

Mechanism Design: How to Implement Social Goals*

Eric S. Maskin**

February 2008

* This article is a revised version of Eric Maskin's Nobel Memorial Prize Lecture delivered on December 8, 2007 in Stockholm

** Albert O. Hirschman Professor of Social Science, Institute for Advanced Study, and Visiting Lecturer, Princeton University. Research support from NSF grant #SES-0318103 is gratefully acknowledged.

The theory of mechanism design can be thought of as the “engineering” side of economic theory. Much theoretical work, of course, focuses on *existing* economic institutions. The theorist wants to explain or forecast the economic or social outcomes that these institutions generate. But in mechanism design theory the direction of inquiry is reversed. We begin by identifying our desired outcome or *social goal*. We then ask whether or not an appropriate institution (mechanism) could be designed to attain that goal. If the answer is yes, then we want to know what form that mechanism might take.

In this paper, I offer a brief introduction to the part of mechanism design called *implementation theory*, which, given a social goal, characterizes when we can design a mechanism whose *predicted* outcomes (i.e., the set of equilibrium outcomes) coincide with the *desirable* outcomes, according to that goal. I try to keep technicalities to a minimum, and usually confine them to footnotes.¹

I. Outcomes, Goals, and Mechanisms

What we mean by an “outcome” will naturally depend on the context. Thus, for a government charged with delivering public goods, an outcome will consist of the quantities provided of such goods as intercity highways, national defense and security, environmental protection, and public education, together with the arrangements by which they are financed. For an electorate seeking to fill a political office, an outcome is simply the choice of a candidate for that office. For an auctioneer selling a collection of assets,

¹ There are many excellent surveys and textbook treatments of implementation theory that go into considerably more detail—both technical and conceptual—than I do here; see in particular: Andrew Postlewaite (1985), Theodore Groves and John Ledyard (1987), John Moore (1992), Thomas Palfrey (1992), chapter 10 of Martin Osborne and Ariel Rubinstein (1994), Beth Allen (1997), Luis Corchon (1996), Matthew Jackson (2001), Palfrey (2002), Roberto Serrano (2004), chapters 2 and 3 of David Austen-Smith and Jeffrey Banks (2005), chapter 6 of James Bergin (2005), chapters 14-16 of Allan Feldman and R. Serrano (2006), chapter 10 of Eric Rasmusen (2006), Sandeep Baliga and Tomas Sjöström (2007) and Corchon (2008). See also Partha Dasgupta, Peter Hammond, and Eric Maskin (1979), Maskin and Sjöström (2002), Baliga and Maskin (2003), and my old survey Maskin (1987).

an outcome corresponds to an allocation of these assets across potential buyers, together with the payments that these buyers make. Finally, in the case of a home buyer and a builder contemplating the construction of a new house, an outcome is a specification of the house's characteristics and the builder's remuneration.

Similarly, the standards by which we judge the “desirability” or “optimality” of an outcome will also depend on the setting. In evaluating public good choices, the criterion of “net social surplus” maximization is often invoked: does the public good decision maximize gross social benefit minus the cost of providing the goods? As for electing politicians, the property that a candidate would beat each competitor in head-to-head competition (i.e., would emerge a *Condorcet winner*) is sometimes viewed as a natural desideratum (see Partha Dasgupta and Eric Maskin, forthcoming). In the auctioning of assets, there are two different criteria by which an outcome is typically judged: (i) whether the assets are put into the hands of bidders who value them the most (i.e., whether the allocation is *efficient*); and alternatively (ii) whether the seller raises the greatest possible revenue from sales (i.e., whether *revenue maximization* is achieved). Finally, for the home buyer and builder, an outcome will ordinarily be considered “optimal” if it exhausts the potential gains from exchange between the parties, i.e., the house specification and remuneration are together Pareto optimal and individually rational.

A *mechanism* is an institution, procedure, or game for determining outcomes. Not surprisingly, who gets to choose the mechanism—i.e., who is the mechanism designer—will, once again, depend on the setting. In the case of public goods, we normally think of the government providing the goods as also choosing the method by which the levels of

provision and financing are determined. Similarly, when it comes to sales of assets—where an *auction* is the typical mechanism—the asset seller often gets to call the shots about the rules, i.e., he is the one who chooses the auction format.

In the case of national political elections, by contrast, a mechanism is an *electoral procedure*, e.g., plurality rule, run-off voting, or the like. Moreover, the procedure is ordinarily prescribed long in advance, indeed sometimes by the country's constitution. Thus, here we should think of the framers of the constitution as the mechanism designers.

