

ESTIMATING DYNAMIC MODELS OF IMPERFECT COMPETITION

BY PATRICK BAJARI, C. LANIER BENKARD, AND JONATHAN LEVIN¹

We describe a two-step algorithm for estimating dynamic games under the assumption that behavior is consistent with Markov perfect equilibrium. In the first step, the policy functions and the law of motion for the state variables are estimated. In the second step, the remaining structural parameters are estimated using the optimality conditions for equilibrium. The second step estimator is a simple simulated minimum distance estimator. The algorithm applies to a broad class of models, including industry competition models with both discrete and continuous controls such as the Ericson and Pakes (1995) model. We test the algorithm on a class of dynamic discrete choice models with normally distributed errors and a class of dynamic oligopoly models similar to that of Pakes and McGuire (1994).

KEYWORDS: Markov perfect equilibrium, dynamic games, incomplete models, bounds estimation.

1. INTRODUCTION

IN MANY BRANCHES OF APPLIED ECONOMICS, it has become common practice to estimate structural models of decision-making and equilibrium. With a few notable exceptions, most of this work has focused on static environments or on single-agent dynamic decision problems. Many economic policy debates, however, turn on quantities that are inherently linked to dynamic competition, such as entry and exit costs, the returns to advertising or research and development, the adjustment costs of investment, or the speed of firm and consumer learning. Estimating these dynamic parameters has been seen as a major challenge, both conceptually and computationally.

One reason for this is the perceived difficulty of incorporating information from a dynamic equilibrium into an estimation algorithm. Research on dynamic competition (e.g., Ericson and Pakes (1995), Pakes and McGuire (1994, 2001), Gowrisankaran and Town (1997), and Benkard (2004)) has shown that computing an equilibrium for even relatively simple industry models is all but prohibitive. For models with the complexity usually required for empirical work, the situation is even bleaker. Even with advancing computer technology, computing equilibria over and over, as would be required in a typical estimation routine, seems out of the question. Moreover, dynamic games often admit a vast multiplicity of equilibria. This multiplicity greatly complicates the application of estimators that require computing equilibria and then matching these equilibria to observed data.

¹We thank the co-editor and three anonymous referees for detailed and constructive suggestions. The paper has benefited from our conversations with Jeremy Fox, Phil Haile, Igal Hendel, Guido Imbens, Phillip Leslie, Ariel Pakes, Peter Reiss, Azeem Shaikh, Elie Tamer, and Ed Vytlačil. Matthew Osborne provided exemplary research assistance. We thank the Bureau of Economic Analysis and the National Science Foundation for financial support.

This paper develops a method for estimating dynamic models of imperfect competition that is straightforward to apply and does not require the ability to compute an equilibrium even once. The approach involves two steps. The first step is to recover the agents' policy functions, as well as the probability distributions determining the evolution of the relevant state variables. This essentially involves regressing observed actions (such as investment, quantity, price, entry, or exit) on observed state variables (such as demand or cost shifters, and firm and product characteristics). In an equilibrium model, agents have correct beliefs about their environment and the behavior of other agents. As a consequence, by estimating the probability distributions for actions and states, one effectively recovers the agents' equilibrium beliefs.²

The second step is to find a set of structural parameters that rationalize the observed policies as a set of optimal decisions or, more precisely, as a set of mutual best responses. This parallels the second feature of equilibrium models, namely that agents maximize expected discounted profits given their beliefs. We represent the conditions for optimality as a system of inequalities that require each agent's observed behavior at each state be weakly preferred to the feasible alternatives. The model's parameters are estimated as the solution to this system of inequalities. In practice, we apply a simple simulated minimum distance estimator that minimizes violations of the optimality conditions.

Our approach builds on a line of research, initiated by Rust (1987), on the estimation of single-agent dynamic discrete choice models.³ Rust showed that dynamic single-agent models could be estimated using a nested algorithm. His idea was to solve the agent's dynamic decision problem for candidate parameter values and search for the value that yields predictions most closely aligned with the data. Hotz and Miller (1993) suggested a computationally convenient two-step alternative. Their estimator exploits the mapping in dynamic discrete choice problems between conditional choice probabilities and "choice-specific" value functions. Hotz and Miller estimated the probability distribution over choices at different states and used this to recover the agent's value function. The estimated value function is then an input to a second-stage estimate of the structural parameters, just as in our estimator.

Hotz and Miller's approach has two features that make it a natural building block for estimating dynamic models of strategic interaction. The first is computational. Extending Rust's method to games requires one to compute the full set of equilibria for candidate parameter values, when finding even a single equilibrium may be costly and difficult. In contrast, the two-step approach does not require equilibrium computations. The two-step approach also somewhat

²Ching (2005) used a similar approach to recover consumer's beliefs about future prices in a model of dynamic demand.

³Geweke and Keane (2001) and Imai, Jain, and Ching (2005) provided two additional approaches that may also be useful for estimating equilibrium models, but we have not investigated them.

mitigates problems caused by multiple equilibria. As long as the data are generated by a single equilibrium, the first-stage estimation recovers the correct value functions for that equilibrium. Therefore, second-stage parameter estimates will be consistent even if the estimated parameters would support other equilibria that are not observed in the data.

