere.

L'economia della felicità utilizza metodi di analisi tipici sia dell'economia sia della psicologia. Si cerca di comprendere quali siano le variabili che possono influenzare la soddisfazione non solo per la vita in generale, ma anche per dimensioni più particolari, come il reddito, le condizioni di salute, le relazioni sociali o familiari.

In un'analisi *cross-section* relativa ad un dato paese, è consueto riscontrare una relazione crescente tra reddito e felicità: i ricchi sono in genere più felici dei poveri. Questa relazione non è però lineare: a bassi livelli di reddito, la felicità cresce rapidamente con esso. Partendo da redditi più elevati, invece, la felicità aumenta lentamente, il che sta ad indicare che l'utilità marginale del reddito è decrescente. Molti altri fattori hanno un effetto significativo sulla felicità: le condizioni di lavoro, i tratti della personalità, la salute, le relazioni familiari.

Uno dei primi studiosi a gettare le basi di questo filone di ricerca, R. Easterlin, rilevò negli anni '70 un interessante «paradosso». Come si è appena visto, il reddito influenza positivamente il livello di felicità individuale. Nel corso del tempo, però, al crescere del livello di reddito di un paese non si accompagna in genere un altrettanto significativo incremento del livello medio di felicità. Nel lungo periodo cresce quindi il reddito *pro capite*, non la felicità media.

Una possibile spiegazione del paradosso di Easterlin consiste nell'idea che gli individui effettuino un confronto tra il proprio reddito e quello degli altri: il loro benessere aumenta solo se c'è un miglioramento relativo rispetto al resto della popolazione. Una semplice crescita del reddito medio nazionale, quindi, non sarebbe sufficiente per produrre un incremento del livello medio di felicità.

Una seconda spiegazione, non necessariamente alternativa, assume che gli uomini posseggano un meccanismo di adattamento delle proprie aspirazioni rispetto a variazioni dell'ambiente in cui vivono: un incremento di reddito, come altri eventi positivi, produce una maggiore felicità solo per un certo intervallo di tempo, per poi tornare (più o meno) allo stesso livello di prima. In senso opposto giocherebbero eventi negativi, come una separazione o una seria malattia. Questo meccanismo di adattamento è stato a lungo studiato in psicologia: la felicità dipende dalla presenza di uno scarto tra aspirazioni e realizzazioni. Dopo che un obiettivo è stato raggiunto, non rimaniamo a lungo molto soddisfatti di esso, perché rivediamo le nostre aspirazioni verso l'alto.

L'adattamento edonico consiste in un qualsiasi processo che riduca l'intensità del (dis)piacere provato a seguito di un evento. Si tratta di un meccanismo naturale che sembra tutti gli uomini possiedano, almeno in parte. La sua utilità può consistere nel fornire alle persone un sistema automatico di difesa da eventuali conseguenze sconvolgenti che sulla psiche umana potrebbero avere fenomeni traumatici, non necessariamente negativi.

In sintesi, la felicità sembra aumentare meno che proporzionalmente rispetto al reddito, o potrebbe anche non aumentare affatto, per due motivi principali: l'influenza del reddito relativo, per cui sono più felice se il mio reddito non aumenta solo in assoluto, ma anche relativamente agli altri; l'effetto di adattamento, per cui al crescere del reddito aumenta anche il livello di quest'ultimo che considero «adeguato», «necessario» o, appunto, soddisfacente.

## 2. Economia con beni pubblici e meccanismi di decisione politica

Il sistema economico sinora considerato prevede l'esistenza di un solo tipo di beni: i beni privati, i beni cioè che vengono solitamente prodotti e scambiati nel mercato, come alimentari, abbigliamento, servizi abitativi, ecc. L'osservazione della società ci mostra però anche la presenza di altri tipi di beni e servizi che vengono prodotti e fruiti, ad esempio la difesa, l'amministrazione della giustizia, i trasporti pubblici, i servizi stradali e autostradali, i servizi di nettezza urbana, i servizi di un parco pubblico, i servizi di comunicazione di massa, radio e Tv, ecc. Quasi sempre i beni e servizi di questo tipo non sono forniti da imprese private, ma sono, seppure in misure e modi diversi, gestiti dall'autorità pubblica, dallo Stato.

L'economia pubblica nasce dall'osservazione di questa realtà e le prime domande che ci si è posti sono le seguenti: quali sono le caratteristiche comuni a questi tipi di beni e servizi? Esistono differenze tra questi beni e servizi e quelli privati? Nell'ambito di una teoria positiva, può la loro spiegazione teorica – vale a dire le ragioni che spiegano la misura in cui tali servizi sono prodotti e la definizione dei meccanismi con cui viene finanziato il loro costo – essere affrontata applicando ad essi le stesse categorie interpretative che gli economisti utilizzano per i beni privati?

Un passo avanti significativo nella costruzione di una teoria dei servizi pubblici viene realizzato allorché gli studiosi riescono a caratterizzare con precisione i connotati del bene pubblico rispetto a quello privato. Sin dagli albori della teoria finanziaria, era stato individuato il concetto di *indivisibilità dei vantaggi* connessi ai servizi pubblici. La natura dei servizi pubblici è tale che non è possibile individuare in quale misura il beneficio del servizio affluisce ad un individuo o ad un altro. Il vantaggio che deriva dalla difesa, ad esempio, non è divisibile. Da tale constatazione si traeva la conclusione che in questi casi non era possibile individuare un prezzo come controprestazione del servizio. La genesi logica dell'imposta veniva allora spiegata partendo dalla seguente considerazione: l'indivisibilità richiede l'introduzione di una forma di pagamento non correlata al beneficio e coercitiva.

Successivamente però la teoria ha messo a fuoco con maggiore precisione i caratteri essenziali dei servizi o beni pubblici rispetto a quelli privati, individuando due aspetti molto importanti:

1. non rivalità;
2. non escludibilità.

Un bene o un servizio è *rivale* quando il consumo da parte di un soggetto non può essere condiviso anche da un altro soggetto. Il classico esempio: se Tizio mangia una mela non può mangiarla Caio. Un bene è non rivale quando il consumo da parte di Tizio non impedisce anche a Caio di godere del bene. Se Tizio ascolta un brano di musica riprodotto da un CD, anche Caio può ascoltarlo senza creare alcun problema al godimento di Tizio. Lo stesso concetto di non rivalità viene illustrato ricorrendo all'idea di *offerta congiunta*: l'offerta di un'unità di bene non rivale risulta disponibile in modo congiunto a tutti i membri della collettività. Si può esprimere il concetto anche in questo modo: il costo marginale di offrire una data quantità di bene pubblico ad un soggetto in più è nullo. Ciò naturalmente non significa che sia nullo il costo marginale di produrre un'unità di bene pubblico in più.

Un bene è *escludibile* se può essere regolamentato il suo consumo, vale a dire se è possibile consentirlo ad un soggetto ma impedirlo ad un altro.

Si noti che il concetto di bene pubblico non ha nulla a che fare, di per sé, con il fatto che esso venga prodotto e offerto dallo Stato; i suoi connotati riguardano esclusivamente due sue caratteristiche intrinseche, potremmo dire tecnologiche. Un esempio classico è l'ordinamento giudiziario: esso produce un servizio, il rispetto dei diritti, che in una vita associata ha la stessa importanza dell'aria per la vita organica. La giustizia è forse l'unico esempio di bene pubblico che storicamente sia sempre stato fornito dall'autorità politica: sarebbe invero assai poco tranquillizzante vivere in una società in cui i giudici sono pagati da imprese private! Altri esempi sono la difesa e i servizi di illuminazione pubblica.

Oltre al caso polare di bene pubblico, esistono però anche categorie di altri tipi di beni che possiedono solo una delle due caratteristiche citate o che le possiedono in misura limitata. Si tratta dei *beni misti*, di fatto i più frequenti nella realtà concreta, che meriterebbero di essere analizzati in dettaglio. Molto importanti, come si vedrà nel capitolo ottavo, sono taluni servizi rivali ed escludibili, ma caratterizzati dalla presenza di esternalità positive. Qui ci limiteremo a richiamare una classica quadripartizione (cfr. tabella 1.2).

Un esempio di bene tariffabile è un'autostrada, i cui servizi, nei limiti della non congestione, sono non rivali, ma possono essere escludibili a costi non proibitivi.

TAB. 1.2. Beni privati, beni pubblici e beni misti

|                 | Rivale       | Non rivale       |
|-----------------|--------------|------------------|
| Escludibile     | Bene privato | Bene tariffabile |
| Non escludibile | Bene comune  | Bene pubblico    |



Un esempio di bene comune è una riserva di pesca. I servizi da essa resi sono rivali (ciò che peschi tu non posso prenderlo io), ma pongono spesso problemi di escludibilità.

## 2.1. L'equilibrio concorrenziale con beni privati e beni pubblici

Come funziona un'economia decentrata di tipo concorrenziale in cui esistono anche beni pubblici oltre che beni privati? Le condizioni di equilibrio che ci consentono di determinare i prezzi e le quantità restano quelle già note, cioè eguaglianza dei tassi marginali di sostituzione con il rapporto dei prezzi, ecc.?

La risposta a questa seconda domanda è negativa: le condizioni cambiano e la loro derivazione è stata opera di Samuelson con due saggi magistrali del 1954 e del 1955. I risultati di Samuelson possono essere così sintetizzati.

In un contesto di teoria positiva, si può dimostrare che se gli individui sono disposti a rivelare le proprie preferenze per i beni pubblici, in un mercato concorrenziale è possibile definire un insieme di quantità e di prezzi rispettivamente per i beni pubblici e per i beni privati, purché sia possibile immaginare, per i beni pubblici, l'esistenza di una sorta di mercato individualizzato, che consenta di determinare prezzi personalizzati (pseudo-mercato). La novità per i beni pubblici è che in equilibrio il prezzo di un'unità di essi non è uguale alla valutazione marginale che del bene dà ciascun individuo, ma sarà pari alla *somma* di tali valutazioni marginali. Poiché un'unità di bene pubblico è consumata in eguale misura da tutti gli individui, il costo marginale di un'unità dovrà essere confrontato con la somma delle valutazioni marginali di tutti gli individui. Un modo semplice di comprendere questo aspetto era stato già messo in luce, in un contesto di equilibrio parziale, e non di equilibrio generale come nell'analisi di Samuelson, da Bowen nel 1943: nel caso di beni privati la domanda aggregata di un bene si ricava sommando *orizzontalmente* le domande individuali, mentre la domanda aggregata di un bene pubblico si ottiene sommando *verticalmente* le domande individuali.

Il grafico di Bowen è proposto nella figura 1.5. Nella parte superiore si considera un bene privato. Sono descritte le domande di due soggetti *A* e *B*, che sono aggregate nella terza parte della figura, ove è rappresentata anche la curva dell'offerta. Nel caso di un bene privato l'aggregazione si realizza per somma orizzontale: nel mercato viene bandito un prezzo e i due soggetti indicano quale quantità del bene sono disposti ad acquistare a quel prezzo. Il punto della scheda di domanda aggregata corrispondente al prezzo bandito è dato dalla somma delle domande dei due individui. L'equilibrio si ha in corrispondenza al prezzo  $P^*$  e alla quantità  $Q^*$ , che viene ripartita tra i due soggetti nell'ammontare  $0Q_B$  e  $Q_BQ^* = 0Q_A$ .

Nella parte inferiore è rappresentato il caso di beni pubblici. Si suppone che i due individui esprimano domanda individuale del bene in modo corretto (non c'è *free riding*). L'aggregazione deve però avvenire per somma verticale in relazione alle caratteristiche di non rivalità. L'offerta di una data quantità sarà infatti goduta da entrambi nella stessa misura e si tratta di determinare

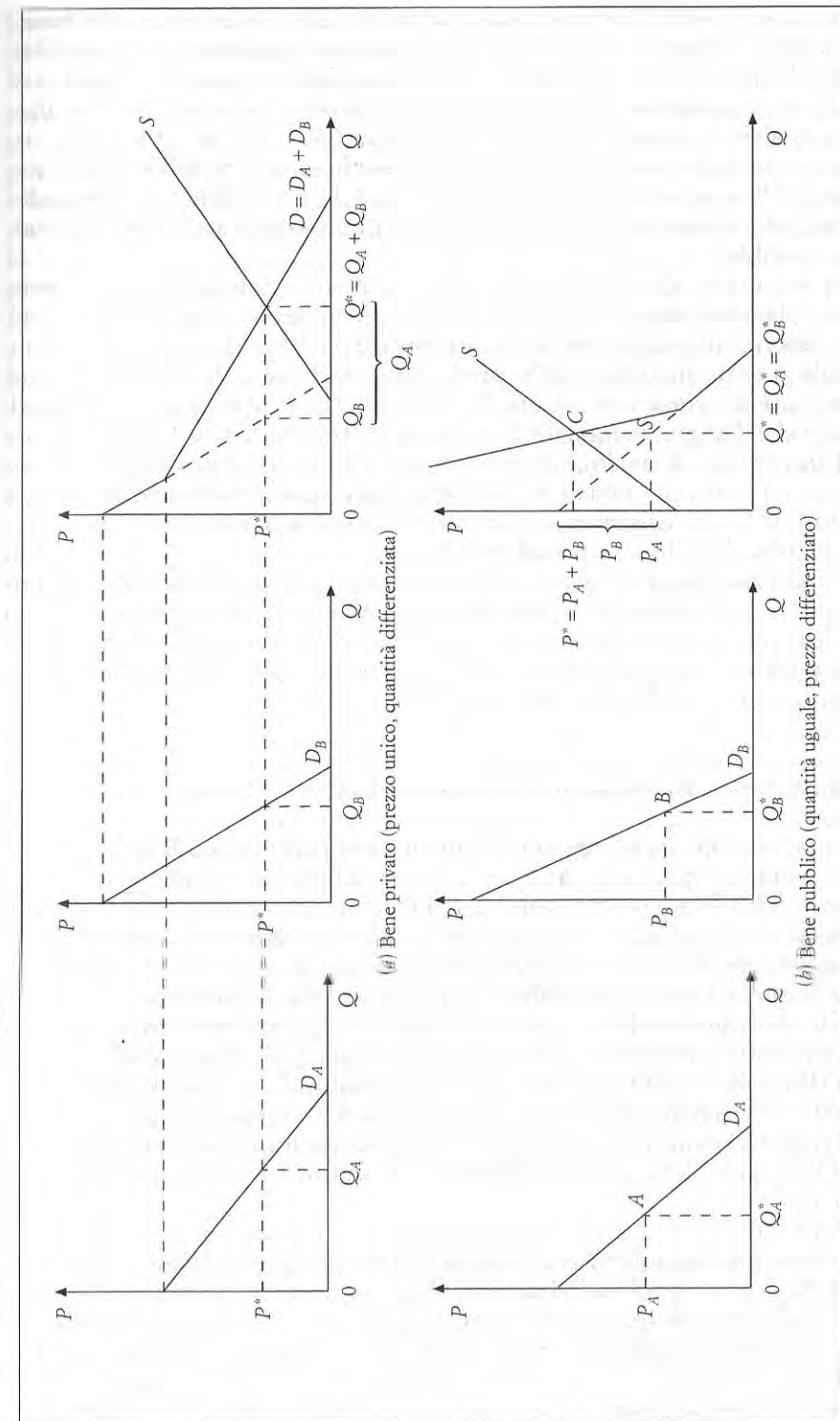


FIGURA 1.5. Aggregazione delle domande di beni privati e pubblici.

il prezzo complessivo che i due consumatori sono disposti a pagare sulla base delle proprie domande individuali. Tale prezzo sarà evidentemente pari alla somma dei prezzi che ciascuno di essi è disponibile a pagare. L'equilibrio si ha in corrispondenza alla quantità  $Q^*$  e ad un prezzo pari a  $P^* = P_A + P_B$ , che va ripartito tra i due individui nella misura  $CS = BQ_B^*$  e  $SQ^* = AQ_A^*$ .