Finally, in the house-building example, a mechanism is a contract between the home buyer and builder and lays out the rights and responsibilities of each. Since these parties are presumably the ones who negotiate this contract, they themselves are the mechanism designers in this last setting.

Now, in the public framework, if the government knows at the outset which choice of public goods is optimal, then there is a simple—indeed, trivial—mechanism for achieving the optimum: the government has only to pass a law mandating this outcome. Similarly, if the auctioneer has prior knowledge of which bidders value the assets most, he can simply award them directly to those bidders (with or without payment).

The basic difficulty—which gives the subject of mechanism design its theoretical interest—is that the government or auctioneer will typically *not* have this information. After all, the net surplus-maximizing choice of public goods depends on citizens' *preferences* for such goods, and there is no particular reason why the government should know these preferences. Likewise, we wouldn't normally expect an auctioneer to know how much different potential buyers value the assets being sold.

Because mechanism designers do not generally know which outcomes are optimal in advance, they have to proceed more indirectly than simply prescribing outcomes by fiat; in particular, the mechanisms designed must generate the information needed as they are executed. The problem is exacerbated by the fact that the individuals who *do* have this critical information—the citizens in the public good case or the buyers in the asset-selling example—have their own objectives and so may not have the incentive to behave in a way that reveals what they know. Thus, the mechanisms must be *incentive compatible*. Much of the work in mechanism design, including my own, has been directed at answering three basic questions:

(A) When is it possible to design incentive-compatible mechanisms for attaining social goals?

(B) What form might these mechanisms take when they exist?

and

(C) When is finding such mechanisms ruled out theoretically?

That is ever possible to design such mechanisms may, at first, seem surprising. How, after all, can a mechanism designer attain an optimal outcome without knowing exactly what he is aiming for? Thus, it may be helpful to consider a simple concrete example.

II. An Example

Consider a society consisting of two consumers of energy, Alice and Bob. An energy authority is charged with choosing the type of energy to be used by Alice and Bob. The list of options—from which the authority must make a single selection—are gas, oil, nuclear power, and coal.

Let us suppose that there are two possible states of the world. In state 1, the consumers place relatively little weight on the future, i.e., they have comparatively high temporal discount rates. In state 2, by contrast, they attach a great deal of importance to the future, meaning that their rates of discount are correspondingly low.

Alice, we will imagine, cares primarily about convenience when it comes to energy. This means that, in state 1, she will rank gas over oil, oil over coal, and coal over nuclear power, because as we move down her ranking, the energy source becomes either messier or more cumbersome to use. In state 2, by contrast, her ranking is

nuclear
gas
coal
oil

because she anticipates that technical advances will eventually make gas, coal, and especially nuclear power much easier to use - - - and, in this state, she lays particular stress on *future* benefits.

Bob is interested particularly in *safety*. This implies that in state 1, when he puts greatest weight on the present, he favors nuclear power over oil, oil over coal, and coal over gas. But if state 2 obtains—so that the future is comparatively important—his ranking is:

oil
gas
coal
nuclear

which reflects the fact that, in the long-run, the problem of disposing of nuclear waste can be expected to loom large, but that oil and gas safety are likely to improve somewhat.

To summarize, the consumers' rankings in the two states are given in Table 1.

State 1		State 2	
<u>Alice</u>	<u>Bob</u>	<u>Alice</u>	<u>Bob</u>
gas	nuclear	nuclear	oil
oil	oil	gas	gas
coal	coal	coal	coal
nuclear	gas	oil	nuclear

Table 1

Assume that the energy authority is interested in selecting an energy source that both consumers are reasonably happy with. If we interpret “reasonably happy” as getting one’s first or second choice, then oil is the optimal choice in state 1, whereas gas is the best outcome in state 2. In the language of implementation theory, we say that the authority’s *social choice rule* prescribes oil in state 1 and gas in state 2. Thus, if f is the social choice rule, it is given by Table 2.²

$$f(\text{state 1}) = \text{oil} \qquad f(\text{state 2}) = \text{gas}$$

Table 2

Suppose, however, that the authority does not know the state (although Alice and Bob do). This means that it does not know which alternative the social choice rule prescribes, i.e., whether oil or gas is the optimum.