Recent work by Pakes, Ostrovsky, and Berry (2007), Pesendorfer and Schmidt-Dengler (2003), and Aguirregabiria and Mira (2007) develops alternative extensions to the Hotz–Miller approach that can be used to estimate dynamic games where actions have a discrete choice structure. Their estimators and ours are conceptually similar, but differ in the specifics of implementation.⁴ Pesendorfer and Schmidt–Dengler’s work is perhaps the most direct application of Hotz and Miller’s estimator; they also provide identification results for dynamic games. Pakes, Ostrovsky, and Berry build their second-stage estimation around pooled moment conditions rather than maximum likelihood. We discuss this idea at length in Section 6.2.3. Aguirregabiria and Mira argue for an iterative application of the Hotz and Miller two-step estimator, in which one uses the first round of estimates to generate new conditional choice probabilities and then reruns the estimator. They show that iteration, if it converges, has attractive properties, something they earlier established for single-agent models (Aguirregabiria and Mira (2002)). A potential benefit of iteration, which we discuss in Section 5, is that it allows one to include serially correlated unobserved state variables in the model.

The main substantive difference between our approach and these papers is that our estimator applies well beyond the discrete choice framework they consider. While some decisions, such as whether to enter or exit a market, are naturally modeled as discrete choices with an independent preference shock attached to each alternative, others are not. Decisions about price setting, capacity, research, or capital investment are better viewed as continuous choices, perhaps with a stochastic component of marginal returns. Our estimation approach incorporates these types of continuous choices, as well as discrete choices, in a unified framework.⁵ We think this has particular value for estimating dynamic models of industry competition. For example, Ryan (2006) applied our estimator to study regulation in a model where firms make both continuous investment and discrete entry and exit decisions.⁶

In simple discrete choice settings, one can use matrix algebra to compute value functions from conditional choice probabilities in the first stage of esti-

⁴Akerberg, Benkard, Berry, and Pakes (2007) provided a detailed discussion and comparison of the techniques in these papers.

⁵Several papers have proposed techniques based on first-order conditions that are specifically designed to estimate dynamic games with continuous decisions. Jofre-Bonet and Pesendorfer (2003) provided an elegant analysis of repeated auctions, and Berry and Pakes (2000) proposed an estimator that makes use of observed data on profits.

⁶See also Beresteanu and Ellickson (2006), Ryan and Tucker (2007), and Sweeting (2006) for other applications of the estimator.

mation. To handle both discrete and continuous decisions, we instead use forward simulation. The use of simulation was suggested by Hotz, Miller, Sanders, and Smith (1994), whose approach we follow. In the process we show how to exploit linearity to substantially reduce computation. If firms' profit functions are linear in the unknown parameters, so too will be their value functions. Therefore, in computing firms' value functions, it is possible to simulate the expected value terms just once and scale them by different parameter values, rather than having to simulate the value functions many times. This significantly reduces the amount of time required for estimation.

Our second stage of estimation also differs somewhat from the papers mentioned above. We use a minimum distance estimator based on minimizing violations of the equilibrium conditions, rather than maximum likelihood or method of moments (although we also discuss a strategy based on moment conditions). Pesendorfer and Schmidt-Dengler (2007) provided a least-squares interpretation of these alternatives. A potential value to our approach is that it extends easily to models that are not point identified. Strategic models with discrete choices, such as entry models, may be only partially identified. That is, even with infinite data, one can only place the parameters within a restricted set (Bresnahan and Reiss (1991), Ciliberto and Tamer (2006), Haile and Tamer (2003), and Pesendorfer and Schmidt-Dengler (2003)). Our estimation algorithm applies in this case, with little alteration, to produce bounds and confidence regions on the parameters of interest.

One general drawback to two-step estimators that is shared by our algorithm is that first-stage estimates do not fully exploit structure that may be imposed by the model. For instance, certain functional forms for policy or value functions may be incompatible with an equilibrium of the underlying model for the relevant parameter range, but estimation will not make use of this information. Our solution is to use a flexible first-stage model. This, however, is not a silver bullet, because the noise from the first stage can lead to finite sample bias when combined with a nonlinear second stage. A related problem, which we discuss in more detail in Section 5, is that two-step estimators appear to be more limited in their ability to handle serially correlated unobserved state variables.

Given this concern, we would like to assess whether the two-step approach can provide good estimates with reasonably sized data sets. To answer this question, we use Monte Carlo experiments to evaluate the efficiency and computational burden of our estimators. We consider two examples: a single-agent discrete choice decision problem similar to that of Rust (1987), and a version of the Ericson and Pakes (1995) dynamic oligopoly model that has both discrete entry and exit decisions and continuous investment decisions. We find that the algorithm has very low computational burden and works surprisingly well even for relatively small data sets. For instance, in the dynamic oligopoly example, we find that even with data sets smaller than one might have in real-world applications, it is possible to recover the entry cost distribution nonparametrically

with acceptable precision. We also provide estimates using a variation of the two stage estimator that is designed to reduce finite sample bias, and find that this version of the estimator performs particularly well.

The paper proceeds as follows. The next section introduces the class of models that we are interested in and provides two specific examples. Sections 3 and 4 outline the estimation algorithm and provide the relevant asymptotic theory. Section 6 details how the algorithm applies to the two examples and provides Monte Carlo evidence on the performance of the estimator. Section 7 concludes the paper.