In un contesto normativo, utilizzando cioè lo stesso impianto dell'Economia del benessere sulla cui base sono stati derivati i due fondamentali teoremi, Samuelson ricava le condizioni di ottimo sociale anche in presenza di beni pubblici.

La sua analisi, di cui riferiamo solo il risultato, afferma che, per i beni pubblici, la condizione di ottimo paretiano si realizza quando la somma dei saggi marginali di sostituzione tra bene privato e pubblico di tutti gli individui è uguale al costo marginale della produzione del bene pubblico. (Nel caso dei beni privati, come si ricorderà, la condizione di efficienza nello scambio è invece che il saggio marginale di sostituzione tra i beni  $X_1$  e  $X_2$  dell'individuo 1 sia identico all'analogo saggio marginale di sostituzione dell'individuo 2 e questi a loro volta uguali al costo marginale della produzione del bene privato.) Alla luce della spiegazione fornita dello schema di Bowen tale conclusione non dovrebbe sorprendere il lettore.

Questo risultato è ottenibile se il governante è in grado di conoscere le valutazioni degli individui, relativamente sia ai beni privati sia ai beni pubblici. Samuelson non si preoccupa però di individuare, con riferimento alla teoria positiva, i meccanismi istituzionali che consentano di determinare i prezzi-imposta che gli individui dovrebbero pagare.

## 2.2. Il fallimento del mercato in un'economia con beni pubblici

La possibilità che un regime di concorrenza perfetta con beni pubblici possa vantare le preziose caratteristiche di ottimalità, sancite dal primo teorema dell'Economia del benessere, ha come presupposto che vi siano imprese disposte ad offrire beni pubblici e individui disposti a rivelare le loro preferenze, cioè le domande individuali (lo pseudo-mercato). In un'economia decentrata reale né l'una né l'altra cosa possono essere realizzate.

Gli individui non hanno interesse a rivelare le loro preferenze: tale fenomeno è noto come *free riding*. Dato che un'unità di bene pubblico, una volta disponibile, sarà godibile, senza possibilità di esclusione, da tutti i soggetti del sistema economico, nessuno si farà avanti per pagare. In altre parole, se si chiedesse ad un soggetto quale valutazione marginale attribuisce al bene pubblico, questo sarebbe tentato di fornire valutazioni inferiori alla verità.

Poiché il prezzo che copre il costo, in un'economia con beni pubblici, deve essere pari alla somma delle valutazioni marginali individuali, la sottovalutazione da parte degli individui fa sì che le imprese non possano produrre i beni pubblici nella quantità ottimale, perché produrrebbero in perdita.

Un modo per risolvere il problema del *free riding* potrebbe essere quello di affidare allo Stato l'offerta e la determinazione dei prezzi-imposta. In tal caso si potrebbe raggiungere una situazione di ottimo sociale, ma tale solu-

zione è assai irrealistica, dato che essa presuppone che lo Stato disponga di una quantità di informazioni enorme, analoga a quella necessaria per fare funzionare un'economia pianificata. Se lo Stato dispone solo di una parte di esse, il problema del *free rider* può essere visto come un problema di asimmetria informativa del tipo *adverse selection* (su tale concetto ci soffermeremo più avanti). Lo Stato può essere configurato come il *Principale* e il cittadino come l'*Agente* all'interno di un rapporto contrattuale di delega. Il cittadino dispone di un'informazione (la propria valutazione del bene pubblico) che lo Stato non ha e quindi quest'ultimo si trova in difficoltà a realizzare una situazione di ottimo sociale. (Si noti che lo stesso problema di *adverse selection* è presente in un'economia senza beni pubblici, allorché si cerchi di applicare il secondo teorema dell'Economia del benessere e di determinare le *lump sum taxes* necessarie per passare da un ottimo paretiano ad un altro. Tali imposte infatti, per non essere distorsive, devono essere commisurate, come si è detto, a fattori fuori dal controllo dell'individuo, come le sue abilità naturali, sulle quali il contribuente dispone di un'informazione maggiore rispetto allo Stato che fissa il prelievo.)

Con riferimento al problema del *free rider*, se considerassimo la decisione di rivelare o meno le proprie preferenze come un gioco tra individui, a cui sono associati dei *payoff* in relazione ai benefici che si traggono dal bene e a quanto si è disposti a pagare, tale situazione si configurerebbe come un gioco non cooperativo simile al Dilemma del prigioniero. La soluzione scelta dai partecipanti al gioco è quella di non dichiarare le proprie preferenze, raggiungendo una situazione (non produzione del bene pubblico) che è meno preferita di un'altra strategia che potrebbe essere superiore dal punto di vista paretiano. Non rivelare le preferenze è cioè una strategia dominante per tutti i soggetti del sistema economico e pertanto il bene pubblico in un'economia decentrata non verrà né offerto né domandato. Si ha così un primo caso di fallimento del mercato, che può richiedere l'intervento pubblico.

Si consideri la matrice di *payoff* della tabella 1.3, che si riferisce a un caso in cui un individuo  $A$  è uno dei tanti che possono trarre vantaggio dalla produzione di un bene pubblico. Se il bene è prodotto,  $A$  ottiene un vantaggio di 10 e gli è richiesto di contribuire per 5. Se gli altri contribuiscono, e  $A$  paga, avrà un vantaggio di  $5 = 10 - 5$ ; ma se non paga avrà un vantaggio di 10. Posto che gli altri non contribuiscono, se  $A$  paga subisce una perdita di 5 (non c'è rimborso); mentre se non paga ha un costo pari a 0. Non contribuire costituisce dunque una strategia dominante per  $A$  (cioè ad  $A$  converrà comunque non contribuire, qualunque sia la scelta degli altri).

Se l'individuo  $A$  non fosse solo uno fra i tanti, il problema potrebbe presentarsi in modo diverso. Supponiamo che gli individui siano solo due,  $A$  e

TAB. 1.3. La domanda di un bene pubblico come gioco non cooperativo

|                      | Gli altri contribuiscono<br>(bene prodotto) | Gli altri non contribuiscono<br>(bene non prodotto) |
|----------------------|---------------------------------------------|-----------------------------------------------------|
| $A$ contribuisce     | 5                                           | -5                                                  |
| $A$ non contribuisce | 10                                          | 0                                                   |



TAB. 1.4. *Domanda di beni pubblici e «piccoli numeri»*

|                    | B contribuisce | B non contribuisce |
|--------------------|----------------|--------------------|
| A contribuisce     | 5, 5           | -5, 0              |
| A non contribuisce | 0, -5          | 0, 0               |

B. Una matrice dei *payoff* potrebbe essere quella indicata nella tabella 1.4, ove i due numeri separati da una virgola (ad es. 5, 5) indicano, per ciascun esito possibile, i *payoff* di A e B.

In questo caso, in cui vi sono due soli individui, è ragionevole supporre che se uno dei due soggetti non contribuisce il bene pubblico non verrà prodotto. Se ci domandiamo quali siano le strategie per A e B, in questo caso vediamo che sono possibili due equilibri di Nash. Il primo (5, 5) è ovviamente preferibile al secondo (0, 0), ma in mancanza di ulteriori vincoli, non è detto che i due individui arrivino a quel risultato. È però chiaro che in questa situazione di «piccoli numeri» si apre uno spazio di contrattazione e ciascun individuo acquisisce una maggiore consapevolezza del peso della propria decisione, assente invece nel caso di «grandi numeri».

### 2.3. I meccanismi di decisione politica

Di fronte a queste difficoltà sorge allora il problema di configurare meccanismi istituzionali diversi dal mercato, che abbiamo visto non può funzionare, per indurre i soggetti a rivelare le loro preferenze. Gli economisti pubblici si sono a lungo esercitati nella ricerca di questa ingegneria istituzionale e alcuni risultati sono stati raggiunti. Da tale ricerca non sono tuttavia emersi nuovi modelli istituzionali praticabili su larga scala nelle istituzioni sociali.

#### ► Il voto come alternativa al mercato

Il fallimento del mercato a causa del problema del *free rider* in un'economia con beni pubblici non è dunque facilmente superabile. Resta pur vero che la scelta di quali e quanti beni pubblici produrre deve essere, e di fatto è storicamente, compiuta. L'istituto che nelle nostre società svolge questo ruolo è il meccanismo del voto, attraverso il quale vengono scelti i governanti che sono delegati a stabilire, anche utilizzando poteri coercitivi, quanti beni pubblici produrre e come ripartirne il costo tra gli elettori.

La sostituzione del voto al mercato come meccanismo allocativo nelle decisioni di produzione di beni pubblici rappresenta un passaggio logicamente molto impegnativo, della cui portata teorica bisogna avere consapevolezza. Il meccanismo allocativo del mercato, che abbiamo descritto nei richiami di economia del benessere, assume infatti come punto di partenza le preferenze individuali e le *dotazioni iniziali*. Queste ultime identificano un assetto dato della distribuzione delle risorse sancita dal rispetto dei diritti di proprietà. Anche nel voto hanno rilievo le preferenze individuali dei cittadini. Ciò che

però caratterizza questo meccanismo allocativo è l'irrilevanza della distribuzione delle risorse, se si suppone, come avviene di solito, e come storicamente si è verificato nei moderni sistemi di rappresentanza politica, che ad ogni testa sia associato un voto. Nel mercato l'esito allocativo è influenzato in misura maggiore da chi ha più risorse; nel voto ogni cittadino ha lo stesso potere nel processo di decisione di produzione di beni pubblici. Questa osservazione lascia intuire perché vi siano punti di vista così contrastanti sul peso relativo che nella società devono avere i due meccanismi allocativi e perché gli individui più ricchi solitamente avversino l'estensione del ruolo dello Stato. In prospettiva storica può essere interessante richiamare che il diritto al voto per ogni cittadino è un connotato dei sistemi politici che è stato raggiunto solo di recente. Nel secolo XIX il diritto al voto era infatti riconosciuto solo a chi aveva un certo censo, riproducendo così anche nella sfera pubblica il ruolo che egli aveva nel mercato. Ancora oggi, anche nelle democrazie più evolute, in cui le costituzioni sanciscono il principio del voto universale, la distribuzione delle risorse non è neutrale rispetto al processo decisionale: organizzare una campagna elettorale richiede infatti risorse finanziarie considerevoli e la possibilità di disporre di canali di comunicazione di massa. Le chance dei più ricchi di influenzare l'esito dei meccanismi elettorali restano quindi molto forti anche in sistemi liberali e democratici.

Gli studiosi di economia pubblica hanno cercato di capire quali siano gli effetti dei diversi possibili meccanismi di voto, un terreno ben noto e coltivato dagli studiosi di scienza della politica, e si sono anche domandati quali possano essere preferibili alla luce della teoria dell'Economia del benessere. L'analisi, ancora una volta, è stata condotta ai due livelli, normativo e positivo. Questo importante campo di ricerca, in cui l'economia invade con i suoi strumenti di analisi il campo della teoria politica e della sociologia, viene solitamente indicato, come abbiamo ricordato nel paragrafo introduttivo, come teoria della *Public Choice*.

La *Public Choice* nasce in parte come reazione all'idea che al cattivo funzionamento del mercato possa porre un rimedio efficace l'intervento pubblico. Nel suo versante di teoria positiva, al di là di questa originaria preoccupazione di tipo ideologico, essa pone alla base della sua analisi il rapporto politico nei suoi aspetti di relazione tra cittadino-elettore e governanti e tra governanti e burocrati, utilizzando gli strumenti correnti dell'analisi economica di tipo neoclassico (comportamenti ottimizzanti). Gli studi di questa scuola hanno posto attenzione al funzionamento delle democrazie rappresentative, cercando di spiegare l'esistenza di certi istituti e i limiti del loro funzionamento: procedimenti di votazione a maggioranza, regole costituzionali, formazione di coalizioni tra partiti, studio dei rapporti tra cittadino e politico, tra politico e burocrate, tra burocrate e cittadino.

Nella prospettiva normativa, questo filone ha cercato di chiarire in quale misura il meccanismo del voto, come surrogato del mercato, possa approssimare risultati efficienti nella fornitura dei beni pubblici. In seguito ci occuperemo soltanto di questo aspetto, prendendo le mosse dall'importante contributo di Arrow.

## ► Il Teorema di Arrow

Dal punto di vista normativo, il problema che maggiormente ha colpito gli economisti pubblici è stato lo studio delle relazioni tra processo del voto e formazione delle scelte sociali, ovvero la costruzione della funzione del benessere sociale, che come abbiamo visto costituisce il coronamento della teoria normativa dell'ottima allocazione delle risorse. In altre parole ci si è domandati se attraverso meccanismi di voto sia possibile pervenire alla definizione di una funzione che sia in grado di ordinare diverse alternative sociali, rispettando principi etici di carattere generale desiderabili per una società di tipo democratico come quella in cui viviamo. L'autore che ha dato il contributo fondamentale a questa ricerca è il premio Nobel K. Arrow. Un risultato molto importante della sua ricerca – ancora una volta, però, di tipo negativo – è il *Teorema dell'impossibilità*, dimostrato nel 1951.

Il suo modo di procedere è, dal punto di vista metodologico, di tipo assiomatico e consiste nell'applicazione alle scelte collettive della medesima impostazione che sta alla base della teoria della domanda del consumatore. Il ragionamento è impostato nel seguente modo. Si immagina che  $K$  individui debbano esprimere le loro valutazioni rispetto a  $N$  situazioni sociali e che ciascun individuo esprima un ordinamento delle proprie preferenze.

Il teorema afferma che attraverso meccanismi di aggregazione delle preferenze individuali, ad esempio meccanismi di voto, non è possibile definire una regola di scelta collettiva che, partendo da ordinamenti di preferenze individuali completi e transitivi, sia in grado di fornire un ordinamento delle alternative completo e transitivo che soddisfi simultaneamente i seguenti ragionevoli assiomi.

1. *Indipendenza dalle alternative irrilevanti*: la scelta sociale fra due alternative,  $A$  e  $B$ , dipende esclusivamente dall'ordinamento delle preferenze individuali circa quelle due alternative. Le preferenze relative ad altre alternative non giocano alcun ruolo. In particolare si esclude la possibilità di valutare l'intensità delle preferenze.

2. *Principio di Pareto debole*: se tutti gli individui preferiscono l'alternativa  $A$  all'alternativa  $B$ , anche a livello sociale deve valere la stessa preferenza.

3. *Dominio non ristretto*: la regola di scelta collettiva deve essere in grado di funzionare per qualunque struttura delle preferenze individuali.

4. *Non dittatorialità*: la regola di scelta collettiva non deve coincidere necessariamente con il sistema di preferenze di un unico individuo della società.

Il risultato, che getta una luce negativa sull'ottimalità del mondo in cui viviamo, non ha potuto essere molto attenuato nonostante gli sforzi di ricerca ad esso dedicati. Non ci è possibile fornire una rassegna di questo lavoro di ricerca, molto raffinato e complesso; basti tuttavia sapere che gli studiosi, preso atto della indiscutibile solidità del teorema di Arrow, hanno tentato di vedere se risultati soddisfacenti dal punto di vista normativo possono essere raggiunti eliminando o attenuando i requisiti degli assiomi sopra elencati. Si è cercato in particolare di attenuare il ruolo di quello che appare meno rilevante eticamente (il punto 1). I passi avanti compiuti non sembrano però molto significativi.

Il teorema dell'impossibilità di Arrow sembra dunque avere posto un ulteriore enorme ostacolo alla possibilità di definire dei criteri di formazione delle scelte collettive coerenti con i presupposti della teoria normativa della moderna Economia del benessere.