Probably the most straightforward mechanism would be for the authority to ask each consumer to announce the state, whereupon it would choose oil if both consumers

² In a more general setting, where Θ is the set of possible states of the world and A is the set of possible outcomes, a social choice rule f is a correspondence (a set-valued function) $f : \Theta \rightarrow \rightarrow A$, where, for any θ , $f(\theta)$ is interpreted as the set of optimal outcomes in state θ (we are allowing for the possibility that more than one outcome might be considered optimal in a given state).

said “state 1,” choose gas if both said “state 2,” and flip a coin between them if it got a mixed response. But notice that in this mechanism Alice has the incentive to say “state 2” regardless of the actual state and regardless of what Bob says, because she prefers gas to oil in both states. Indeed, by saying “state 2” rather than “state 1,” she raises the probability of her preferred outcome from 0 to .5 if Bob says “state 1,” and from .5 to 1 if Bob says “state 2.” Hence, we would expect Alice to report “state 2” in both states. Similarly, Bob would always report “state 1,” because he prefers oil to gas in either state. Taken together, Alice’s and Bob’s behavior implies that, in each state, the outcome is a 50-50 randomization between oil and gas. That is, there is only a 50% chance that the outcome is optimal, and so this mechanism is demonstrably too naïve.

Let us suppose, therefore, that the authority has the consumers participate in the mechanism given by Table 3:

		Bob	
		Left	Right
Alice	Top	oil	coal
	Bottom	nuclear	gas

Table 3

That is, Alice chooses “Top” or “Bottom” as her strategy; simultaneously, Bob chooses “Left” or “Right” as his strategy; and the outcome of those choices is given in the corresponding entry of the matrix.³

³ More generally, a mechanism for a society with n individuals is a mapping $g : S_1 \times \cdots \times S_n \rightarrow A$, where, for all i , S_i is individual i ’s strategy space and $g(s_1, \dots, s_n)$ is the outcome prescribed by the mechanism if individuals play the strategies $(s_1, \dots, s_n)$.

Observe that, in state 1, Bob is better off choosing Left regardless of what Alice does: if she plays Top, then Left leads to oil as the outcome (which Bob prefers), whereas Right gives rise to coal. If she plays Bottom, then nuclear power (Bob's preferred outcome) is the consequence of going Left, while Right leads to gas. That is, Left is the “dominant strategy” for Bob in state 1. Moreover, given that Bob is going Left, Alice is better off choosing Top rather than Bottom, because she prefers oil to nuclear power. Thus, in state 1, the clear prediction is for Alice to play Top and for Bob to play Left, i.e., (Top, Left) is the unique Nash equilibrium.⁴ Furthermore—and this is the critical point—the resulting outcome, oil, is optimal in state 1.

Turning to state 2, we see that Bottom is the dominant strategy for Alice in that state. If Bob plays Left, then she is better off with Bottom than Top because she prefers nuclear power to oil. And if Bob goes Right, then Bottom leads to gas, which she prefers to the Top outcome, coal. With Alice choosing Bottom, Bob is better off going Right, because gas is better for him than nuclear power. Hence, in state 2, the (unique) Nash equilibrium is (Bottom, Right): Alice plays Bottom and Bob goes Right. Furthermore, this results in the optimal outcome, gas.

We have seen that in either state, the mechanism of Table 3 achieves the optimal outcome even though (i) the mechanism designer (the energy authority) does not even know the actual state, and (ii) Alice and Bob are interested only in their own preferences, not those of the authority. More precisely, because the Nash equilibrium outcomes of the

⁴ In general, a Nash equilibrium is a specification of strategies—one for each individual—from which no individual has the incentive to deviate unilaterally. Thus, if $u_i(a, \theta)$ is individual i 's payoff from outcome a in state θ , strategies $(s_1, \dots, s_n)$ constitute a Nash equilibrium of mechanism g in state θ if $u_i(g(s_1, \dots, s_i, \dots, s_n), \theta) \geq u_i(g(s_1, \dots, s'_i, \dots, s_n), \theta)$ for all i and all $s'_i \in S_i$.

Table 3 mechanism coincide with the optimal outcomes in each state, we say that the mechanism *implements* the authority's social choice rule in Nash equilibrium.^{5,6}

III A Brief History of Mechanism Design

The intellectual history of mechanism design theory goes back at least to nineteenth-century utopian socialists such as Robert Owen and Charles Fourier. Repulsed by what they viewed as the evils of the burgeoning capitalist system, these thinkers argued that socialism offered a more humane alternative and sometimes became involved in setting up experimental communities such as New Harmony, Indiana.