2. A DYNAMIC MODEL OF COMPETITION

We start with a model of dynamic competition between oligopolistic competitors. The defining feature of the model is that actions taken in a given period may affect both current profits and, by influencing a set of commonly observed state variables, future strategic interaction. In this way, the model can permit aspects of dynamic competition such as entry and exit decisions, dynamic pricing or bidding, and investments in capital stock, advertising, or research and development. The model includes single-agent dynamic decision problems as a special case.

There are N firms, denoted $i = 1, \dots, N$, that make decisions at times $t = 1, 2, \dots, \infty$. Conditions at time t are summarized by a commonly observed vector of state variables $\mathbf{s}_t \in S \subset \mathbb{R}^L$. Depending on the application, relevant state variables might include the firms' production capacities, their technological progress up to time t , the current market shares, stocks of consumer loyalty, or simply the set of incumbent firms.

Given the state $\mathbf{s}_t$, firms choose actions simultaneously. These actions might include decisions about whether to enter or exit the market, investment or advertising levels, or choices about prices and quantities. Let $a_{it} \in A_i$ denote firm i 's action at time t and let $\mathbf{a}_t = (a_{1t}, \dots, a_{Nt}) \in A$ denote the vector of time t actions.

We assume that before choosing its action, each firm i receives a private shock ν_{it} , drawn independently across agents and over time from a distribution $G_i(\cdot | \mathbf{s}_t)$ with support $\mathcal{V}_i \subset \mathbb{R}^M$.⁷ The private shock might derive from variability in marginal costs of production, due for instance to the need for plant maintenance, or from variability in sunk costs of entry or exit. We denote the vector of private shocks as $\boldsymbol{\nu}_t = (\nu_{1t}, \dots, \nu_{Nt})$.

Each firm's profits at time t can depend on the state, the actions of all the firms, and the firm's private shock. We denote firm i 's profits by $\pi_i(\mathbf{a}_t, \mathbf{s}_t, \nu_{it})$.

⁷The assumption of independence across agents can be relaxed in some settings. One example is if each firm has a single dimension of private information and takes an action that is strictly increasing in its shock, as in a repeated private value auction. The added complication is that one must account for the correlation of opponent actions in first-stage policy function estimation.

Profits include variable returns as well as fixed or sunk costs incurred at date t , such as entry costs or the sell-off value of an exiting firm. We assume firms share a common discount factor $\beta < 1$.

Given a current state $\mathbf{s}_t$, firm i 's expected future profit, evaluated prior to realization of the private shock, is

$$\mathbb{E} \left[\sum_{\tau=t}^{\infty} \beta^{\tau-t} \pi_i(\mathbf{a}_\tau, \mathbf{s}_\tau, \nu_{i\tau}) \middle| \mathbf{s}_t \right].$$

The expectation is over i 's private shock and the firms' actions in the current period, as well as future values of the state variables, actions, and private shocks.

The final aspect of the model is the transition between states. We assume that the state at date $t+1$, denoted $\mathbf{s}_{t+1}$, is drawn from a probability distribution $P(\mathbf{s}_{t+1} | \mathbf{a}_t, \mathbf{s}_t)$. The dependence of $P(\cdot | \mathbf{a}_t, \mathbf{s}_t)$ on the firms' actions $\mathbf{a}_t$ means that time t behavior, such as entry/exit decisions or long-term investments, may affect the future strategic environment. Not all state variables necessarily are influenced by past actions; for instance, one component of the state could be an independent and identically distributed shock to market demand.

To analyze equilibrium behavior, we focus on pure strategy Markov perfect equilibria (MPE). In a MPE, each firm's behavior depends only on the current state and its current private shock. Formally, a Markov strategy for firm i is a function $\sigma_i: S \times \mathcal{V}_i \rightarrow A_i$. A profile of Markov strategies is a vector, $\boldsymbol{\sigma} = (\sigma_1, \dots, \sigma_n)$, where $\boldsymbol{\sigma}: S \times \mathcal{V}_1 \times \dots \times \mathcal{V}_N \rightarrow A$.

If behavior is given by a Markov strategy profile $\boldsymbol{\sigma}$, firm i 's expected profit given a state $\mathbf{s}$ can be written recursively:

$$V_i(\mathbf{s}; \boldsymbol{\sigma}) = \mathbb{E}_\nu \left[\pi_i(\boldsymbol{\sigma}(\mathbf{s}, \boldsymbol{\nu}), \mathbf{s}, \nu_i) + \beta \int V_i(\mathbf{s}'; \boldsymbol{\sigma}) dP(\mathbf{s}' | \boldsymbol{\sigma}(\mathbf{s}, \boldsymbol{\nu}), \mathbf{s}) \middle| \mathbf{s} \right].$$

Here V_i is firm i 's ex ante value function in that it reflects expected profits at the beginning of a period before private shocks are realized. We will assume that V_i is bounded for any Markov strategy profile $\boldsymbol{\sigma}$.