## ► I meccanismi di votazione

Lo studio delle caratteristiche dei meccanismi di votazione è stato assai coltivato anche in un passato molto lontano. Tracce di questi studi si possono trovare in tutta la letteratura politica sulla democrazia, a partire dall'età della Grecia classica, con intensificazioni nel pensiero politico del '700, che ha posto le basi dei sistemi rappresentativi emersi con la Rivoluzione francese. Tra gli studiosi più vicini alle problematiche dell'economia pubblica, vanno segnalati Wicksell, Dodgson (alias Lewis Carroll, l'autore di *Alice nel paese delle meraviglie*) nel secolo XIX e nel XX Black, Downs, Breton, Buchanan e Tullock.

L'analisi di questo problema è stata spesso svolta in contesti molto semplici, immaginando procedure decisionali chiamate «decisioni dei comitati», in cui un gruppo di individui deve prendere decisioni su proposte alternative. A questo tipo di decisioni si riferiranno gli esempi che vedremo in seguito, anche se va ricordato che i processi di votazione politica concreti devono soddisfare requisiti in parte diversi da quelli richiesti dalle decisioni di comitato. Nelle votazioni politiche non si votano né singole proposte di fornitura di servizi pubblici, né solo decisioni di tipo economico, ma viene definita una rappresentanza politica con mandato molto generale sulla base di valutazioni assai complesse.

*Il principio di unanimità di Wicksell*. Una breve rassegna dei meccanismi deve prendere le mosse dal contributo di Wicksell. Al termine della sua analisi dell'offerta dei servizi collettivi, egli giunge alla conclusione, che apparirà forse ovvia ma che resta comunque importante, che solo un sistema di voto che realizzi l'unanimità può essere rispettoso del principio di Pareto, in quanto l'esito del voto non comporta un danno per nessuno dei partecipanti. Naturalmente Wicksell era ben consapevole della difficoltà di poter raggiungere decisioni all'unanimità e pertanto la sua proposta era quella di assumere come criterio decisionale un metodo di voto che imponga un'*unanimità relativa*, vale a dire una maggioranza molto qualificata. Si tratta di un'indicazione alquanto indeterminata e ci si potrebbe chiedere che cosa significhi maggioranza qualificata e se essa non possa coincidere con la maggioranza assoluta.

La proposta di Wicksell andrebbe tuttavia letta all'interno della sua visione dell'articolazione del bilancio pubblico che prevedeva la votazione di singoli progetti fiscali a cui corrispondevano specifiche forme di entrate (vale a dire imposte di scopo o *earmarked taxes*). Essa va ricordata perché fa riflettere sull'opportunità di realizzare, attraverso la discussione e la continua revisione delle proposte, soluzioni di compromesso che possano ricevere il massimo grado di consenso.



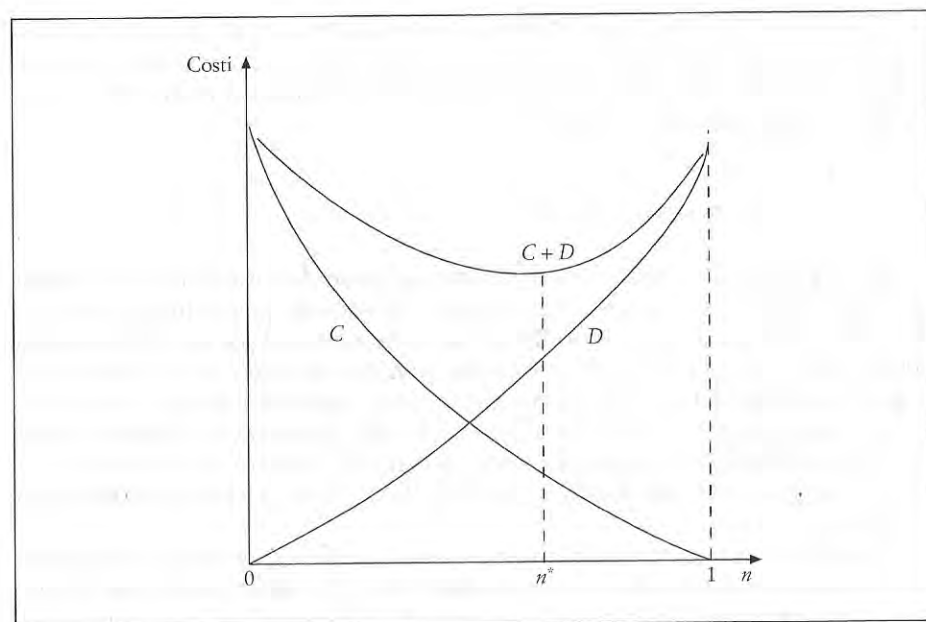


FIGURA 1.6. Costi esterni e costi interni del voto.

Come ha messo in rilievo Buchanan, la scelta stessa del metodo di votazione dovrebbe rispondere a due criteri, che sono riconducibili ad un problema di minimizzazione di costi.

Se si adotta un sistema con una maggioranza bassa (al limite decide solo il dittatore) si produrranno costi connessi al fatto che la volontà di uno solo sarà con ogni probabilità in conflitto con il volere degli altri. All'aumentare del *quorum* richiesto per l'approvazione tale costo tende a diminuire, e nel caso di unanimità esso sarà nullo. Indicheremo tale costo come costo *esterno* del meccanismo di voto.

D'altro canto il processo di raggiungimento di un consenso sempre più ampio comporta costi di altro tipo, in termini di tempo di decisione, ecc. Indicheremo tale costo come costo *interno* del meccanismo di decisione.

Rappresentiamo nella figura 1.6 con  $C$  i costi esterni e con  $D$  i costi interni del meccanismo di voto.  $C$  tende dunque a diminuire al crescere della proporzione  $n$  degli elettori favorevoli richiesti dalla regola di voto, fino ad annullarsi in caso di unanimità ( $n = 1$ ), mentre  $D$  seguirà un andamento opposto. Una collettività che agisca razionalmente sceglierà allora quella regola di voto che minimizza il costo totale ( $C + D$ ) interno ed esterno del processo decisionale collettivo.  $n^*$  individua quindi la regola di voto ottimale: infatti in corrispondenza di tale proporzione di votanti favorevoli la somma verticale delle due curve raggiunge il minimo. Va sottolineato come questa maggioranza ottimale non necessariamente coincide con la regola più utilizzata nelle moderne democrazie che è la maggioranza assoluta.

*La votazione a maggioranza.* È il sistema più utilizzato nelle decisioni collettive, ma, come si vedrà, il suo favore politico non trova sostegno adeguato

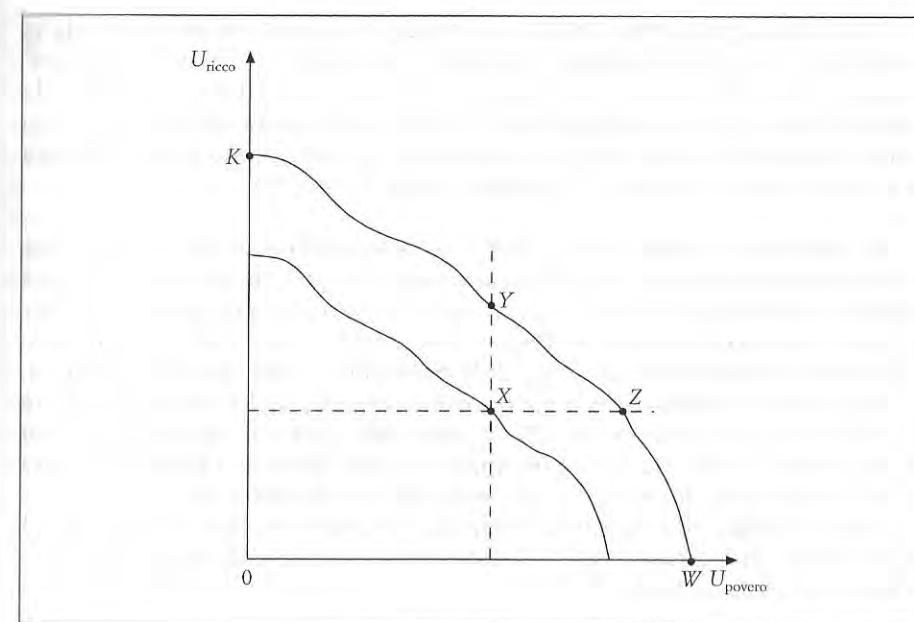


FIGURA 1.7. Voto a maggioranza e principio di Pareto.

nelle teorie del voto, che mostrano soprattutto i limiti di tale meccanismo di decisione.

Un primo modo per mostrare la contraddizione è un esempio formulato da Mueller, autore di un'importante opera, *Public Choice* (1989), successivamente ampliata e aggiornata, secondo cui la regola di voto a maggioranza non necessariamente restringe i possibili risultati alle alternative Pareto ottimali o, più debolmente, a quelle *potenzialmente* Pareto ottimali (vale a dire a situazioni in cui c'è uno che guadagna e uno che perde, ma esiste la possibilità di potenziale compensazione del secondo, ad opera del primo).

Nella figura 1.7 si immagina un'economia in cui sono presenti due gruppi di individui, ricchi e poveri. Si considera dapprima una situazione in cui non vi siano beni pubblici e si traccia la grande frontiera dell'utilità relativa a questa situazione. Si supponga di trovarsi nel punto  $X$  e si introduca la possibilità di produrre un bene pubblico desiderato da entrambi i gruppi. A causa della natura di esternalità presente nel bene pubblico, la sua produzione fa espandere la grande frontiera dell'utilità, che diviene la curva  $KW$ . Quale nuovo punto sarà selezionato sulla nuova frontiera? L'insieme di punti a cui corrisponde un miglioramento paretiano è  $YZ$ . Ma è facile vedere che se i ricchi sono in maggioranza sceglieranno un meccanismo di finanziamento del bene pubblico che li avvantaggi in modo netto: con ogni probabilità sceglieranno un punto nel tratto  $YK$ . Viceversa i poveri. È quindi molto probabile che le alternative cooperative, quelle sul tratto  $YZ$ , vengano escluse dalla scelta.

Il voto a maggioranza può avvenire secondo diverse modalità, che presentano però limiti e problemi di rilievo.

Il problema principale da cui conviene prendere le mosse riguarda la transitività. Come si è ricordato parlando del teorema di Arrow, un meccanismo di votazione, e cioè una regola per le scelte collettive, partendo da sistemi di preferenze individuali transitivi dovrebbe a sua volta garantire un ordinamento delle scelte collettive transitivo. La votazione a maggioranza non è però sempre in grado di garantire questo risultato.

1. *Votazione su coppie di alternative.* Il metodo di votazione che permette di conoscere la graduazione delle scelte collettive con riferimento all'intera gamma di alternative è il metodo dei confronti a coppie: ogni alternativa viene singolarmente confrontata con ciascuna delle altre e per ciascun confronto si stabilisce l'alternativa vincente. Tale procedura, studiata da Condorcet, un importante esponente della Rivoluzione francese, nel 1785, può portare all'individuazione di un *Condorcet-winner*, vale a dire di una proposta che batta sempre (e cioè in tutti i confronti a coppie) le altre. Questa proposta rappresenterà la scelta sociale fra le alternative in discussione.

Consideriamo, ad esempio (cfr. tabella 1.5), una società composta da 43 votanti divisi in quattro gruppi omogenei ( $\alpha$ ,  $\beta$ ,  $\gamma$ ,  $\delta$ ) chiamati a scegliere (votare) fra tre alternative: A, B, C.

TAB. 1.5. Schema di voto n. 1

| Votanti              | Tipologia di votanti |              |               |               |
|----------------------|----------------------|--------------|---------------|---------------|
|                      | Tipo $\alpha$        | Tipo $\beta$ | Tipo $\gamma$ | Tipo $\delta$ |
|                      | 18                   | 10           | 7             | 8             |
| Ordine di preferenza |                      |              |               |               |
| Primo                | A                    | B            | B             | C             |
| Secondo              | B                    | C            | A             | B             |
| Terzo                | C                    | A            | C             | A             |

Nel confronto fra A e B, B vince con 25 voti contro 18.

Nel confronto fra A e C, A vince con 25 voti contro 18.

Nel confronto fra C e B, B vince con 35 voti contro 8.

È possibile allora individuare un ordinamento completo e transitivo per le preferenze sociali,  $B > A > C$  (il simbolo  $>$  indica «preferito a») e un *Condorcet-winner*: l'alternativa B che vince, nei confronti a coppia, sia A che C.

L'ordinamento che emerge dalla votazione su coppie di alternative può però non essere transitivo: in questo caso non si ha un *Condorcet-winner*, e il metodo di votazione proposto non conduce ad alcuna scelta sociale. Si manifesta cioè il *paradosso di ciclicità del voto a maggioranza*.

Si veda il seguente esempio (cfr. tabella 1.6), in cui si hanno sempre 43 votanti, divisi in quattro gruppi omogenei ( $\alpha$ ,  $\beta$ ,  $\gamma$ ,  $\delta$ ) chiamati a scegliere fra tre alternative A, B, C.

Nel confronto fra A e B, B vince con 25 voti contro 18.

Nel confronto fra A e C, A vince con 25 voti contro 18.

Nel confronto fra C e B, C vince con 26 voti contro 17.

TAB. 1.6. Schema di voto n. 2

| Votanti              | Tipologia di votanti |              |               |               |
|----------------------|----------------------|--------------|---------------|---------------|
|                      | Tipo $\alpha$        | Tipo $\beta$ | Tipo $\gamma$ | Tipo $\delta$ |
|                      | 18                   | 10           | 7             | 8             |
| Ordine di preferenza |                      |              |               |               |
| Primo                | A                    | B            | B             | C             |
| Secondo              | C                    | C            | A             | B             |
| Terzo                | B                    | A            | C             | A             |

L'ordinamento delle preferenze sociali che emerge da questi confronti a coppie non è transitivo:  $B > A > C > B$ . La ciclicità deriva dal fatto che non si arriva mai ad una scelta: B è vincitore nel confronto con A, che però vince su C, che a sua volta vince su B. La votazione non conduce a nessuna scelta perché non si ha un *Condorcet-winner*.

Il metodo di votazione prescelto si rivela quindi inefficace.

La rilevanza di questo risultato è stata valutata dalla letteratura economica che si è proposta, innanzitutto, di individuare le condizioni in presenza delle quali il paradosso del voto a maggioranza non si manifesta.

Un requisito importante, messo in luce da Black, è che il profilo delle preferenze individuali di coloro che partecipano alla votazione sia unimodale (o, per usare l'espressione inglese, *single-peaked*, caratterizzato cioè da un'unica punta). Il profilo delle preferenze individuali può essere rappresentato graficamente ponendo in ordinata il livello (l'ordine) delle preferenze attribuibili a ciascuna alternativa (primo, secondo, terzo, ecc.) e in ascissa le singole alternative disposte in uno qualsiasi degli ordini possibili (ad esempio A, B, C oppure C, A, B, ecc.). È possibile allora individuare dei punti sul piano che associano ad ogni alternativa il posto che essa occupa nell'ordinamento del singolo votante. Collegando questi punti da sinistra a destra si ottiene una curva spezzata che rappresenta il profilo di preferenza del singolo votante o della singola tipologia di votanti.

Nella figura 1.8, che rappresenta i profili di preferenze indicati nella tabella 1.6, preferenze unimodali sono quelle corrispondenti ai votanti di tipo  $\beta$  (la moda è B),  $\gamma$  (la moda è B) e  $\delta$  (la moda è C); al contrario le preferenze dei votanti di tipo  $\alpha$  non sono unimodali ma hanno due mode locali (due punte) A e C.

I profili di preferenza indicati nella tabella 1.5 sono invece tutti unimodali. La presenza di unimodalità è condizione sufficiente per escludere la ciclicità del voto, è cioè condizione sufficiente per l'esistenza di un *Condorcet-winner*.