A more direct influence on the modern theory was the Planning Controversy, which reached its greatest intensity in the 1930s. The principal antagonists on one side were Oskar Lange and Abba Lerner, who argued forcefully that, done right, central planning could replicate the performance of free markets (Lange 1936 and Lerner 1944). Indeed, they suggested, planning could correct serious “market failures”—notably those on display in the Great Depression—and thereby potentially surpass markets. On the other side, Friedrich von Hayek and Ludwig von Mises staunchly denied the possibility that a planned system could ever approach the success of the free market (von Hayek 1944 and von Mises 1920).

The controversy was important and fascinating, but for certain onlookers such as Leonid Hurwicz, it was also rather frustrating. This was because it lacked conceptual

⁵ In a more general setting, mechanism g implements social choice rule f in Nash equilibrium if

$$f(\theta) = NE_g(\theta) \quad \text{for all } \theta,$$

where $NE_g(\theta)$ is the set of Nash equilibrium outcomes of g in state θ .

⁶ Nash equilibrium is a prediction of how individuals in a mechanism will behave. But a number of other predictive concepts—i.e., equilibrium concepts—have been considered in the implementation literature, among them subgame perfect equilibrium (Moore and Rafael Repullo 1988), undominated Nash equilibrium (Palfrey and Sanjay Srivastava 1991), Bayesian equilibrium (Postlewaite and David Schmeidler 1986), dominance solvability (Hervé Moulin 1979), trembling-hand perfect equilibrium (Sjöström 1993), and strong equilibrium (Bhaskar Dutta and Arunava Sen 1991)

precision: critical terms such as “decentralization” were left undefined. Moreover, the arguments adduced on either side often were often highly incomplete. In part, this was because they simply lacked the technical apparatus—in particular, game theory and mathematical programming—to generate truly persuasive conclusions.

This is where Leo Hurwicz entered the picture. Inspired by the debate, he attempted to provide unambiguous definitions of the central concepts, and this effort culminated in his two great papers, Hurwicz (1960) and (1972), where he also introduced the critical notion of incentive compatibility.

The work inspired by Hurwicz and others has produced a broad consensus among economists that von Hayek and von Mises were, in fact, correct—the market *is* the “best” mechanism—in settings where (i) there are large numbers of buyers and sellers, so that no single agent has significant market power; and (ii) there are no significant externalities, that is, an agent’s consumption, production, and information does not affect others’ production or consumption.⁷ However, mechanisms improving the market are generally possible if either assumption is violated.⁸

Hurwicz’s work gave rise to an enormous literature, which has largely branched in two different directions. On the one hand, there is work that makes use of special, highly structured settings to study particular questions such as how to allocate public goods, how to design auctions, and how to structure contracts. On the other hand, there are studies obtaining results at a general, abstract level, that is, they make as few

⁷ See, for example, Peter Hammond (1979)—who shows, roughly, that the competitive market is the only incentive compatible-mechanism producing individually rational and Pareto optimal outcomes— and James Jordan (1982)—who shows the same thing when “incentive compatible” is replaced by “information efficient”, under assumptions (i) and (ii).

⁸ See, for instance, Theodore Groves (1973) and Edward Clarke (1971) for the case of public goods and Jean-Jacques Laffont (1985) for the case of informational externalities.

assumptions as possible about preferences, technologies, and so on. My own work has fallen into both categories at different times. But, in this paper, I will emphasize general results.

IV. Implementation of Social Choice Rules

Above I set out three central questions (A) – (C) about incentive-compatible mechanisms. Rephrased in the language of implementation theory these questions become:

(A') Under what conditions can a social choice rule be implemented?

(B') What form does an implementing mechanism take?

(C') Which social choice rules cannot be implemented?

In the mid-1970s I struggled with these questions. Eventually, I discovered that a property called *monotonicity* (now sometimes called Maskin-monotonicity) is the key to implementability in Nash equilibrium. Suppose that outcome a is optimal in state θ according to the social choice rule f in question, that is, $f(\theta) = a$. Then, if a doesn't fall in anyone's ranking relative to any other alternative in going from state θ to state θ' , monotonicity requires that a also be optimal in state θ' : $f(\theta') = a$. However, if a does fall relative to some outcome b in someone's ranking, monotonicity imposes no restriction.⁹

To see what monotonicity means more concretely, let's consider our energy example from before (see Tables 1 and 2). Recall that oil is the optimal outcome in state

⁹ In a more general setting in which f can be set-valued, monotonicity requires that, for all states θ, θ' and all outcomes a , if $a \in f(\theta)$ and $u_i(a, \theta) \geq u_i(b, \theta)$ implies $u_i(a, \theta') \geq u_i(b, \theta')$ for all i and b , then $a \in f(\theta')$.