The profile $\boldsymbol{\sigma}$ is a Markov perfect equilibrium if, given the opponent profile $\boldsymbol{\sigma}_{-i}$, each firm i prefers its strategy σ_i to all alternative Markov strategies σ'_i . That is, $\boldsymbol{\sigma}$ is a MPE if for all firms i , states $\mathbf{s}$, and Markov strategies σ'_i ,

$$\begin{aligned} V_i(\mathbf{s}; \boldsymbol{\sigma}) &\geq V_i(\mathbf{s}; \sigma'_i, \boldsymbol{\sigma}_{-i}) \\ &= \mathbb{E}_\nu \left[\pi_i(\sigma'_i(\mathbf{s}, \nu_i), \boldsymbol{\sigma}_{-i}(\mathbf{s}, \boldsymbol{\nu}_{-i}), \mathbf{s}, \nu_i) \right. \\ &\quad \left. + \beta \int V_i(\mathbf{s}'; \sigma'_i, \boldsymbol{\sigma}_{-i}) dP(\mathbf{s}' | \sigma'_i(\mathbf{s}, \nu_i), \boldsymbol{\sigma}_{-i}(\mathbf{s}, \boldsymbol{\nu}_{-i}), \mathbf{s}) \middle| \mathbf{s} \right]. \end{aligned}$$

Doraszelski and Satterthwaite (2007) provided conditions for equilibrium existence in a closely related model. Here, we simply assume that a MPE exists, noting that there could be many such equilibria.

The structural parameters of the model are the discount factor β , the profit functions $\pi_1, \dots, \pi_N$, the transition probabilities P , and the distributions of the private shocks $G_1, \dots, G_N$. For econometric purposes, we treat the discount factor β as known and estimate the transition probabilities P directly from the observed state transitions. We assume the profit functions and the private shock distributions are known functions indexed by a finite parameter vector θ : $\pi_i(\mathbf{a}, \mathbf{s}, \nu_i; \theta)$ and $G_i(\nu_i | \mathbf{s}; \theta)$. The goal of estimation is to recover the true value of θ , denoted θ_0 , under the assumption that the data are generated by equilibrium behavior. Before describing our approach to estimation, however, we introduce two examples that we will use to illustrate our procedure.

2.1. Example: Dynamic Discrete Choice

Our first application is to dynamic discrete choice settings. In a discrete choice model, each agent i chooses its action at date t from a finite set of actions A_i . The state $\mathbf{s}_t$ includes variables observed by the economist, while ν_{it} is typically assumed to be a vector of choice-specific preference shocks $\{\nu_{it}(a_i)\}_{a_i \in A_i}$ that are additively separable in the profit function so $\pi_i(\mathbf{a}_t, \mathbf{s}_t, \nu_{it}; \theta) = \tilde{\pi}_i(\mathbf{a}_t, \mathbf{s}_t; \theta) + \nu_{it}(a_{it})$.

A classic example of dynamic discrete choice is the machine replacement problem in Rust (1987). A single manager owns a machine that operates each period. The state variable is the machine's age $s_t \in \{1, \dots, M\}$. At date t , the manager chooses whether to replace the machine ($a_t = 1$) or maintain it ($a_t = 0$), in which case the machine ages by one period. If the machine reaches age M , it remains M years old until it is replaced.

The manager's per-period profit is (dropping i subscripts because there is a single agent)

$$\pi(a, s, \nu; \theta) = \begin{cases} -\mu s + \nu(0), & \text{if } a = 0, \\ -R + \nu(1), & \text{if } a = 1, \end{cases}$$

where $\mu s - \nu(0)$ is the cost of maintaining the machine, $R - \nu(1)$ is the cost of replacement, and the parameters are $\theta = (\mu, R)$. A Markov policy $\sigma(s, \nu)$ specifies whether or not to replace the machine given the state s and shock ν . An optimal policy has a cutoff form: $\sigma(s, \nu) = 1$ if and only if $\nu(1) - \nu(0) \geq \eta(s)$, where $\eta(\cdot)$ is decreasing in the machine age s , consistent with the idea that older machines are more often replaced.

2.2. Example: Dynamic Oligopoly

Our second application is to a class of dynamic oligopoly models based on Ericson and Pakes (1995) and Pakes and McGuire (1994). In these mod-

els, there is a set of incumbent firms and a large number of potential entrants. Incumbent firms are heterogeneous, with each firm described by its state $z_{it} \in \{1, 2, \dots, \bar{z}\}$. Depending on the application, these states might represent product quality, capital stock, or productivity. Potential entrants have $z_{it} = 0$. In a given period, each incumbent can make an investment $I_{it} \geq 0$ to improve its state for the next period. Investment outcomes are random and a firm's investment has no direct effect on the state of the other firms.

Firms earn a profit by competing in a spot market. If firm i is an incumbent in period t , it earns

$$(1) \quad q_{it}(p_{it} - mc(q_{it}, \mathbf{s}_t; \mu)) - C(I_{it}, \nu_{it}; \xi),$$

where p_{it} is firm i 's price, $q_{it} = q_i(\mathbf{s}_t, \mathbf{p}_t; \lambda)$ is the quantity it sells, $mc(\cdot)$ is the marginal cost of production, ν_{it} represents a private shock to the cost of investment, and λ, μ, ξ are parameters. The state $\mathbf{s}_t$ includes at least the individual firm states and perhaps additional variables such as a common demand shock.

The typical assumption is that prices and quantities do not influence the evolution of the state variables. Instead they are determined in a static equilibrium conditional on the current state. In our Monte Carlo experiments in Section 6.2, we model the spot equilibrium as Bertrand–Nash in prices.

The model also allows for entry and exit. In period t , each incumbent firm can choose to exit the market at the end of the period and receive a fixed scrap value, ϕ . In addition, one potential entrant, selected randomly, has the opportunity to enter at a privately observed cost ν_{et} , drawn from a distribution G_e .