Questa affermazione richiede due precisazioni.

a) Nella costruzione della figura 1.8, sull'asse delle ascisse le alternative sono state poste nell'ordine A, B, C. Si tratta, come si è detto, di una scelta arbitraria. Si potrebbero costruire figure analoghe, adottando un ordinamento diverso. La presenza di unimodalità deve essere verificata con riferimento a tutti i possibili ordinamenti delle alternative. In altri termini, le preferenze dell'esempio della tabella 1.5 sono unimodali perché è possibile



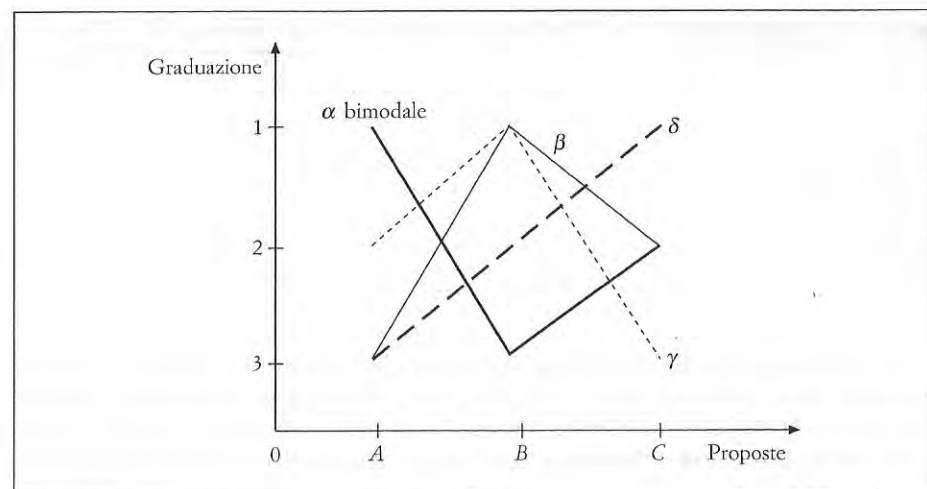


FIGURA 1.8. Profili di preferenze individuali.

darne una rappresentazione unimodale. Le preferenze della tabella 1.6 non sono unimodali perché, qualunque sia l'ordine con cui le alternative vengono rappresentate sull'asse delle ascisse, il profilo di preferenza di almeno una tipologia di votanti risulterà bimodale.

Casi di bimodalità sorgono frequentemente in scelte pubbliche dove emergono preferenze per le posizioni estreme: se  $A$ ,  $B$  e  $C$  del nostro esempio rappresentassero posizioni politiche di sinistra, di centro e di destra, essi si verificherebbero, ad esempio, se soggetti favorevoli alle posizioni della sinistra preferissero, come seconda alternativa, le posizioni di destra, piuttosto che quelle di centro.

b) La presenza di unimodalità è condizione sufficiente ma non necessaria per l'esistenza di un *Condorcet-winner*. Si consideri ad esempio lo schema di voto di cui alla tabella 1.7.

I profili delle preferenze individuali di questo schema di voto sono i medesimi di cui alla tabella 1.6. Le preferenze non sono unimodali. È però diverso, rispetto alla tabella 1.6, il peso relativo delle quattro tipologie di votanti. Dall'analisi dei confronti a coppie emerge con chiarezza l'esistenza di un *Condorcet-winner*: l'alternativa  $A$ , che vince sia contro  $C$  sia contro  $B$ .

TAB. 1.7. Schema di voto n. 3

| Votanti              | Tipologia di votanti |              |               |               |
|----------------------|----------------------|--------------|---------------|---------------|
|                      | Tipo $\alpha$        | Tipo $\beta$ | Tipo $\gamma$ | Tipo $\delta$ |
| Votanti              | 39                   | 1            | 1             | 1             |
| Ordine di preferenza |                      |              |               |               |
| Primo                | $A$                  | $B$          | $B$           | $C$           |
| Secondo              | $C$                  | $C$          | $A$           | $B$           |
| Terzo                | $B$                  | $A$          | $C$           | $A$           |

Questo risultato discende dalla fortissima omogeneità nei profili di preferenza dei votanti: 39 su 42 la pensano esattamente nello stesso modo. È chiaro che in condizioni di questo tipo l'aggregazione delle preferenze individuali non pone particolari problemi (si pone semmai un problema di tutela delle minoranze).

Analogamente, imporre unimodalità nelle preferenze significa imporre che esse non siano eccessivamente disperse, che esista cioè un'alternativa che, pur non essendo quella universalmente preferita, non trova opposizioni di rilievo all'interno della comunità. Se si esaminano le preferenze unimodali di cui alla tabella 1.5, si nota infatti che l'alternativa vincente,  $B$ , non trova opposizioni sociali di rilievo (in particolare, diversamente dalle alternative  $A$  e  $C$ , non occupa il terzo posto in nessuno dei profili di preferenze dei diversi gruppi di votanti). In altri termini, assumere che le preferenze siano unimodali significa assumere che vi sia un sufficiente grado di omogeneità nelle preferenze dei componenti la collettività. Assumere che le preferenze siano unimodali risolve quindi il problema del paradosso del voto a maggioranza. Va sottolineato però che ciò avviene al prezzo di violare uno degli assiomi del teorema di Arrow: l'assioma di Dominio non ristretto che, come si è ricordato, richiede che la procedura di aggregazione prescelta sia in grado di funzionare per qualsiasi struttura delle preferenze individuali.

All'ipotesi di unimodalità è stata data molta importanza in letteratura perché essa permette di collegare il risultato di un voto a maggioranza e il profilo di preferenze di quel particolare votante che rappresenta l'*elettore mediano*. Tale connessione è evidenziata nel cd. teorema dell'*elettore mediano*, che vale appunto nel caso di preferenze *single-peaked*. Supponiamo che le alternative tra cui scegliere (ad esempio, differenti livelli di spesa pubblica) possano essere tra loro ordinate secondo un qualche criterio (ad esempio, l'ordine crescente di ammontare). L'elettore mediano è quel votante rispetto al quale il numero degli individui che preferiscono alternative di ammontare inferiore (ad esempio, un livello di spesa pubblica inferiore) è esattamente uguale al numero degli individui che preferiscono alternative di ammontare superiore (ad esempio, un livello di spesa pubblica superiore).

Nella tabella 1.8 sono rappresentati gli ordinamenti di preferenza di cinque individui relativamente a cinque livelli di spesa pubblica (con  $A < B < C < D < E$ ).

La figura 1.9 illustra graficamente i corrispondenti profili di preferenze, tutti unimodali:  $\gamma$  corrisponde al profilo di preferenza dell'elettore mediano.

TAB. 1.8. Schema di voto n. 4

| Votanti              | Tipologia di votanti |              |               |               |                    |
|----------------------|----------------------|--------------|---------------|---------------|--------------------|
|                      | Tipo $\alpha$        | Tipo $\beta$ | Tipo $\gamma$ | Tipo $\delta$ | Tipo $\varepsilon$ |
| Ordine di preferenza |                      |              |               |               |                    |
| Primo                | $A$                  | $B$          | $C$           | $E$           | $D$                |
| Secondo              | $B$                  | $C$          | $D$           | $D$           | $E$                |
| Terzo                | $C$                  | $D$          | $B$           | $C$           | $C$                |
| Quarto               | $D$                  | $E$          | $E$           | $B$           | $B$                |
| Quinto               | $E$                  | $A$          | $A$           | $A$           | $A$                |

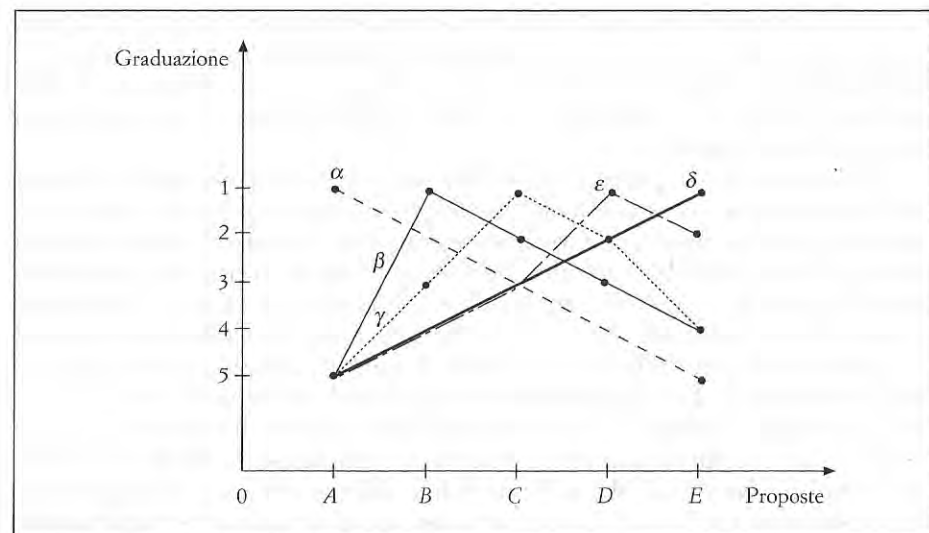


FIGURA 1.9. Profili di preferenze e teorema dell'elettore mediano.

Il teorema dell'elettore mediano afferma che in un voto a maggioranza l'alternativa di equilibrio (ad esempio, il livello di spesa pubblica di equilibrio), sempre che tale alternativa di equilibrio esista, coincide con l'alternativa (il livello di spesa pubblica) preferita dall'elettore mediano. Con riferimento alla figura 1.9 si può dimostrare che l'alternativa vincente nella procedura di voto sarà C, che è precisamente quella preferita dall'elettore mediano. C rappresenta infatti un *Condorcet-winner*: l'alternativa che vince in tutti i confronti a coppie. Come detto, tuttavia, la validità del teorema dell'elettore mediano richiede che un'alternativa di equilibrio per il voto a maggioranza esista. Come si è discusso sopra, tale esistenza è garantita dalla proprietà di unimodalità delle preferenze degli individui che partecipano al voto.

Il teorema dell'elettore mediano è molto importante anche in un contesto di teoria positiva. Esso consente ad esempio di spiegare perché in un sistema politico basato sulla competizione tra due coalizioni di partiti, posto che le preferenze degli elettori si esprimano tra posizione di destra e di sinistra rispettando l'unimodalità (non c'è cioè un votante che piuttosto del centro preferisce sia l'estrema destra, sia l'estrema sinistra), nell'arena politica avranno un ruolo decisivo le posizioni dell'elettore mediano. Non stupisce quindi che nel dibattito politico l'attenzione alle posizioni di centro sia sempre molto viva.

**2. Sistemi ad eliminazione.** Per eliminare il problema della ciclicità del voto, senza imporre unimodalità delle preferenze, nel caso di votazioni a maggioranza, sono stati elaborati metodi di votazione che prevedono la progressiva eliminazione delle alternative sconfitte.

**2.1. Votazione su coppie di alternative, con eliminazione dell'alternativa sconfitta.** Si consideri l'esempio della tabella 1.6, caratterizzato, come si è

detto, dall'assenza di un *Condorcet-winner*. Supponiamo che la votazione avvenga tramite confronti a coppie ma con eliminazione delle alternative di volta in volta sconfitte. Supponiamo che la sequenza delle votazioni preveda il confronto tra A e B, e del vincitore con C. B batte A (25 a 18). Nella votazione tra C e B, C batte B 26 a 17 e viene quindi scelta. Questo sistema di votazione conduce, quindi, inequivocabilmente ad una scelta. Esso non richiede infatti un ordinamento completo delle alternative sociali (che riproporrebbe il problema delle ciclicità delle preferenze) ma si limita ad individuare una modalità di selezione delle stesse. Il limite di questa procedura è che *l'esito del voto non è univocamente definito ed è quindi manipolabile*. Nel caso considerato, l'esito del voto dipende crucialmente dall'ordine in cui le diverse alternative sono poste in votazione.

Se consideriamo infatti una seconda possibile sequenza di votazioni: A contro C e il vincitore contro B, vediamo che: A batte C, 25 a 18. Nel confronto tra A e B, vince B, 25 a 18. Infine con la sequenza B contro C e il vincitore contro A, si avrà: B contro C, vince C, 26 a 17; C contro A, vince A, 25 a 18. A seconda della sequenza prescelta l'esito della votazione può quindi essere C, B o A. Questo risultato lascia indubbiamente insoddisfatti circa le qualità del metodo utilizzato ma mette anche in luce il grande potere che può avere il presidente di un'assemblea, che ha il diritto di decidere l'ordine con cui effettuare le votazioni.

**2.2. Maggioranza semplice.** Tutte le alternative vengono messe in votazione insieme e ciascuno vota su una sola di esse. Vince quella che ha più voti: le altre vengono quindi simultaneamente eliminate. Anche in questo caso si individua un'alternativa vincente, al prezzo però di rinunciare ad avere un ordinamento completo delle preferenze. Il votante è posto nella condizione di dovere dividere le alternative in due gruppi: quella preferita e le altre, considerate tutte insieme. L'ordinamento di queste ultime è irrilevante ai fini dell'esito della votazione. L'inconveniente principale di questa regola è ancora una volta la possibilità di manipolazione dell'agenda. Si consideri ancora una volta l'esempio della tabella 1.6. Con il voto a maggioranza semplice A risulterebbe vincente in quanto preferito, e quindi votato, da 18 individui, B da 17 = 10 + 7, C da 8. Se però la proposta C non ci fosse, vincerebbe B (25 a 18). Quindi qualcuno potrebbe essere indotto a inventare proposte fittizie e includerle nel processo di voto al solo scopo di distrarre voti dalla scelta tra A e B.

**2.3. Doppio turno.** Il doppio turno e la modalità di voto di cui al punto 2.4 sono modalità di votazione che combinano quelle precedenti, con la finalità di ridurre la manipolabilità. Nel caso del doppio turno, si hanno due distinte votazioni. Nella prima votazione si vota con il metodo della maggioranza semplice e si eliminano tutte le alternative sconfitte ad esclusione di quella che ha ricevuto il maggior numero di voti. Nella seconda votazione, tale alternativa è posta a confronto (ballottaggio) con quella che aveva vinto nel primo turno. Ancora nell'esempio della tabella 1.6, si può vedere che al primo turno passano le alternative A e B. Nel secondo turno A raccoglie 18 voti e B 25. Vince quindi B.



2.4. *Maggioranza semplice con eliminazione progressiva dell'alternativa con minor numero di voti.* Si vota a maggioranza semplice (nel senso che ciascun partecipante può dare un solo voto ad una sola alternativa), e si elimina ogni volta solo l'alternativa che riceve il minore numero dei voti. Questo sistema è più costoso perché richiede più votazioni e quindi può essere attuato da comitati ristretti, ma non si presta per elezioni più complesse.

È alquanto sorprendente e preoccupante constatare l'importanza che la scelta dei sistemi di votazione ha sugli esiti delle decisioni. Nell'esempio della tabella 1.6, si vede che, mentre con il metodo 2.1 può vincere qualsiasi alternativa a seconda dell'ordine seguito per le votazioni, con la maggioranza semplice vince *A*, con i procedimenti di cui ai punti 2.3 e 2.4 vince *B*. Ancora più preoccupante è notare, con riferimento all'esempio di tabella 1.5, che un metodo di votazione spesso utilizzato, la maggioranza semplice, potrebbe non fare emergere come alternativa vincente il *Condorcet-winner* che pure esiste (vincerebbe infatti l'alternativa *A*).