1. Notice too that oil *falls* in Alice’s ranking, relative to both coal and nuclear power, in going from state 1 to state 2 (Alice ranks oil higher than coal and nuclear in state 1, but just the opposite is true in state 2). Thus, the fact that gas—not oil—is optimal in state 2 does not violate monotonicity. Similarly, observe that gas falls in Bob’s ranking, relative to both coal and nuclear power, in going from state 2 to state 1. Hence, even though gas is optimal in state 2, the fact that it is not optimal in state 1 is also not in conflict with monotonicity. Indeed, these verifications establish that the authority’s social choice rule satisfies monotonicity (and thus the possibility of implementing it, which was shown earlier, does not contradict Theorem 1 below).

But suppose we modify the example somewhat, so that rankings and optimal outcomes are given by Table 4. With these changes, the social choice

State 1		State 2	
<u>Alice</u>	<u>Bob</u>	<u>Alice</u>	<u>Bob</u>
gas	nuclear	gas	nuclear
oil	oil	oil	oil
coal	coal	nuclear	coal
nuclear	gas	coal	gas
oil optimal		nuclear optimal	

Table 4

rule is no longer monotonic. Specifically, observe that although oil is optimal in state 1, it is not optimal in state 2, despite the fact that it falls in neither Alice’s nor Bob’s rankings between states 1 and 2 (given that oil doesn’t fall, monotonicity would require it to remain optimal in state 2). Hence, we can conclude that there is *no* mechanism that implements the social choice rule of Table 4. More generally, we have:

Theorem 1 (Maskin 1977): *If a social choice rule is implementable, then it must be monotonic.*

To see why the social choice rule in Table 4 is not implementable, suppose to the contrary that there *were* an implementing mechanism. Then, in particular, the mechanism would necessarily contain a pair of strategies (s_A, s_B) —for Alice and Bob, respectively—that result in outcome oil and constitute a Nash equilibrium in state 1.

I claim that (s_A, s_B) must also constitute a Nash equilibrium in state 2. To understand this claim, note first that Bob has no incentive to deviate unilaterally from s_B in state 2, since (i) he has no such incentive in state 1 (by definition of Nash equilibrium) and (ii) his preference ranking is the same in both states. Furthermore, Ann has no incentive to deviate from s_A in state 2. To see this, observe that if, contrary to the claim, Ann gained from deviating unilaterally from s_A in state 2, she must thereby be inducing the outcome gas (because this is the only outcome she prefers to oil in state 2). But Alice also prefers gas to oil in state 1, and so would benefit from the same deviation in that state, contradicting the assumption that (s_A, s_B) constitutes a Nash equilibrium in state 1.

Hence, (s_A, s_B) is indeed a Nash equilibrium in state 2. But the outcome it generates—oil—is not optimal in that state, establishing that the social choice rule is not implementable after all.

As we have seen, Tables 1 and 2 provide an example of a social choice rule that is monotonic and also implementable. However, it is not true that *all* monotonic social choice rules are implementable; see Maskin (1977) for a counterexample. Nevertheless, such counterexamples are rather contrived, and if an additional, often innocuous

condition is imposed, monotonicity *does* guarantee implementability, if there are at least three individuals in society.¹⁰

The additional condition is called *no veto power*. Suppose that all individuals, except possibly one, agree that a particular outcome *a* is *best*, meaning that they all put *a* at the *top* of their preference rankings. Then, if the social choice rule satisfies no veto power, *a* must be optimal. In other words, the remaining individual cannot “veto” it.

No veto power is especially innocuous—indeed, it imposes no restriction at all—when outcomes entail a distribution of economic goods across individuals. In that case, each individual will prefer a bigger share of those goods for himself or herself. So, no two of them can agree that a given outcome *a* is best: they cannot both get the biggest share. This means that, if there are three or more individuals, the hypothesis posited by the no veto power condition *cannot be satisfied*, and so logically the condition holds *automatically*.