In equilibrium, incumbents make investment and exit decisions to maximize expected profits, so each incumbent i uses an investment strategy $I_i(\mathbf{s}, \nu_i)$ and exit strategy $\chi_i(\mathbf{s}, \nu_i)$. Because there are a large number of potential entrants, each follows a strategy $\chi_e(\mathbf{s}, \nu_e)$ that calls for it to enter if the expected profit from doing so exceeds its entry cost. In our Monte Carlo experiments, we focus on symmetric equilibria, meaning each incumbent's strategy is the same function of its own state and opponent states, and is invariant to permutations of opponent states.

3. FIRST-STAGE ESTIMATION: POLICY FUNCTIONS, STATE TRANSITIONS, AND VALUE FUNCTIONS

Our estimation approach proceeds in two steps. The goal of the first stage is to estimate the state transition probabilities $P(\mathbf{s}'|\mathbf{a}, \mathbf{s})$ and equilibrium policy functions $\boldsymbol{\sigma}(\mathbf{s}, \boldsymbol{\nu})$. The second stage uses the equilibrium conditions described above to estimate the structural parameters θ .

In this section, we start by describing the assumptions required to connect the data to the theoretical model. We then discuss the first-stage estimation and describe a simulation procedure that uses the estimated state transitions and policy functions to recover the equilibrium value functions

$V_1(\mathbf{s}; \boldsymbol{\sigma}), \dots, V_N(\mathbf{s}; \boldsymbol{\sigma})$. The next section of the paper explains how to estimate the parameters of the payoff functions and private shock distributions using the estimated value functions.

We should also note that in many cases some of the parameters of the profit function can be estimated without using the approach described below. For instance, in our dynamic oligopoly application, recovering the costs of investment, entry, and exit requires the full dynamic model, but the demand and marginal cost parameters can be estimated using standard static methods (e.g., Berry (2004), Berry, Levinsohn, and Pakes (1995)). For expositional purposes, we will not consider this possibility explicitly, although it is easily handled by simply assuming that profit function parameters that can be estimated using static methods are included in the vector of first-stage parameters.

3.1. *The Data Generating Process*

Data for estimation may come from a single market or a set of similar markets. For each market, we assume that the data at time t consist of the actions $\mathbf{a}_t$ and the complete vector of state variables $\mathbf{s}_t$. We discuss the possibility of unobserved state variables in Section 5. Our main assumption is that for each market, the data are generated by the same Markov perfect equilibrium of the above model.

ASSUMPTION ES—Equilibrium Selection: *The data are generated by a single Markov perfect equilibrium profile $\boldsymbol{\sigma}$.*

This assumption is relatively unrestrictive if the model has a unique equilibrium. It is stronger if the model has many equilibria and particularly if one wishes to pool data from many markets. Provided the state space is not expanded to include an identifier for each market, Assumption ES requires that equilibrium selection is consistent across markets.

3.2. *Estimating the Policy Functions and State Transitions*

The first step of our approach is to estimate the policy functions $\sigma_i: S \times \mathcal{V}_i \rightarrow A_i$ and state transition probabilities $P: A \times S \rightarrow \Delta(S)$. In practice, the best methods for doing this will depend on the application, so here we provide only an outline of the approach and details for two important cases.

Assuming the full vector of state variables and actions is observed, the main choice in estimating the transition probabilities is how flexible a specification to adopt. Because the transition probabilities are a model primitive, they would often be parameterized and estimated using parametric methods such as maximum likelihood. Alternatively, nonparametric methods can be used if one has little prior knowledge on the form of the state transitions. For instance, if the number of actions and states is small, the observed transition frequencies can be used as an estimate of P .

Estimating the policy functions may be more demanding, both because they are functions of the unobserved private shocks and because they result from equilibrium play. As a result, the optimal strategy and techniques may depend heavily on the specifics of the dynamic model. Here we discuss the two cases we believe will be encountered most often. The first is discrete choice with choice-specific payoff shocks, where we follow the literature on single-agent dynamic estimation, especially Hotz and Miller (1993) and Hotz, Miller, Sanders, and Smith (1994). The second is continuous choice where the optimal action is monotone in a private shock, where our approach is new.

One general point to emphasize about first-stage estimation is that an overly restrictive parametric assumption may yield policy functions that are inconsistent with an equilibrium of the underlying model. This means there is a benefit to flexible estimation. With a large number of states, flexibility may be difficult to achieve in practice, although dynamic models often have general structure—such as symmetry, constant returns, or monotonicity of optimal policies in certain state variables—that can be exploited in estimation.

3.2.1. Policy function estimates: Discrete choice

If firms choose from a finite set of actions, it is natural to adopt the traditional discrete choice framework outlined in Section 2.1. The key assumption for this approach is that each firm's private shock takes the form of a vector of choice-specific shocks that enter additively into the profit function.

ASSUMPTION DC—Discrete Choice: *For each firm i , $A_i = \{0, 1, \dots, K_i\}$, the profit shock v_i is a vector of choice-specific shocks $(v_i(a_i))_{a_i \in A_i}$, and the profit function is additively separable: $\pi_i(\mathbf{a}, \mathbf{s}, v_i) = \tilde{\pi}_i(\mathbf{a}, \mathbf{s}) + v_i(a_i)$.*

Note that Assumption **DC** together with our earlier assumptions implies Rust's (1994) assumptions of additive separability and conditional independence.