3. *Sistema di Borda.* Una modalità alternativa per superare il problema del paradosso del voto a maggioranza e arrivare comunque ad una scelta, senza imporre unimodalità e senza rinunciare alla costruzione di un ordinamento completo delle preferenze, è il metodo di Borda. Esso permette ai votanti di assegnare punteggi decrescenti alle alternative in votazione in relazione all'ordine che esse occupano nel suo sistema di preferenze. Ad esempio: se ci sono *N* alternative, si assegnano *N* punti a quella che ciascun individuo preferisce su tutte; *N* - 1, alla seconda, e così via e si sommano i punti. Nell'esempio della tabella 1.6 si attribuisce un punteggio di 3 all'alternativa preferita, 2 alla seconda e 1 alla terza. Il punteggio complessivo ottenuto da *A* è allora pari a  $18 \times 3 + 10 \times 1 + 7 \times 2 + 8 \times 1 = 86$ , analogamente quello di *B* è 85, e quello di *C* è 87. È quindi possibile stabilire un ordinamento completo e transitivo delle preferenze sociali  $C > A > B$ , e individuare univocamente un'alternativa vincente: *C*.

Questo metodo presenta però due limiti di rilievo.

a) In primo luogo esso viola uno degli assiomi di Arrow: l'assioma dell'Indipendenza dalle alternative irrilevanti che richiede, come si è ricordato, che la scelta fra due alternative dipenda esclusivamente dalle preferenze individuali circa quelle due alternative. Nell'esempio considerato, invece, l'alternativa *A* perde rispetto all'alternativa *C*, in un confronto a coppie, solo perché nel valutarla si tiene conto anche dell'esistenza dell'alternativa *B*, che pure non è in quel momento oggetto di scelta. L'esistenza dell'alternativa *B* condiziona infatti il punteggio assegnato alle alternative oggetto di votazione.

In termini più generali, l'attribuzione di punteggi nel metodo di Borda altro non è che una modalità per valutare e confrontare interpersonalmente l'intensità delle preferenze, cosa assolutamente in contrasto con l'approccio assiomatico del teorema di Arrow.

b) D'altro lato, proprio la possibilità riconosciuta al votante di assegnare un peso alle proprie preferenze rende il sistema di votazione di Borda particolarmente esposto al rischio di manipolazioni. I votanti possono infatti essere

indotti a non rivelare correttamente le proprie preferenze pur di conseguire il risultato voluto. Possono ad esempio attribuire un peso molto basso ad alternative che minacciano particolarmente quella preferita, anche se ciò non corrisponde alla propria vera preferenza.

### 3. Altre cause di fallimento del mercato

In un corso di scienza delle finanze la presenza di beni pubblici è la causa di fallimento del mercato su cui è naturale concentrare maggiormente l'attenzione. Fallimenti del mercato possono però manifestarsi anche per altre ragioni, che è utile cercare di identificare per vedere in quali casi l'intervento pubblico può rivelarsi una soluzione appropriata. Nella letteratura economica e quindi nei libri di testo si trovano modi diversi di classificare le cause di fallimento del mercato secondo uno schema logico. Qui seguiremo la stessa impostazione di Gravelle e Rees, autori di un noto testo di microeconomia.

Il mercato può essere definito come il luogo in cui si scambiano «diritti di proprietà e di uso» di beni o servizi. Il mercato fallisce se per qualche ragione non è possibile sfruttare la possibilità di raggiungere, attraverso lo scambio, posizioni Pareto efficienti. In termini generali, le cause di fallimento possono essere ricondotte a tre grandi categorie.

1. Difficoltà per le parti che operano nel mercato di trovare un accordo per uno scambio che pure sarebbe potenzialmente vantaggioso per entrambe.
2. Mancanza di controllo pieno sui beni o sulle risorse e sui modi di utilizzarle.
3. Mancanza o incompletezza delle informazioni necessarie allo scambio, o presenza di costi per ottenerle.

#### 3.1. Il monopolio

Un esempio di fallimento che rientra nella prima categoria – difficoltà per le parti a trovare un accordo – è costituito dalla presenza di posizioni di *monopolio*. Il monopolista pone il prezzo al livello a cui corrisponde l'eguaglianza tra ricavo marginale e costo marginale. Tale prezzo di equilibrio è superiore al costo marginale che corrisponderebbe all'ottimalità paretiana.

La figura 1.10 illustra il nostro caso, nell'ipotesi di costi marginali e costi medi costanti.

È del tutto chiaro che la soluzione di monopolio,  $Q_m$ , non è Pareto efficiente dato che vi sarebbe spazio per aumentare il benessere di qualcuno senza diminuire quello degli altri. I consumatori sarebbero infatti disposti a pagare, per il bene oggetto dello scambio, un prezzo superiore al costo unitario di produzione. Se fosse possibile stabilire un accordo tra monopolista e consumatori tale da spingere l'offerta fino al valore  $Q^*$ , al prezzo  $P^*$ , i consumatori ne ricaverebbero un vantaggio, misurato in termini di surplus del consumatore dal trapezio  $P_m ACP^*$ , superiore al costo del monopolista (riduzione del suo profitto) che è pari a  $P_m ABP^*$ . I consumatori potrebbero

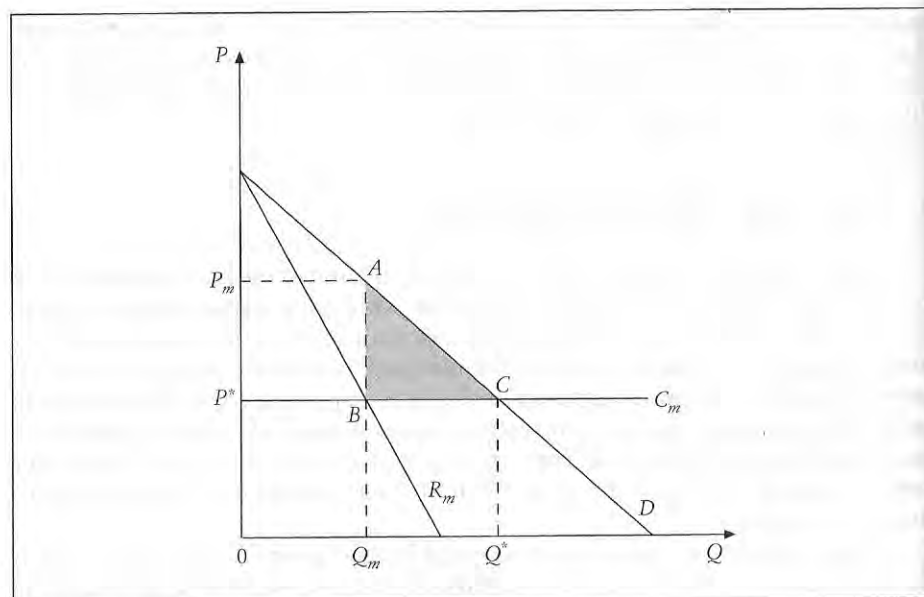


FIGURA 1.10. Monopolio e fallimento del mercato.

avere interesse a rifondere al monopolista la perdita da lui subita e a trarre comunque un vantaggio pari all'area del triangolo ABC.

Perché non si sfrutta questa possibilità? Raggiungere un accordo su come spartire i guadagni corrispondenti alla situazione Pareto ottimale non è semplice. Bisognerebbe trovare meccanismi per ripartire fra tutti i consumatori il risarcimento al monopolista, che per essere realizzati comportano costi o problemi di *free rider*. Tale difficoltà non si sarebbe presentata in una situazione di concorrenza, perché in tal caso produttori e consumatori assumono il prezzo come un dato e non vi è quindi possibilità da parte del produttore di scegliere il livello del prezzo in relazione all'andamento della domanda. Le posizioni monopolistiche portano dunque ad una sottoproduzione e ad un minore benessere ( $0Q_m < 0Q^*$ ).

L'intervento pubblico in tal caso può essere giustificato e le forme di intervento possono essere molto diverse. La più semplice sarebbe quella di rimuovere le cause che portano ad un solo offerente, promuovendo regole del gioco che consentano la massima estensione della possibilità di entrata da parte di più imprese, imponendo in tal modo condizioni che si avvicinano a quelle della concorrenza.

Vi sono però casi in cui questa politica può rivelarsi inefficace. Nell'economia pubblica hanno assunto un interesse particolare quelle situazioni di mercato in cui il formarsi di posizioni monopolistiche risulta inevitabile. Tali situazioni sono chiamate di *monopolio naturale*. Più esattamente un monopolio naturale tende a formarsi allorché il costo di fornire una data quantità da parte di una sola impresa è inferiore alla somma dei costi che potrebbero sopportare imprese di dimensioni minori, ciascuna delle quali contribuisca solo parzialmente all'offerta complessiva (principio di *subadditività dei costi*).

In simboli, se  $C_A(Q)$  rappresenta il costo della grande impresa e  $C_i(Q_i)$  i costi delle imprese più piccole, si avrà:

$$C_A(\sum Q_i) < \sum C_i(Q_i)$$

ove  $\sum Q_i = Q$ .

È chiaro che se la produzione del bene di cui si sta discutendo presenta *costi medi decrescenti* per l'intero tratto rilevante della domanda, l'offerta da parte di una sola impresa avverrà a costi minori dell'offerta ripartita fra più imprese. Poiché le caratteristiche sopra indicate della funzione dei costi si manifestano certamente nel caso di un'impresa monoprodotto caratterizzata da economie di scala, nella manualistica meno aggiornata il concetto di monopolio naturale è associato alla presenza di economie di scala. In realtà la presenza di economie di scala è solo una condizione sufficiente, ma non necessaria perché esista un monopolio naturale che risponda alla definizione generale che abbiamo sopra esposto. Il principio della subadditività può essere verificato anche se la domanda taglia i costi medi quando questi hanno già un andamento crescente, se il costo medio delle unità frazionate si mantiene comunque superiore al costo medio dell'impresa monopolista. Se però la domanda tende ad espandersi, essa taglierà la curva dei costi medi in tratti via via più crescenti ed è quindi sempre più probabile che il principio della subadditività dei costi non sia verificato e quindi cessi di esistere la posizione di monopolio naturale. Le economie di scala non sono quindi una condizione necessaria per il monopolio naturale. La presenza di un monopolio naturale può anche dipendere dal livello della domanda: quanto più questa è elevata, tanto più è probabile che la condizione di monopolio naturale cessi di esistere.

Il problema si complica se si considerano imprese multiprodotto o a produzione congiunta e la possibile esistenza di economie di scopo, che consentono economie di costo grazie alla varietà di produzione realizzata ottimizzando l'uso degli input produttivi. Nel caso, ad esempio, di produzione congiunta di due beni, si dice che esistono economie di scopo se il costo della produzione congiunta di una data quantità dei beni è inferiore al costo della loro produzione separata. Si può dimostrare che in questi casi l'esistenza di economie di scala non è neppure una condizione sufficiente perché si possa parlare di monopolio naturale, ma diviene rilevante accertare l'intensità delle economie di scopo presenti nella produzione. Questi affinamenti della definizione di monopolio naturale mostrano che l'accertamento della situazione di monopolio naturale è complesso e dipende non solo da fattori tecnologici ma anche dal livello della domanda.

Qualora la presenza di un monopolio naturale sia accertata, politiche volte a favorire la concorrenza non sarebbero corrette, perché l'offerta da parte di una sola impresa risulta, per quanto detto, più efficiente dell'offerta atomizzata. D'altro canto, l'offerta realizzata da un'impresa sola risulterebbe subottimale per le ragioni già dette ( $R_m = C_m$  e non  $P = C_m$ ).

Beni e servizi la cui produzione avviene a costi decrescenti e che, alle condizioni dette, tendono a produrre monopoli naturali sono in generale le



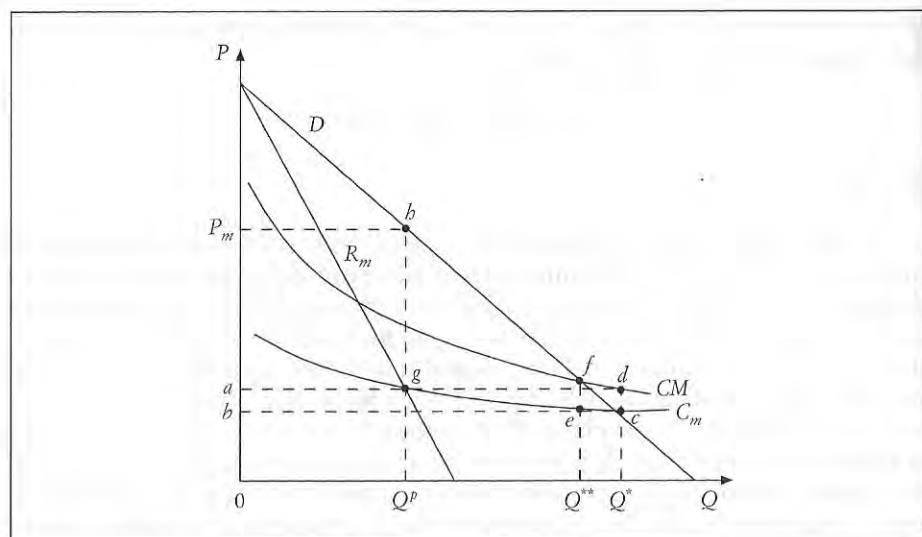


FIGURA 1.11. Fallimenti del mercato e monopolio naturale.

produzioni a cui corrisponde un volume di costi fissi molto elevati. Tra questi si distinguono i servizi a rete, vale a dire quei servizi che per essere prodotti necessitano di impianti distribuiti territorialmente: rotaie per le ferrovie, rete elettrica per l'energia elettrica, condutture per i servizi del gas e dell'acqua, cavi per i servizi telefonici, ripetitori per i servizi televisivi.

La presenza di monopolio naturale rappresenta quindi un caso in cui è spesso indicato un intervento pubblico volto sia ad impedire l'entrata inefficiente, sia a promuovere un'offerta adeguata. Si deve però sottolineare che il fondamento dell'intervento pubblico non può essere dato per scontato una volta per tutte: l'evoluzione della tecnologia (sotto forma di modificazione della presenza di economie di scala e di scopo) e della domanda può modificare il quadro e rendere razionale l'abbandono dell'intervento pubblico.

Posto che un monopolio naturale esista e si sia deciso di optare per la soluzione della produzione pubblica, la figura 1.11 ci consente di mostrare come l'intervento pubblico possa operare al fine di garantire l'efficienza paretiana che la gestione monopolistica privata non consentirebbe di realizzare. Il monopolista privato fisserebbe infatti la quantità al livello  $Q^p$ , dove il ricavo marginale è uguale al costo marginale, mentre lo Stato che intende realizzare un ottimo paretiano dovrebbe fissarla al livello  $Q^*$ , dove il costo marginale è uguale al prezzo. Si noti però che in corrispondenza di tale punto, il costo medio supera nella misura del segmento  $dc$  il costo marginale, e pertanto il monopolista che adotti questa politica (sia esso pubblico o privato) incorrerà in una perdita pari a  $abcd$ , corrispondente ai costi fissi, che la politica Pareto ottimale non consente di coprire. Affinché le condizioni di ottimalità siano preservate è quindi necessario supporre che lo Stato copra questa perdita con imposte in somma fissa, al fine di evitare distorsioni. Sappiamo però che l'applicabilità di imposte in somma fissa è assai problematica: anche l'intervento pubblico difficilmente potrà realizzare una situazione di First

Best. Se non è possibile introdurre *lump sum taxes*, ci si dovrà accontentare di soluzioni diverse, e molte possono essere individuate, che hanno comunque la caratteristica di essere di Second Best.

Una prima soluzione potrebbe essere quella di coprire i costi con imposte, ancorché di tipo diverso dalle imposte *lump sum* e quindi distorsive.

Una seconda soluzione, molto accreditata, è quella di fissare il prezzo in misura pari al costo medio, e quindi di fissare il livello di produzione pari a  $Q^*$ . I costi fissi sono ora coperti, ma si manifesta una perdita di benessere che può essere misurata dall'area  $cef$ . Tale perdita è comunque inferiore a quella che si sarebbe verificata in caso di monopolio privato ( $cgh$ ).