A general result on the possibility of implementing social choice rules is the following:

Theorem 2 (Maskin 1977): *Suppose that there are at least three individuals. If the social choice rule satisfies monotonicity and no veto power, then it is implementable.*

Proofs of Theorem 2 are beyond the scope of this paper (see Repullo 1987 for an especially elegant argument), but I should mention that they are usually *constructive*.

That is, given the social choice rule to be implemented, a proof lays out an explicit recipe for the construction of a mechanism that does the trick.

¹⁰ That is not to say that implementation is impossible with just two individuals - - indeed, our energy example of Tables 1 and 2 had only two individuals. However, as we will see below, implementation is facilitated by there being three or more individuals.

It is worth pointing out why Theorem 2 posits at least three individuals. Often in economics, moving from two to three persons makes things more difficult.¹¹ But, for implementation theory, three individuals actually make matters easier. To understand why, remember that the underlying idea of a mechanism is to give individuals the incentive to behave in a way that ensures an optimal outcome. This entails “punishing” an individual for deviating from his prescribed (i.e., equilibrium) strategy. But if there are only two individuals, Alice and Bob, and one of them has deviated, it may be difficult to determine whether it was Alice who deviated and Bob who complied or vice versa. This problem of identification is resolved once there are three people: a deviator sticks out more obviously when two or more other individuals are complying with equilibrium.

VI Concluding Remarks

This has been only a very brief introduction to implementation theory (which itself constitutes only part of the field of mechanism design). I have concentrated on work that was done over thirty years ago, which perhaps gives a misleadingly “antique” flavor to the paper. In fact, an especially gratifying aspect of the theory is that almost fifty years after Hurwicz (1960), the subject remains intellectually vibrant and important: new implementation papers appearing all the time. It will be interesting to see where the field goes in the next fifty years.

¹¹ Zero-sum games provide a classic example of this phenomenon. The minimax theorem—which greatly simplifies the analysis of behavior in games—applies to *two*-person zero-sum games, but not, in general, to the case of *three* or more players.

References

- Allen, Beth (1997), "Implementation Theory with Incomplete Information," in S. Hart and A. Mas-Colell, (eds.), *Cooperation: Game Theoretic Approaches*, Berlin: Springer.
- Austen-Smith, David and Jeffrey Banks (2005), *Positive Political Theory II*, Ann Arbor: University of Michigan Press.
- Baliga, Sandeep and Eric Maskin (2003), "Mechanism Design for the Environment," in K.G. Mäler and J. Vincent, eds., *Handbook of Environmental Economics*, Vol. 1, Amsterdam: North-Holland, pp. 305-324.
- Baliga, S. and Tomas Sjöström (2007), "Mechanism Design: Recent Developments," L. Blume and S. Durlauf (eds.), *The New Palgrave Dictionary of Economics*, 2nd Edition, London:McMillan.
- Bergin, James (2005), *Microeconomic Theory*, Oxford: Oxford University Press.
- Clarke, Edward (1971), "Multipart Pricing of Public Goods," *Public Choice*, pp. 19-33.
- Corchon, Luis (1996), *The Theory of Implementation of Socially Optimal Decisions in Economics*, London: Macmillan.
- Corchon, L. (2008), "The Theory of Implementation," *The Encyclopedia of Complexity and System Science*, Berlin: Springer.
- Dasgupta, Partha, Peter Hammond, and E. Maskin (1979), "The Implementation of Social Choice Rules: Some General Results on Incentive Compatibility," *Review of Economic Studies*, 46, pp. 185-216.
- Dasgupta, P. and E. Maskin (forthcoming), "On the Robustness of Majority Rule", *Journal of the European Economic Association*.
- Dutta, Bhaskar and Arunava Sen (1991), "Implementation under Strong Equilibrium: A Complete Characterization," *Journal of Mathematical Economics*, 20, pp. 46-67.
- Feldman, Allan and Roberto Serrano (2006), *Welfare Economics and Social Choice Theory*, Berlin: Springer.
- Groves, Theodore (1973), "Incentives in Teams," *Econometrica*, 41, pp. 617-631.
- Groves, T. and John Ledyard (1987), "Incentive Compatibility since 1972," in T. Groves, R. Radner, and S. Reiter (eds.), *Information, Incentives and Economic Mechanisms*, Minneapolis: University of Minnesota Press, pp. 48-111.

Jordan, James (1982), "The Competitive Allocation Process is Informationally Efficient Uniquely," *Journal of Economic Theory*, 28, pp. 1-18.