Given Assumption **DC**, let $v_i(a_i, \mathbf{s})$ denote the *choice-specific value function*

$$v_i(a_i, \mathbf{s}) = \mathbb{E}_{\boldsymbol{\nu}_{-i}} \left[\tilde{\pi}_i(a_i, \boldsymbol{\sigma}_{-i}(\mathbf{s}, \boldsymbol{\nu}_{-i}), \mathbf{s}) + \beta \int V_i(\mathbf{s}'; \boldsymbol{\sigma}) dP(\mathbf{s}' | a_i, \boldsymbol{\sigma}_{-i}(\mathbf{s}, \boldsymbol{\nu}_{-i}), \mathbf{s}) \right].$$

With this notation, firm i optimally chooses an action a_i that satisfies

$$(2) \quad v_i(a_i, \mathbf{s}) + v_i(a_i) \geq v_i(a'_i, \mathbf{s}) + v_i(a'_i) \quad \forall a'_i \in A_i.$$

Equation (2) represents the policy rule to be estimated in the first stage. Clearly this requires estimating the choice-specific value functions $v_i(a_i, \mathbf{s})$ for

each action a_i and state $\mathbf{s}$. One can apply well known methods to estimate these functions from observed choice data. In particular, Hotz and Miller (1993) showed how to recover the choice-specific value functions by inverting the observed choice probabilities at each state. To illustrate their approach, suppose the private shocks have a type 2 extreme value distribution, so that we have a logit choice model. Then for any two actions $a_i, a'_i \in A_i$,

$$v_i(a'_i, \mathbf{s}) - v_i(a_i, \mathbf{s}) = \ln(\Pr(a'_i|\mathbf{s})) - \ln(\Pr(a_i|\mathbf{s})),$$

where $\Pr(a_i|\mathbf{s})$ is the probability of observing choice a_i in state $\mathbf{s}$, which is easily estimated.

As a general rule, discrete choice methods such as the Hotz–Miller inversion only permit one to recover differences in the choice-specific value functions. Identifying the levels of these functions requires a normalization, such as setting $v_i(a_i = 0, \mathbf{s}) = 0$ for all $\mathbf{s} \in S$. To estimate the policy rules σ_i , however, it suffices to recover the differences, as is apparent from (2).

Two observations may be useful. First, in the logit example, the distribution of firm i 's choice-specific shocks $G_i(\cdot)$ is assumed to be known (i.e., it is not a function of θ). As a result, the Hotz–Miller inversion does not depend on any unknown parameters. More generally, the inversion may depend on unknown parameters of the distribution of private shocks, in which case the policy function could be recovered only as a function of these parameters. We discuss this issue further below.

Second, in many industrial organization applications, one may want to allow for relatively large state spaces. In this case, a state-by-state inversion approach is likely to generate very noisy estimates of the policy functions. Instead it may be preferable to smooth the estimates across states. This could be done by modeling the choice-specific value functions $v_i(\cdot)$ as flexibly parameterized functions of the actions and states. Alternatively, Hotz and Miller use kernel methods. As noted above, theoretical restrictions from the model, such as a focus on symmetric equilibria, may allow for much better use of limited data.

3.2.2. Policy function estimates: Continuous choice

Many applications of our general model, such as models of dynamic pricing or bidding and models of investment in product quality or advertising, involve choices that do not naturally have the kind of discrete choice structure described above. Instead, optimal policies have the form $a_i = \sigma_i(\mathbf{s}, \nu_i)$, where $a_i \in A_i \subset \mathbb{R}$ is an action such as investment or price and ν_i is a single-dimensional private shock. Typically policies satisfy the additional property that the optimal policy is increasing in the private shock. We propose a method for estimation that covers this leading case and we discuss extensions to multi-dimensional actions (e.g., investment in quality and advertising) below.

ASSUMPTION MC—Monotone Choice: *For each agent i , $A_i, \mathcal{V}_i \subset \mathbb{R}$ and $\pi_i(\mathbf{a}, \mathbf{s}, \nu_i)$ has increasing differences in (a_i, ν_i) .*

This assumption, which is equivalent to $\partial^2 \pi_i / (\partial a_i \partial v_i) \geq 0$ if π_i is twice-differentiable, implies that v_i affects the marginal return to higher actions (e.g., it might affect the marginal cost of investment). By Topkis' theorem, it implies that firm i 's optimal policy $\sigma_i(\mathbf{s}, v_i)$ will be increasing in v_i .⁸

To exploit this monotonicity in estimation, let $F_i(a_i|\mathbf{s})$ denote the probability that firm i takes an action less than or equal to a_i at state $\mathbf{s}$. This quantity is observed in the data. Because $\sigma_i(\mathbf{s}, v_i)$ is increasing in v_i , $F_i(a_i|\mathbf{s}) = \Pr(\sigma_i(\mathbf{s}, v_i) \leq a_i|\mathbf{s}) = G_i(\sigma_i^{-1}(a_i; \mathbf{s})|\mathbf{s}; \theta)$. Substituting $a_i = \sigma_i(\mathbf{s}, v_i)$ and rearranging, we have for all $(\mathbf{s}, v_i)$,

$$\sigma(\mathbf{s}, v_i) = F_i^{-1}(G_i(v_i|\mathbf{s}; \theta)|\mathbf{s}).$$

Therefore, to estimate the policy function at a given state $\mathbf{s}$, one needs only to estimate the distribution of actions at each state $F_i(a_i|\mathbf{s})$ and have knowledge of G_i . As an example, consider the dynamic oligopoly model. In that model, each firm has an investment policy function given by $I(\mathbf{s}, v_i)$. Suppose further that $v_i \sim \mathcal{N}(0, \sigma^2)$. Then the policy function $I(\cdot)$ could be estimated by estimating $F(I|\mathbf{s})$ using the data, inverting it at each point $\mathbf{s}$, and evaluating it at $\Phi(v_i/\sigma)$, where Φ denotes the normal cumulative distribution function.