L'aspetto redistributivo, inevitabilmente presente in un mondo di Second Best, appare chiaro nei due casi esemplificati. Nel primo, l'onere dell'eccesso dei costi sui ricavi è addossato a tutta la collettività; nel secondo, esso grava invece solo su coloro che fanno domanda del bene interessato.

Gli interventi relativi a questo tipo di fallimento hanno prodotto un vasto campo di studi, del quale due linee meritano di essere ricordate.

– La *teoria delle imprese pubbliche* studia le modalità di fissazione delle tariffe da parte di imprese pubbliche che operino in regime di monopolio naturale, compatibili con il massimo dell'efficienza paretiana.

– Lo studio dei meccanismi di *regolamentazione e controllo della produzione di servizi pubblici*, offerti dal settore pubblico, ma prodotti da imprese private, appare oggi particolarmente rilevante, dato che il problema delle privatizzazioni è all'ordine del giorno.

Queste due aree di studio saranno affrontate nel capitolo settimo.

### 3.2. Le esternalità

Al secondo tipo di cause di fallimento del mercato (mancanza di controllo pieno sulle risorse) possiamo associare il fenomeno delle *esternalità*, che sta assumendo un ruolo sempre più importante nello studio delle caratteristiche ottimali dei sistemi economici decentrati.

Un'esternalità esiste quando alcune delle variabili che influenzano il costo di un produttore o l'utilità di un consumatore sono direttamente influenzate dalla decisione di produzione o di consumo di un altro soggetto, e tale effetto non è valutato o compensato.

È possibile distinguere quattro tipi di effetti esterni, a seconda che il soggetto colpito dall'esternalità e il soggetto causa dell'esternalità siano un consumatore o un produttore. Le esternalità possono poi essere ulteriormente distinte in positive o negative a seconda del segno dell'effetto prodotto. Possiamo fare alcuni esempi.

– Produttore/consumatore negativa: un'impresa inquina l'aria di una zona residenziale o l'acqua di una zona balneare; il disboscamento della foresta amazzonica produce squilibri ecologici e malattie.

– Produttore/produttore negativa: un'industria inquina l'acqua utilizzata da un'impresa agricola.

– Produttore/produttore positiva: un'impresa sviluppa un software o attività di R&S di cui si appropria un'altra impresa.

- Consumatore/produttore negativa: il traffico autostradale di automobili provoca rallentamenti al trasporto di merci delle imprese.
- Consumatore/produttore positiva: il giardino pieno di fiori e di alberi da frutta di una villa avvantaggia l'attività di un apicoltore vicino.
- Consumatore/consumatore negativa: il vicino tiene troppo alto lo stereo.
- Consumatore/consumatore positiva: un consumatore coltiva un bel giardino che allietta la vista del vicino.

Si è giustamente osservato che gli effetti esterni positivi del tipo consumatore/consumatore presentano affinità con il concetto di bene pubblico. Un bene pubblico può infatti essere interpretato come un caso limite di effetto esterno positivo. Si ha un'externalità positiva se il consumo di un'unità di bene produce utilità non solo al soggetto che compie l'atto di consumo, ma anche, seppure in misura presumibilmente inferiore, ad altro soggetto. Nel caso limite in cui il vantaggio del bene che produce un'externalità positiva risulti identico per tutti i soggetti, ci troviamo nella stessa situazione caratteristica di non rivalità nel consumo propria di un bene pubblico.

Il concetto di externalità prevede anche che i costi o benefici esterni non siano valutati o compensati, vale a dire che il soggetto che produce un'externalità non tenga conto degli effetti esterni nelle sue decisioni di consumo o di produzione. Questa circostanza può dipendere dal fatto che l'allocatione dei diritti o, più in generale, gli usi e costumi sociali consentano al soggetto che produce externalità di considerare nelle sue decisioni solo i costi e i benefici privati.

Anche se molti degli esempi fatti possono apparire scolastici, si può cogliere la rilevanza del concetto di externalità in un campo di studio molto importante come l'*economia ambientale*, in cui si presentano soprattutto externalità negative del tipo produttore/consumatore o produttore/produttore.

Qui considereremo solo un caso, produttore/produttore, prendendo come esempio quello dell'attività inquinante di un'impresa industriale *A* che danneggia un'attività agricola *B*. Formalmente l'externalità negativa può essere rappresentata in questo modo:

$$C_B = C_B(Q_A, Q_B) \quad \text{con} \quad \partial C_B / \partial Q_A > 0$$

ove  $Q_A$  e  $Q_B$  sono i livelli di produzione di *A* e *B* e  $C_B$  la funzione di costo di *B*.

Il costo di *B* dipende positivamente dalla quantità prodotta dall'impresa *A*, configurandosi in tal modo l'esistenza di un'externalità negativa.

La figura 1.12 mostra la produzione di equilibrio del produttore *A* (punto *B*, dove il costo marginale è uguale al prezzo) e il costo marginale esterno ( $C_{MAE}$ ) causato al produttore *B*, che abbiamo supposto lineare e crescente: al margine, il costo marginale inflitto da *A* al produttore *B* cresce al crescere del livello produttivo di *A*. Il nostro ragionamento non muterebbe anche per diversi andamenti del  $C_{MAE}$ .

La produzione offerta da *A*, in corrispondenza del punto in cui  $P_A = C_{MA}$ ,  $Q_A$ , non rappresenta tuttavia una situazione Pareto efficiente, perché l'effetto

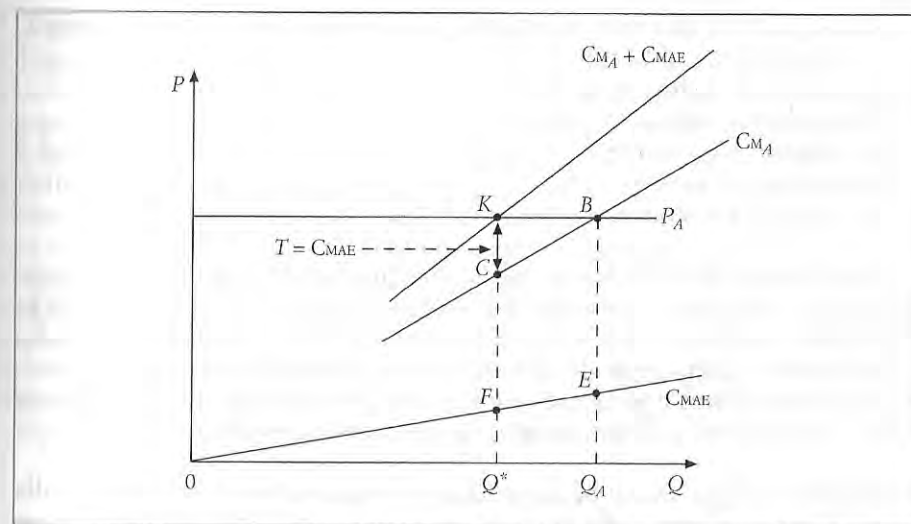


FIGURA 1.12. Externalità negativa produttore/produttore: decisioni di produzione dell'impresa *A* che inquina.

esterno negativo non è stato considerato dall'impresa *A* nella sua decisione di produzione. Ricordiamo la regola generale: un'allocatione è Pareto efficiente se il costo marginale è uguale al beneficio marginale. Mentre nel caso di un bene che non produce effetti esterni i costi e i benefici marginali da considerare sono solo quelli dei soggetti che comprano o vendono il bene, nel caso di un bene che produce un'externalità c'è differenza tra costi e benefici marginali sociali e costi e benefici marginali privati dei soggetti coinvolti direttamente nella transazione. Nel nostro caso il beneficio marginale sociale è dato dal prezzo, mentre il costo marginale sociale è dato dalla somma tra il costo marginale interno, sostenuto dall'impresa che produce il bene, ed il costo marginale esterno, subito dal soggetto danneggiato dall'externalità. La quantità del bene sarà dunque efficiente quando il suo costo marginale sociale uguaglia il beneficio marginale sociale.

Dal punto di vista sociale l'ottimo sarebbe nel punto *K*, a cui corrisponde una produzione pari a  $Q^*$ . In tale punto il prezzo (che rappresenta la valutazione marginale sociale del bene prodotto da *A*) eguaglia  $KQ^*$ , pari alla somma del costo marginale interno di  $Q^*$ ,  $CQ^*$ , e dell'externalità negativa  $KC$ , con  $KC = FQ^*$ .

Il caso descritto mostra che il mercato tende a generare una situazione di sovrapproduzione rispetto a quella ottimale.

È importante notare che la soluzione Pareto ottimale non corrisponde ad un livello di inquinamento nullo. Il punto  $Q^*$  definisce invece un *grado di inquinamento ottimale*.

I rimedi a situazioni di questo tipo sono molteplici e spesso possono essere realizzati attraverso interventi dello Stato. In generale essi consistono nell'individuazione di meccanismi ed incentivi che portino gli operatori a prendere decisioni che tengano conto anche dei costi esterni, vale a dire ad «internalizzare» le externalità.



Tra le possibili soluzioni si possono ricordare succintamente le seguenti:

1. produzione pubblica;
2. fusione delle imprese;
3. regolamentazione;
4. imposte pigouviane;
5. teorema di Coase;
6. diritti di inquinamento trasferibili.

*Produzione pubblica.* Tra le misure che comportano un intervento pubblico la più radicale è costituita dall'assunzione da parte dello Stato della produzione dell'attività che causa l'esternalità, al fine di realizzare il livello di produzione Pareto ottimale. Si tratta naturalmente di una soluzione estrema, che non sempre può essere opportuna e che se applicata su larga scala condurrebbe a forti ridimensionamenti di un'economia decentrata.

*Fusione delle imprese.* Se la produzione dei due beni facesse capo alla stessa impresa, questa nelle decisioni produttive sarebbe indotta a tenere conto dei costi esterni, che risulterebbero così «internalizzati». Una fusione delle due attività sarebbe quindi auspicabile. Non è semplice raggiungere tale risultato, ma lo Stato potrebbe intervenire con incentivi economici.

*Regolamentazione.* Si tratta di disposizioni di legge che impongono alle imprese di limitare la produzione di sostanze inquinanti entro parametri definiti. La reazione da parte dell'impresa a tali divieti può essere la riduzione della produzione, in corrispondenza al valore di inquinamento imposto dalla legge. In alternativa, l'impresa potrebbe realizzare investimenti in impianti di depurazione per limitare la produzione di sostanze nocive. Questa soluzione è quella più adottata nella pratica; talvolta essa è l'unica possibile, ad esempio quando si voglia vietare l'emissione di sostanze molto tossiche. Ma, come vedremo, essa è in generale meno efficiente di altre soluzioni economiche (imposte pigouviane e scambio dei diritti di inquinamento). Con la regolamentazione si impone infatti lo stesso comportamento a tutti gli operatori inquinanti, indipendentemente dal contesto ambientale e dalle condizioni in cui operano e ciò rende difficile conseguire un risultato efficiente. Supponiamo, ad esempio, che si imponga a due imprese che scaricano la stessa sostanza inquinante in due fiumi di diversa portata lo stesso limite alla quantità di scarica effettuabile. In questo caso si otterrebbe un grado di inquinamento molto diverso per i due fiumi, che non corrisponde al risultato che si vuole perseguire. Si potrebbe ovviare a questo inconveniente solo prevedendo limiti diversi per lo scarico delle due imprese, così da rendere pari o inferiore ad un livello prefissato l'inquinamento dei due fiumi. Questa soluzione è però difficilmente praticabile. Neppure sul piano applicativo la regolamentazione risulta superiore alle altre soluzioni indicate: anche essa richiede infatti controlli continui e l'applicazione di sanzioni agli inadempienti.

*Imposte pigouviane.* Una soluzione alternativa alla regolamentazione delle quantità di inquinamento è l'introduzione di imposte pari al costo marginale esterno, valutato nel punto efficiente, che inducano il produttore

al livello efficiente di produzione. Tale soluzione è stata proposta da A. Pigou nella sua importante opera *Economia del benessere* del 1920. Nel caso della figura 1.12, il livello dell'imposta è pari a  $KC$ . L'adozione di imposte pigouviane non costituisce tuttavia un rimedio di facile applicazione. Non è infatti chiaro come lo Stato possa acquisire le informazioni necessarie per definire il corretto livello dell'imposta, relativo al punto di equilibrio *ex post*, che consente di raggiungere l'ottimo paretiano. È però interessante ricordare l'esempio della *carbon tax*, un'imposta ecologica promossa dall'Unione europea, che colpisce i prodotti energetici in proporzione alla produzione di ossido di carbonio conseguente al loro consumo, che si ispira appunto a questo strumento.

Tale soluzione, che comporta un intervento pubblico, è stata criticata in modo più radicale da Coase, recente premio Nobel dell'economia, le cui argomentazioni sono note come teorema di Coase.

*Teorema di Coase.* Secondo questo studioso, le conclusioni degli economisti in tema di fallimento del mercato attribuibile ad effetti esterni negativi non sono corrette, perché non tengono in adeguato conto le possibilità di reazione degli individui alle esternalità stesse. L'intervento dello Stato, ad esempio utilizzando imposte pigouviane, non sarebbe necessario, data la natura «reciproca» dell'esternalità: un fumatore che inquina l'aria lede il diritto del suo vicino a respirare aria pura; nello stesso tempo, la proibizione di fumare impedisce al fumatore di esercitare un'attività che gli procura piacere. Secondo Coase il solo problema è quello di definire con esattezza l'attribuzione dei diritti ai soggetti interessati, se cioè *A* abbia o no il diritto di produrre inquinamento o viceversa se spetti a *B* il diritto di non subire gli effetti esterni di *A*. Una volta fissati i diritti e indipendentemente dalla scelta compiuta (se cioè il diritto di inquinare spetti ad *A* o no), è sufficiente lasciare alle parti la possibilità di contrattare, cioè lasciare operare ulteriormente il mercato, al fine di raggiungere una situazione efficiente. Se *A* ha il diritto di inquinare, *B* sarà disposto ad offrire un compenso per indurlo a limitare la produzione inquinante e a convincerlo a produrre la quantità ottimale.

Considerando la figura 1.12, se *A* ha il diritto di inquinare, la soluzione del mercato  $Q_A$  sarà solo un punto di partenza per una contrattazione privata con *B*. Quest'ultimo avrà infatti interesse ad offrire ad *A* un compenso per indurlo a produrre una quantità inferiore. Per ogni unità in meno prodotta da *A*, *B* consegue un vantaggio marginale pari al  $CMAE$ , in corrispondenza del nuovo livello produttivo, superiore alla riduzione del profitto marginale di *A*, che è misurata dalla differenza tra  $P_A$  e il costo marginale. Tale convenienza si verifica fino al raggiungimento della quantità  $Q^*$ . Oltre questo punto l'incentivo viene meno perché la riduzione marginale del beneficio che potrebbe derivare da un'ulteriore riduzione di un'unità della produzione diventa maggiore del  $CMAE$ .

Alla stessa soluzione – e comunque senza alcun bisogno di un intervento pubblico – si arriverebbe anche se a *B* venisse attribuito il diritto di non essere inquinato. In tal caso la situazione di partenza sarebbe un livello di produzione pari a zero, poiché in corrispondenza di una produzione nulla,

in cui si ha una totale assenza di inquinamento,  $B$  raggiunge il massimo benessere.  $A$  avrebbe però interesse ad offrire a  $B$  un compenso, inducendolo ad accettare un livello di produzione pari a  $Q^*$ . Il profitto marginale che egli otterrebbe da ogni incremento di un'unità di produzione fino a  $Q^*$  sarebbe infatti superiore all'esternalità subita da  $B$  e quindi sufficiente a ripagare quest'ultima.