Hammond, Peter (1979), "Straightforward Individual Incentive Compatibility in Large Economies," *Review of Economic Studies*, 46, pp. 263-282.

Hurwicz, Leonid (1960), "Optimality and Informational Efficiency in Resource Allocation Processes" in Kenneth Arrow, S. Karlin and P. Suppes, (eds.), *Mathematical Methods in Social Sciences*, Stanford: Stanford University Press, pp. 27-46.

Hurwicz, L. (1972), "On Informationally Decentralized Systems," in C. McGuire, and R. Radner, (eds.), *Decision and Organization*, Amsterdam: North-Holland, pp. 297-336.

Jackson, Matthew (2001), "A Crash Course in Implementation Theory," *Social Choice and Welfare*, 18, pp. 655-708.

Laffont, Jean-Jacques (1985), "On the Welfare Analysis of Rational Expectations Equilibria with Asymmetric Information," *Econometrica*, 53, pp. 1-29.

Lange, Oskar (1936), "On the Economic Theory of Socialism," *Review of Economic Studies*, 4, pp. 53-71.

Lerner, Abba (1944), *The Economics of Control*, New York: McMillan.

Maskin, Eric (1977, published 1999), "Nash Equilibrium and Welfare Optimality," *Review of Economic Studies*, pp. 23-38.

Maskin, E. (1985), "The Theory of Implementation in Nash Equilibrium: A Survey," in L. Hurwicz, D. Schmeidler, and H. Sonnenschein (eds.), *Social Goals and Social Organization*, Cambridge: Cambridge University Press.

Maskin, E. and T. Sjöström (2002), "Implementation Theory," in K. Arrow, A. Sen, and K. Suzumura, (eds.), *Handbook of Social Choice and Welfare*, Vol. I, Amsterdam: Elsevier, pp. 237-288.

Moore, John (1992), "Implementation, Contracts, and Renegotiation in Environments with Complete Information," in J.J. Laffont (ed.), *Advances in Economic Theory*, Vol. 1, Cambridge: Cambridge University Press, pp. 182-282.

Moore, J. and Rafael Repullo (1988), "Subgame Perfect Implementation," *Econometrica*, 56, pp. 1191-1220.

Moulin, Hervé (1979), "Dominance Solvable Voting Schemes," *Econometrica*, 47, pp. 1337-1351.

Osborne, Martin and Ariel Rubinstein (1994), *A Course in Game Theory*, Cambridge: MIT Press.

Palfrey, Thomas (1992), "Implementation in Bayesian Equilibrium: The Multiple Equilibrium Problem in Mechanism Design," in J.J. Laffont (ed.), *Advances in Economic Theory*, Vol. 1, Cambridge: Cambridge University Press, pp. 283-323.

Palfrey T. (2001), "Implementation Theory," in R. Aumann and S. Hart, (eds.), *Handbook of Game Theory*, vol. 3, Amsterdam: North-Holland, pp. 2271-2326.

Palfrey, T. and Sanjay Srivastava (1991), "Nash Implementation using Undominated Strategies," *Econometrica*, 59, pp. 479-501.

Postlewaite, Andrew (1985), "Implementation via Nash Equilibria in Economic Environments," in L. Hurwicz, D. Schmeidler, and H. Sonnenschein (eds.) *Social Goals and Social Organization*, Cambridge: Cambridge University Press, pp. 205-228

Postlewaite, A. and David Schmeidler (1986), "Implementation in Differential Information Economies," *Journal of Economic Theory*, 39, pp. 14-33.

Rasmusen, Eric (2006), *Games and Information: An Introduction to Game Theory*, Oxford: Blackwell Publishing.

Repullo, Rafael (1987), "A Simple Proof of Maskin's Theorem on Nash Implementation," *Social Choice and Welfare*, 4, pp. 39-41.

Serrano, Roberto (2004), "The Theory of Implementation of Social Choice Rules," *SIAM Review*, 46, pp. 377-414.

Sjöström, Tomas (1993), "Implementation in Perfect Equilibria," *Social Choice and Welfare*, 10, pp. 97-106.

von Hayek, Friedrich (1944), *The Road to Serfdom*, London: Routledge.

von Mises, Ludwig (1935), "Die Wirtschaftsrechnung im Sozialistischen Gemeinwesen," in F. von Hayek (ed.), *Collectivist Economic Planning*, London: Routledge.