The two observations made in the discrete choice case apply here as well. First, the policy function estimates may depend on any unknown parameters of G_i . Second, in practice it may be desirable to smooth the estimates of F_i across states and to impose any known theoretical restrictions on F_i such as symmetry.

3.3. Estimating the Value Functions

The purpose of estimating the equilibrium policy functions is that they allow us to construct estimates of the equilibrium value functions, which can be used in turn to estimate the structural parameters of the model. In this section, we show how forward simulation can be used to estimate firms' value functions for given strategy profiles (including the equilibrium profile) given an estimate of the transition probabilities P . Our approach is inspired by Hotz, Miller, Sanders, and Smith (1994), although it differs slightly in that we draw states and private shocks each period and compute actions, while they directly simulated states and actions, and relied on the specific structure of the logit model to compute the contribution of the private shock to payoffs.

⁸The policy function may be only weakly increasing rather than strictly increasing. For expositional purposes, we consider the strictly increasing case, but the approach also works for the weakly increasing case with minor notational adjustments. One example where the optimal action is weakly increasing is when $\mathcal{A}_i$ is discrete but ordered and $\mathcal{V}_i$ is continuous; our method applies here as well.

Let $V_i(\mathbf{s}; \boldsymbol{\sigma}; \theta)$ denote the value function of firm i at state $\mathbf{s}$, assuming firm i follows the Markov strategy σ_i and rival firms follow the strategy $\boldsymbol{\sigma}_{-i}$. Then

$$V_i(\mathbf{s}; \boldsymbol{\sigma}; \theta) = \mathbb{E} \left[\sum_{t=0}^{\infty} \beta^t \pi_i(\boldsymbol{\sigma}(\mathbf{s}_t, \boldsymbol{\nu}_t), \mathbf{s}_t, \boldsymbol{\nu}_{it}; \theta) \mid \mathbf{s}_0 = \mathbf{s}; \theta \right],$$

where the expectation is over current and future values of the private shocks $\boldsymbol{\nu}_t$ and states $\mathbf{s}_t$. Note that the expectation may depend on θ both through the profit function π_i and through the distributions of private shocks, $G_1, \dots, G_N$.

Given a first-stage estimate $\hat{P}$ of the transition probabilities, we can use simulation to estimate the value function $V_i(\mathbf{s}; \boldsymbol{\sigma}; \theta)$ for any strategy profile $\boldsymbol{\sigma}$ and parameter vector θ . A single simulated path of play can be obtained as follows:

- Step 1. Starting at state $\mathbf{s}_0 = \mathbf{s}$, draw private shocks $\boldsymbol{\nu}_{i0}$ from $G_i(\cdot \mid \mathbf{s}_0, \theta)$ for each firm i .
- Step 2. Calculate the specified action $a_{i0} = \sigma_i(\mathbf{s}_0, \boldsymbol{\nu}_{i0})$ for each firm i and the resulting profits $\pi_i(\mathbf{a}_0, \mathbf{s}_0, \boldsymbol{\nu}_{i0}; \theta)$.
- Step 3. Draw a new state $\mathbf{s}_1$ using the estimated transition probabilities $\hat{P}(\cdot \mid \mathbf{a}_0, \mathbf{s}_0)$.
- Step 4. Repeat Steps 1–3 for T periods or until each firm reaches a terminal state with known payoff (e.g., exit from the market).⁹

Averaging firm i 's discounted sum of profits over many simulated paths of play yields an estimate of $V_i(\mathbf{s}; \boldsymbol{\sigma}; \theta)$, which we denote $\hat{V}_i(\mathbf{s}; \boldsymbol{\sigma}; \theta)$. Such an estimate can be obtained for *any* $(\boldsymbol{\sigma}, \theta)$ pair, including $(\hat{\boldsymbol{\sigma}}, \theta)$, where $\hat{\boldsymbol{\sigma}}$ is the policy profile that results from first-stage estimation. It follows that $\hat{V}_i(\mathbf{s}; \hat{\boldsymbol{\sigma}}; \theta)$ is an estimate of firm i 's payoff from playing $\hat{\sigma}_i$ in response to opponent behavior $\hat{\boldsymbol{\sigma}}_{-i}$ and that $\hat{V}_i(\mathbf{s}; \sigma_i, \hat{\boldsymbol{\sigma}}_{-i}; \theta)$ is an estimate of its payoff from playing σ_i in response to $\hat{\boldsymbol{\sigma}}_{-i}$, in both cases conditional on parameters θ . Below we show how such estimates, combined with the equilibrium conditions of the model, permit estimation of the underlying structural parameters.

3.3.1. Using linearity to reduce computation

The forward simulation procedure allows a relatively low-cost estimate of the value functions for different strategy profiles given a parameter vector θ . In an estimation algorithm that searches over parameters, however, the procedure must be repeated for each candidate parameter value (even if the same simulation draws are used throughout). Fortunately, there is one particularly useful case where this additional computation can be avoided.