Il teorema di Coase presenta tuttavia alcuni limiti. Quello più grave è costituito dal fatto che una contrattazione privata tra le parti interessate non è sempre agevole o possibile. Essa può funzionare nell'esempio qui esposto, ma non funzionerebbe in un caso di danno ambientale analogo a quello degli scarichi di fattori inquinanti nel Po, a causa della presenza di un numero molto grande di soggetti danneggiati e di inquinatori e per la conseguente difficoltà di mettere in contatto le parti interessate. In ogni caso, poi, la soluzione di Coase lascia indeterminato il problema redistributivo che inevitabilmente si deve affrontare quando si allocano i diritti di proprietà.

*Diritti di inquinamento trasferibili.* Si è visto che la soluzione ottimale dal punto di vista economico comporta comunque una qualche misura di inquinamento. Anche se lo Stato fosse in grado di individuare, compito invero di non facile soluzione, tale livello ottimale, in un'economia decentrata resterebbe aperto il problema di definire come si distribuisce tra le imprese esistenti il diritto di inquinare fino alla misura globale prefissata dall'intervento pubblico. Con un ragionamento più articolato e realistico di quello sottostante all'analisi condotta con la figura 1.12, si può immaginare che le imprese possano ridurre l'inquinamento non solo riducendo la produzione, ma anche attraverso investimenti (ad esempio in impianti di depurazione), e che l'efficienza delle imprese in queste attività di investimento, e quindi i costi marginali di abbattimento dell'inquinamento, non siano uguali per tutte le imprese. In questo contesto lo scambio dei diritti di inquinamento fornisce una risposta al problema sollevato in questo paragrafo. Lo Stato potrebbe cioè emettere dei *vouchers* (o diritti di inquinamento) per un ammontare pari all'obiettivo e distribuire, secondo determinati criteri, tali buoni alle imprese. Queste possono effettuare una produzione massima con un inquinamento pari al valore del buono ricevuto o, alternativamente, possono vendere tale diritto ad altre imprese. Si creerà un mercato nuovo che avrà come esito un'allocazione ottimale fra le imprese della produzione del livello prefissato di inquinamento, minimizzando i costi di abbattimento di questo.

In sostanza questo meccanismo consente di conservare gli aspetti positivi sia della regolamentazione, sia delle imposte pigouviane, grazie alla creazione di un nuovo mercato e lasciando gli operatori economici liberi, come è implicito nella soluzione di Coase, di trarre vantaggio dalle opportunità di scambi vantaggiosi così create. Tale soluzione evita poi alcuni aspetti negativi, come l'arbitrarietà nell'allocazione dei diritti di inquinamento nel caso della regolamentazione. Anch'essa ha però implicati effetti redistributivi, connessi alle modalità con cui avviene la distribuzione iniziale dei diritti.

Nel corso dell'ultimo decennio le applicazioni di questo strumento si stanno sempre più diffondendo. Scambi di diritti di inquinamento tra paesi sono uno degli strumenti previsti nel Protocollo di Kyoto del 1997 e sono

anche parte dal 2005 di un programma della Commissione europea (*European Union Emissions Trading Scheme*) che riguarda le emissioni di  $\text{CO}_2$ .

### 9.3. Asimmetrie informative: rischi sociali, mercato assicurativo e intervento pubblico

La terza causa di fallimento del mercato è costituita dalla presenza di carenze o di asimmetrie informative, che si suppongono assenti nel modello dell'Economia del benessere che consente di suffragare il primo teorema. È assai frequente il caso in cui i soggetti che entrano in relazioni contrattuali per realizzare uno scambio di beni o di servizi non dispongano di tutte le informazioni rilevanti sulla natura e sulle caratteristiche del rapporto contrattuale che stanno per mettere in essere. Queste carenze e/o asimmetrie informative danno origine a costi, noti come *costi di transazione*, vale a dire costi di negoziazione, monitoraggio e operatività del contratto. L'esistenza o meno di costi di transazione si traduce nella possibilità o meno di siglare *contratti completi*. I contratti sono completi quando è possibile prevedere e disciplinare tutte le contingenze che possono verificarsi nei diversi stati del mondo. Il contratto è cioè in grado di precisare gli impegni reciproci dei contraenti in relazione ad ogni possibile situazione che si venga a creare nel periodo post contrattuale di vigenza del contratto. Quando ciò non avviene i contratti si dicono *incompleti*, situazione che si verifica, ad esempio, ogni qualvolta qualche variabile rilevante ai fini della contrattazione non è osservabile (almeno da parte di un contraente) o anche se osservabile dalle parti, non è valutabile da terzi e cioè da soggetti estranei al contratto, quali le corti di giustizia, a cui ci si rivolge e che devono sindacare sull'esistenza di eventuali inadempienze contrattuali. In un contesto di incompletezza contrattuale si pone infatti il problema di chi abbia la facoltà di prendere decisioni in circostanze che non siano state specificate dal contratto. Il soggetto che può decidere in situazioni che non erano previste dal contratto viene definito come soggetto proprietario dei *diritti residuali di controllo*. L'allocazione di questi diritti è fondamentale e può influire sui costi e quindi sull'efficienza *ex post* degli scambi che vengono compiuti. Questi concetti assumono un ruolo particolare allorché si debba scegliere tra produzione pubblica o privata di un bene o servizio, o in termini più generali si debba definire, in contesti caratterizzati da incertezza, la forma migliore di governance di una determinata attività economica.

Nell'ambito delle situazioni caratterizzate da carenze informative, molta attenzione è stata dedicata al caso in cui l'informazione non sia egualmente distribuita tra i partecipanti allo scambio. Si parla in questi casi di *asimmetria informativa*, che è responsabile di particolari forme di fallimento del mercato (del tipo *moral hazard* e *adverse selection*, di cui i mercati assicurativi sono un caso esemplare).

In molti casi l'intervento dello Stato nell'offerta o nel finanziamento di servizi pubblici è giustificato dalla presenza del rischio che caratterizza talune attività umane, a cui i soggetti non sono in grado di fare fronte attraverso i tradizionali meccanismi assicurativi. I rischi che comunemente richiedono



TAB. 1.9. I vantaggi del «pooling» dei rischi

| Soluzione individuale         |             |                |             |             |
|-------------------------------|-------------|----------------|-------------|-------------|
|                               | Individuo 1 |                | Individuo 2 |             |
|                               | Costo       | Probabilità    | Costo       | Probabilità |
| Stato 1: nessuna perdita      | 0           | 0,5            | 0           | 0,5         |
| Stato 2: perdita              | -100        | 0,5            | -100        | 0,5         |
| Valore atteso                 | -50         |                | -50         |             |
| Varianza                      | 2.500       |                | 2.500       |             |
| Soluzione di «pooling»        |             |                |             |             |
|                               | Costo       | Costo unitario | Probabilità |             |
| Stato 1: nessuna perdita      | 0           | 0              | 0,25        |             |
| Stato 2: perdita solo per 1   | -100        | -50            | 0,25        |             |
| Stato 3: perdita solo per 2   | -100        | -50            | 0,25        |             |
| Stato 4: perdita per entrambi | -200        | -100           | 0,25        |             |
| Valore atteso                 |             | -50            |             |             |
| Varianza                      |             | 1.250          |             |             |

una protezione sono moltissimi: il furto, l'incendio, una malattia che riduce la capacità di lavoro, la perdita del valore della moneta, le variazioni del cambio, incidenti automobilistici, danni causati dal maltempo, ecc.

Come si può cercare di eliminare gli effetti negativi di possibili eventi rischiosi? Nessuno può evitare che si manifestino eventi negativi che dipendono dal caso: è però possibile, in alcune circostanze, ridurre la variabilità attesa degli eventi rischiosi, rendere il futuro meno incerto, anche se ciò comporta necessariamente un costo.

In alcuni casi i soggetti possono realizzare questo obiettivo attraverso accordi privati. Si consideri ad esempio il caso, descritto nella tabella 1.9, di due soggetti ai quali si prospetti, in modo del tutto indipendente, un'identica situazione futura rischiosa, caratterizzata dalla possibilità di una perdita pari a  $C = -100$  con probabilità  $p = 0,5$ . L'evento dannoso ha per entrambi i soggetti, isolatamente considerati, un valore atteso di  $-50$  e una varianza pari a  $2.500$ . I due soggetti possono però accordarsi nel senso di condividere i costi dell'evento rischioso. Se, ad esempio, la perdita di  $100$  si manifesta solo per un soggetto, essa viene equamente ripartita tra entrambi. È questo il principio del *pooling* dei rischi che è alla base della spiegazione dell'esistenza delle imprese di assicurazione. Sulla base di tale accordo la variabile casuale che descrive gli eventi futuri presenta ora quattro possibili determinazioni, a cui sono associate probabilità definibili sulla base della regola della probabilità composta, illustrate nella parte inferiore della tabella. Nella terza colonna è anche indicato il costo unitario sopportato da ciascun soggetto nell'ipotesi di condivisione dei costi. Come si vede, per entrambi i soggetti, il valore atteso della soluzione di *pooling* resta immutato, pari a  $-50$ , ma la varianza si è ora dimezzata ( $1.250$ , anziché  $2.500$ ).

Se immaginassimo un altro esempio con 3 soggetti, si potrebbe mostrare che la varianza si riduce ancora. La varianza si annulla al tendere all'infinito del numero dei soggetti che partecipano al *pooling* dei rischi.

Ma è allora chiaro che se l'evento futuro rischioso presenta una varianza nulla, ciò significa che la situazione rischiosa è stata trasformata in una situazione certa.

L'esempio fatto è abbastanza elementare (costi e probabilità identiche per tutti i soggetti). In casi più complicati la ricerca di un accordo tra soggetti potrebbe essere meno facile e dare origine a costi di transazione molto elevati. È però da questo principio che nasce la spiegazione economica di un'impresa di assicurazione e l'offerta di contratti di assicurazione, il cui ruolo è una funzione di intermediazione e di coordinamento di scelte individuali che consistono nel mettere in comune (*pooling*) i rischi.

Quando un'assicurazione può essere offerta dal mercato e quando invece fallisce? Per affrontare questo argomento, generalmente si ipotizza che i soggetti siano *avversi al rischio*, e cioè che posti davanti a due alternative che hanno lo stesso valore atteso, ma caratterizzate da diverso grado di rischio, preferiscano quella meno rischiosa. Si ipotizza anche che l'assicuratore sia invece indifferente rispetto alle due alternative, sia cioè *neutrale* rispetto al rischio. In queste ipotesi, l'assicurazione è un contratto attraverso il quale ad un soggetto, *avverso al rischio* ed esposto ad eventi rischiosi che assumono la forma di variabile casuale, è offerta la possibilità di trasformare, dietro corrispettivo di un prezzo, detto *premio*, gli eventi rischiosi in una situazione certa. Consideriamo un soggetto al tempo 0 al quale si prospettino, per il periodo 1 successivo, due possibili «stati del mondo». Stato 1: perdita di reddito pari a  $C$ , derivante, ad esempio, dai costi di cura in connessione all'evento rischioso di contrarre una malattia, che ha la probabilità di verificarsi pari a  $p < 1$ ; stato 2: perdurare dello stato di salute e quindi assenza di costi. La variabile casuale in tal caso assume valore  $C$  con probabilità  $p$  e valore 0 con probabilità  $(1 - p)$ . Con un contratto di assicurazione il soggetto può al tempo 0 trasformare questa situazione rischiosa in una situazione certa. L'assicurazione si impegna a fornire al tempo 1 al soggetto la somma  $C$  nel caso in cui si verifichi l'evento dannoso, ma esige in cambio, in ogni caso, il pagamento di un premio pari ad  $H$ .

Da cosa dipende l'interesse delle imprese assicuratrici ad offrire tali contratti? Il requisito essenziale è che l'evento di cui si tratta (nel nostro caso il rischio di malattia), incerto per l'individuo, sia pressoché certo dal punto di vista aggregato. Dal punto di vista individuale non sappiamo se si verificherà lo stato 1 o lo stato 2, ma prendendo in considerazione una collettività sufficientemente ampia, composta da  $n$  individui identici al soggetto qui considerato, è possibile affermare che nel periodo 1 si avranno  $pn$  casi di soggetti che risulteranno colpiti dalla malattia e  $(1 - p)n$  individui che si saranno conservati sani. Se l'impresa di assicurazione è neutrale nei confronti del rischio, prescindendo da costi di gestione e in un contesto di mercato assicurativo di tipo concorrenziale, essa potrà nel periodo 0 offrire il contratto di assicurazione sopra descritto, purché sia in grado di effettuare un numero sufficientemente ampio di contratti. L'assicurazione non sarebbe possibile se non vi fosse un numero sufficientemente grande di soggetti che fanno domanda di assicurazione, su cui distribuire il rischio. Se, invece, tali soggetti esistono, in un mercato concorrenziale, il premio  $H$  si fisserà a quel livello in cui il profitto atteso:  $E\pi = pn(H - C) + (1 - p)nH$ , è nullo.

Ciò si verifica quando  $H = pC$ . Ma  $pC$  è, per ciascun agente, il valore atteso dell'evento:  $pC + (1 - p) \times 0 = pC$ .

Con un premio definito in questa misura un soggetto avverso al rischio avrà convenienza a fare domanda di assicurazione. Se l'individuo non si assicura ha una perdita attesa pari a  $pC$ , che, come abbiamo appena visto, è uguale al premio che l'assicurazione avrebbe richiesto; se si assicura dovrà comunque sostenere un costo pari a  $H$ . La prima situazione è caratterizzata da rischio, mentre la seconda non è rischiosa, ed è quindi preferita.

Le ipotesi fatte riguardo alla propensione al rischio dei soggetti assegnano al contratto di assicurazione un'importante funzione per il raggiungimento di posizioni di ottimo paretiano. Esso consente di realizzare, rispetto allo *status quo*, una situazione che è preferita dall'assicurato e non causa un peggioramento dell'assicuratore. L'attività di assicurazione è cioè un processo che consente di aumentare il benessere economico *ex ante*. Tale situazione Pareto ottimale risulta caratterizzata dal fatto, decisivo per gli argomenti che svolgeremo, che il rischio venga *integralmente* sopportato dal soggetto neutrale nei confronti del rischio, mentre il soggetto che presenta avversione per il rischio potrà trovarsi in una situazione Pareto ottimale solo se è in grado di assicurare *integralmente* il rischio a cui è sottoposto.

Le modalità con cui viene fissato il premio di assicurazione mostrano che il fondamento di un'attività di assicurazione privata è costituito dalla possibilità di applicare criteri attuariali in modo agevole. Perché un mercato assicurativo privato funzioni è necessario, in altre parole, che siano presenti alcune condizioni.

1. La probabilità dell'evento assicurato per ciascun individuo deve essere indipendente da quella per qualsiasi altro individuo. L'assicuratore è infatti disposto ad offrire un contratto di assicurazione se nel complesso del fenomeno è possibile prevedere il numero dei casi favorevoli e di quelli sfavorevoli. Se il fenomeno casuale all'origine dell'assicurazione è prodotto da uno shock comune a tutti i soggetti (ad esempio, disoccupazione, o una malattia molto contagiosa), nessun assicuratore potrà svolgere la funzione di ripartire il rischio.

2. La probabilità dell'evento rischioso deve essere inferiore o non molto prossima all'unità. Si tratta di una precisazione ovvia, ma se la probabilità fosse pari ad 1, dato che l'attività delle assicurazioni comporta dei costi amministrativi che devono essere coperti dal premio, questo sarebbe maggiore del danno assicurato. Ciò spiega, ad esempio, il fatto che non vi siano assicurazioni disposte a coprire con contratto i rischi di malattie congenite o croniche, o ad assicurare contro la morte soggetti che abbiano già avuto infarti, o, in certe zone d'Italia, ad offrire assicurazioni contro il furto delle auto.

3. Le probabilità devono essere stimabili in modo sufficientemente certo. Uno dei motivi per cui non si può ottenere un'assicurazione contro l'inflazione futura è connesso alla difficoltà di calcolare la probabilità di tassi di inflazione futuri alternativi. Per usare una nota distinzione di Knight, il fenomeno oggetto di assicurazione deve ricadere nella categoria del *rischio* (a cui possono essere applicati criteri probabilistici) e non in quella dell'*incertezza* (a cui corrispondono fenomeni casuali ai quali non è possibile associare una probabilità di accadimento).

4. Non deve esistere asimmetria informativa, vale a dire una situazione in cui l'assicuratore disponga di una quantità di informazione, rilevante per definire il contratto di assicurazione, inferiore a quella dell'assicurato. Si possono configurare due tipi di asimmetria informativa: il *moral hazard* e l'*adverse selection*, che illustreremo in dettaglio nei prossimi paragrafi.

Sul tema generale, possiamo dire che tutte le volte in cui viene meno una delle condizioni indicate il mercato fallisce e possono sorgere istituzioni diverse dal mercato per mitigare gli effetti negativi dell'assenza di assicurazione, istituzioni cioè che svolgono una funzione di assicurazione, anche se non sulla base dei criteri strettamente attuariali a cui si attengono le assicurazioni private. Quasi sempre queste istituzioni sono pubbliche e pertanto i fallimenti del mercato attribuibili alla presenza di rischi non assicurabili *in toto* o in parte forniscono una spiegazione (e una giustificazione, in chiave normativa) dell'intervento pubblico.

Questo tipo di ragionamenti giustifica l'intervento dello Stato non solo per garantire livelli minimi di servizi necessari alla sussistenza, ma interventi ben più ampi come assicurazioni contro la disoccupazione, particolari tipi di malattia, la perdita del potere di acquisto dei risparmi accumulati per la vecchiaia.

I sistemi di sicurezza sociale che si sono sviluppati nei paesi occidentali a partire dalla fine del secolo scorso, giunti al loro massimo stadio evolutivo, prevedono forme di assicurazione per queste grandi categorie di fenomeni: malattia, invalidità, disoccupazione, vecchiaia. Le ragioni che hanno portato alla loro nascita e al loro sviluppo (e oggi alla loro messa in discussione, per il predominio di ideologie liberiste) non sono solo di natura economica. Altre motivazioni, di carattere etico, hanno avuto un grande rilievo (nel linguaggio dell'Economia del benessere, fattori che attengono più alle caratteristiche della funzione del benessere sociale che al soddisfacimento di criteri di efficienza paretiana). Le argomentazioni fondate su criteri di efficienza sono tuttavia molto importanti, in quanto forniscono ragioni che possiamo considerare relativamente indipendenti da giudizi di valore. Le considerazioni svolte sinora ci consentono di fornire spiegazioni e giustificazioni di importanti aspetti delle assicurazioni pubbliche contro la disoccupazione, la malattia e l'invalidità, che saranno riprese e approfondite nel capitolo ottavo.

#### ► «Adverse selection» e «moral hazard»

L'*adverse selection* (AS) e il *moral hazard* (MH) rientrano nel campo dell'analisi dei comportamenti economici in presenza di asimmetria informativa. Tali situazioni si verificano in molteplici contesti in cui due o più soggetti intendano definire un contratto (di lavoro, di assicurazione, di delega in generale) e uno dei contraenti ignori alcune informazioni rilevanti per il contratto che si intende stipulare, che sono invece note all'altro soggetto. Il soggetto che si trova nella situazione di informazione incompleta viene chiamato principale (P), mentre il soggetto che dispone di un'informazione privata è chiamato agente (A). Il rapporto che intercorre tra i soggetti è solitamente indicato come rapporto principale/agente. Il nostro scopo è di



spiegare perché la presenza di *As* e di *MH* possa essere causa di fallimenti del mercato, prendendo come esempio il caso dei contratti che si stipulano nei mercati assicurativi. Negli esempi che vedremo l'asimmetria informativa è rappresentata dal fatto che l'assicuratore non conosce alcune caratteristiche degli assicurati o non è in grado di controllare e conoscere alcuni loro comportamenti, che sono rilevanti per stabilire la relazione contrattuale. L'assicuratore sarà pertanto il principale *P*, mentre l'assicurato è l'agente *A*.

*Adverse selection.* Si riferisce a una situazione di asimmetria informativa in cui *P* non è in grado di conoscere alcune caratteristiche o informazioni in possesso di *A* (asimmetria informativa ad informazione nascosta) rilevanti ai fini della stipulazione del contratto di agenzia. Si immagina che l'insieme degli *A* non sia omogeneo, ma sia invece composto da sottoinsiemi (per semplificare ne supporremo due) che si differenziano per una caratteristica rilevante per la fissazione delle condizioni contrattuali. Tale caratteristica non è sotto il controllo di *A*, ma sussiste originariamente. Un esempio di contratto che ben si attaglia a questo caso è quello riguardante i contratti assicurativi del ramo vita, in cui la caratteristica qualitativa che differenzia i due gruppi di *A* sia ad esempio costituita dallo stato di salute di *A*. Soggetti malati sono ad alto rischio in quanto ad essi è associata una maggiore probabilità che si verifichi l'evento assicurato, rispetto a soggetti sani che sono a basso rischio. Se non ci fosse asimmetria informativa (impossibilità da parte di *P* di controllare, senza sostenere costi, se il soggetto che gli chiede di essere assicurato appartenga al gruppo dei sani o dei malati), *P* potrebbe offrire due tipi di contratti, uno con premi bassi per i soggetti sani, a basso rischio e uno con premi alti per quelli malati ad alto rischio, realizzando in tal caso un ottimo paretiano di First Best, in cui *P* non realizza perdite e tutti gli *A* (soggetti sia a basso che ad alto rischio) possono essere integralmente assicurati. Tale esito è però per ipotesi escluso nel caso di asimmetria informativa del tipo *As*, dato che *P* non è in grado di separare i sani dai malati. In presenza di informazione incompleta, quali tipi di contratto *P* potrebbe essere indotto ad offrire? Se *P* offrisse un solo contratto con premio basso, i soggetti ad alto rischio lo sottoscriverebbero volentieri, ma alla fine *P* incorrerebbe in perdite. Tale offerta non può quindi costituire un equilibrio duraturo. Se invece *P* offrisse un contratto con premio alto, i soggetti sani non avrebbero interesse a sottoscriverlo. Si realizzerebbe in tal modo un equilibrio in cui solo i soggetti ad alto rischio fanno domanda di assicurazione. Per gli individui sani si configurerebbe un fallimento del mercato: la presenza di asimmetria informativa fa venire meno un mercato (l'offerta di assicurazione integrale agli individui sani).

*Pooling equilibrium.* Potremmo porci la seguente domanda: esistono forme di contratti unici a copertura integrale in grado di soddisfare tutte le parti (vale a dire equilibri di tipo *pooling*)? Nel caso qui descritto la risposta è negativa. Già si è detto che tale condizione non può essere realizzata offrendo premi alti o bassi. *P* potrebbe essere indotto ad offrire un contratto con un premio unico medio che gli consenta di ottenere un equilibrio finanziario. Sulla misura di tale premio influiscono i diversi valori di *p* associati ai

soggetti a basso e ad alto rischio e l'ampiezza del primo gruppo rispetto al secondo, che dobbiamo supporre sia nota a *P*. Questa soluzione è però solo apparentemente accettabile per *P*. Come nel caso del premio alto citato sopra, anche in questo caso i soggetti a basso rischio non sarebbero disposti ad accettare un contratto con un premio superiore a quello equo per la propria categoria e quindi non sottoscriverebbero. Si assicurerebbero solo soggetti a rischio elevato. Alla fine *P* avrebbe perdite e sarebbe indotto a non offrire tale contratto. Una soluzione con un unico contratto per tutti non è quindi raggiungibile dal mercato.

*Separating equilibrium.* Una soluzione che consenta di separare i soggetti ad alto da quelli a basso rischio potrebbe essere quella di offrire, oltre a un contratto di assicurazione integrale con premio elevato, calcolato cioè sulla base della probabilità relativa ai soggetti ad alto rischio, anche un contratto di assicurazione *parziale* con un premio basso, calcolato sulla base della probabilità dei soggetti a basso rischio. In questo secondo contratto, se *L* è il danno che deriva ad *A* nel caso in cui si verifichi l'evento rischioso, *P* potrebbe offrire una copertura contrattuale pari a  $C < L$ . Anche se il premio è relativamente basso, i soggetti ad alto rischio potrebbero essere disincentivati a sottoscriverlo, in quanto, essendo il contratto a copertura parziale, dovrebbero comunque in qualche misura partecipare al rischio e la perdita attesa potrebbe essere superiore al costo di sottoscrivere un contratto a copertura integrale, sia pure pagando un premio più elevato. A seconda del valore dei parametri in gioco (un ruolo importante è svolto dal rapporto numerico tra soggetti ad alto e a basso rischio) può quindi realizzarsi un equilibrio in cui sono offerti due tipi di contratto, realizzando un *separating equilibrium*, cioè una situazione in cui spontaneamente i due tipi di soggetti scelgono contratti diversi. In tale situazione i soggetti ad alto rischio sceglieranno un contratto a copertura completa, con un premio più elevato di quello pagato dai soggetti a basso rischio, mentre questi ultimi sceglieranno un contratto a copertura parziale. Il contratto risulta così essere incentivo-compatibile, nel senso che nessuno dei due tipi di soggetti (nel nostro caso, in particolare, i malati) ha incentivo a «fingere» di appartenere all'altro tipo. *P* in ogni caso deve essere in grado di raggiungere un equilibrio finanziario e la situazione deve essere tale che non esistano ulteriori tipi di contratti che consentano ad un assicuratore di aumentare i profitti.

Un aspetto molto interessante di questa soluzione è che essa individua un meccanismo in cui l'offerta di assicurazione effettuata da *P* è tale da spingere i soggetti ad autoselezionarsi (*self-selection*) e quindi a rendere palese la categoria a cui appartengono attraverso il loro comportamento (sottoscrizione dell'uno o dell'altro contratto). In questo modo *P* otterrebbe l'informazione di cui difettesse.

La situazione descritta rappresenta un equilibrio di mercato, ma non configura una situazione Pareto ottimale per i soggetti a basso rischio, in quanto ad essi non è offerta effettivamente la possibilità di ottenere una copertura completa.

In conclusione, in presenza di *As* possono aversi o casi di inesistenza di equilibrio o equilibri non Pareto ottimali per i soggetti a basso rischio.

Il caso di As è molto importante nell'applicazione alla sanità, in quanto fornisce una spiegazione teorica, fondata esclusivamente su motivi di efficienza, dell'opportunità dell'intervento pubblico nell'offerta di assicurazione della salute. Anche se lo Stato non è in grado di eliminare alla radice il problema dell'asimmetria informativa, potrebbe tuttavia decidere di offrire gratuitamente a tutti i soggetti l'assicurazione contro la malattia, evitando in tal modo il fallimento derivante dall'assenza di un mercato. Perché tale soluzione (comunque diversa, in termini di benessere, rispetto a quella che si creerebbe in assenza di asimmetria informativa) possa essere Pareto ottimale, il servizio dovrebbe naturalmente essere finanziato con forme di imposizione non distorsive.

*Moral hazard.* Il *moral hazard* (espressione tradotta con «rischio morale» o, meglio, «comportamento sleale») può verificarsi in una situazione di asimmetria informativa in cui  $P$  non ha il controllo completo di un'azione di  $A$  connessa alla prestazione contrattuale (asimmetria informativa ad azione nascosta). Con riferimento ad esempio ad un'assicurazione contro l'incendio, il MH potrebbe essere rappresentato dalla sollecitudine che  $A$  mette nell'evitare che si verifichi l'evento assicurato. Si suppone che tale cura comporti qualche costo per  $A$  e che questo, una volta assicurato, sia tentato di non compiere tutti gli sforzi necessari per evitare l'evento dannoso, come invece fa chi non ha sottoscritto l'assicurazione. La presenza di questo incentivo fa sì che il comportamento di  $A$  (caratterizzato appunto da MH) influisca sulla probabilità del verificarsi dell'evento assicurato (l'incendio) e quindi alteri i termini del problema che in sua assenza influirebbero sulla determinazione del premio di assicurazione. È chiaro infatti che la proposta contrattuale che  $P$  può fare ad  $A$  (vale a dire quale somma assicurare in cambio di quale premio) è influenzata dall'ipotesi che  $A$ , dopo avere sottoscritto il contratto, agisca in modo leale o sleale.

Se  $P$  fosse in grado di controllare l'azione di  $A$  (ad esempio di verificare se  $A$  abbia installato un certo dispositivo antincendio, che ha l'effetto di modificare la probabilità  $p$ ),  $P$  potrebbe offrire due contratti differenziati: uno con un premio più basso a chi abbia installato il dispositivo antincendio, uno con un premio più alto per chi non l'abbia fatto. Ma la circostanza che  $P$  non sia in grado di controllare, senza incorrere in costi eccessivi, se  $A$  ha messo il dispositivo, fa sì che tale soluzione non sia possibile, perché  $A$ , sapendo che  $P$  non lo controlla, avrà sempre interesse a sottoscrivere il contratto con premio più basso e a non installare il dispositivo. Tale comportamento avrà l'effetto di causare perdite a  $P$ , dato che il premio è stato calcolato sull'ipotesi che la probabilità  $p$  fosse più bassa. Quindi alla lunga tale contratto non verrà offerto. L'assicuratore potrebbe avere interesse ad offrire solo il contratto con il premio più elevato. In tal caso però i soggetti con rischio più basso (o che hanno già installato dispositivi antincendio) non avranno interesse ad assicurarsi perché dovrebbero pagare un premio superiore a quello equo. La conseguenza è in ogni caso un'inefficienza allocativa, in definitiva attribuibile alla presenza di un insanabile conflitto in cui si trova  $A$  tra il suo desiderio di assicurarsi e l'incentivo a tenere comportamenti di MH.

$P$  potrebbe però anche proporre altri tipi di contratti. Uno di questi potrebbe consentire ad  $A$  di assicurare solo una parte della possibile perdita (e ciò è quanto in verità avviene nella stragrande maggioranza dei casi concreti). Se il danno derivante dall'incendio è  $L$ ,  $P$  potrebbe consentire che venga assicurata solo una somma  $C < L$ , offrire cioè un contratto di assicurazione parziale. Di conseguenza  $A$ , in analogia a quanto visto nel caso di As, risulterebbe partecipe del rischio. Il fatto che  $A$  concorra al rischio consente di rendere non ineluttabile l'incentivo al comportamento di MH. Con MH,  $A$  infatti aumenta la probabilità del danno e quindi il suo costo di partecipazione allo stesso. Per determinati valori dei parametri in gioco, può accadere che un contratto di assicurazione parziale risolva il problema di  $P$ , posto che l'utilità attesa di  $A$  derivante dal sopportare i costi relativi a comportamenti leali (dispositivo antincendio, cura, ecc.) sia superiore a quella ottenibile in caso opposto. Tale contratto potrebbe essere sottoscritto da  $A$ , in quanto sarebbe comunque preferibile alla situazione di non assicurazione. La teoria dell'informazione ci permette così di individuare soluzioni in cui può essere incentivato un comportamento non opportunistico da parte dell'agente. Tale soluzione non è però Pareto ottimale, di First Best, perché  $A$  in qualche misura deve partecipare al rischio e, come abbiamo sopra illustrato, una situazione di First Best implica che i soggetti con avversione al rischio ottengano una copertura completa.

In conclusione in caso di MH il mercato fallisce o perché non è in grado di realizzare un contratto di assicurazione accettabile per entrambe le parti, o perché è in grado di determinare solo tipi di contratti che comunque non sono Pareto ottimali in quanto impediscono una completa copertura per i soggetti con avversione al rischio.