Suppose the distribution of private shocks is known and the profit function is linear in the unknown parameters θ so that

$$\pi_i(\mathbf{a}, \mathbf{s}, \boldsymbol{\nu}_i; \theta) = \Psi_i(\mathbf{a}, \mathbf{s}, \boldsymbol{\nu}_i) \cdot \theta,$$

⁹Our asymptotic results will assume $T \rightarrow \infty$ to drive the simulation error to zero, but in practice one can stop when β^T becomes insignificantly small relative to the simulation error generated by averaging over only a finite number of paths.

where $\Psi_i(\mathbf{a}, \mathbf{s}, \nu_i)$ is an M -dimensional vector of “basis functions” $\psi_i^1(\mathbf{a}, \mathbf{s}, \nu_i), \dots, \psi_i^M(\mathbf{a}, \mathbf{s}, \nu_i)$. If π_i is linear in $(\mathbf{a}, \mathbf{s}, \nu_i)$, the functions ψ_i^j are simply the elements of $(\mathbf{a}, \mathbf{s}, \nu_i)$, but more generally the basis functions could be polynomials or more complicated functions of $(\mathbf{a}, \mathbf{s}, \nu_i)$.

Under these assumptions, we can write the value functions as

$$(3) \quad V_i(\mathbf{s}; \boldsymbol{\sigma}; \theta) = \mathbb{E} \left[\sum_{t=0}^{\infty} \beta^t \Psi_i(\boldsymbol{\sigma}(\mathbf{s}_t, \nu_t), \mathbf{s}_t, \nu_{it}) \middle| \mathbf{s}_0 = \mathbf{s} \right] \cdot \theta = \mathbf{W}_i(\mathbf{s}; \boldsymbol{\sigma}) \cdot \theta.$$

Again the expectation is over current and future values of the private shocks and states. The useful simplification is that $\mathbf{W}_i = (W_i^1, \dots, W_i^M)$ does not depend on the unknown parameters θ . Therefore, for any strategy profile $\boldsymbol{\sigma}$, one can use the forward simulation procedure once to estimate each $\mathbf{W}_i$ and then obtain V_i easily for any value of θ .

The main restriction needed to achieve this computational savings is that the profit function be linear in the unknown parameters. A number of standard models have this feature. Entry, exit, and fixed cost parameters enter additively into a firm’s profit function. Similarly, investment cost and marginal cost parameters typically enter profits linearly. So, for instance, the value functions in the dynamic oligopoly model described above will be linear in the parameters that are unknown at the time of value function estimation. We take advantage of this property in our Monte Carlo experiments in Section 6.

The derivation of (3) above assumes that $G_1, \dots, G_N$ are known. If $G_1, \dots, G_N$ are functions of unknown parameters, it may still be possible to obtain an expression for V_i that is linear in parameters. For example, if ν_{it} is $N(\mu, \kappa^2)$, where μ and κ are unknown parameters, and π_i contains the term $a_{it}\nu_{it}$, this term can be written as $a_{it}(\mu + \kappa\omega_{it})$, where ω_{it} is $N(0, 1)$. Linearity also can be exploited if V_i is linear in some parameters, but not others. The idea here would be to maximize the second-stage objective function in two steps, with an inner step that maximizes the objective over the linearized parameters and an outer step that maximizes the inner step objective over the nonlinear parameters. This kind of nested optimization procedure is likely to significantly reduce computational burden.

Of course, it is also easy to write down models that are not linear in parameters, so we will not assume linearity in what follows. The general case is conceptually straightforward and the computational burden may often still be low using our methods.

3.4. First-Stage Estimation: Discussion and Extensions

So far we have suggested methods for dealing with either a single discrete choice or a single continuous choice in isolation. In many applications, such as

our dynamic oligopoly example, firms may make several simultaneous choices and receive several private shocks. This is straightforward to handle if the decisions can be cleanly separated, but potentially more complicated if multiple shocks enter a given policy function. In that case, the policy function may not be identified from the data.¹⁰ Our method requires that one be able to consistently estimate each firm's policy function, so this may limit our ability to estimate certain models.

An important point regarding the first-stage estimation concerns *which* realizations of the state variables are observed in the data. Under Assumptions ES and either DC or MC, the firms' policy functions can be estimated consistently at any state s that is the support of the data generating process. Similarly the transition probabilities P can be consistently estimated for these states. Consistent estimation of the value functions $V_1, \dots, V_N$ at a state s requires somewhat more, namely that the transition probabilities and policy functions can be estimated not only at s , but also at all states reachable from s with positive probability in equilibrium.

A potential difficulty in obtaining consistent first-stage estimates is that some states s may never be reached in equilibrium. For instance, in our dynamic oligopoly example, the number of incumbent firms may never exceed some fixed upper bound in equilibrium. More precisely, the general model above generates a Markov process on the state space that may have an invariant distribution with a support that is significantly smaller than the entire state space. Depending on what type of data is available, it may not be possible to consistently estimate the policy functions or value functions at states outside this support. In practice one can try to overcome this in two ways. One possibility is to make a parametric assumption on the form of the policy functions and transition probabilities, and then extrapolate to states that are not observed in the data. Another possibility is to recognize the lack of identification and, in our second stage, omit the equilibrium conditions for states that do not appear in the data.¹¹

Finally, we have so far ruled out the possibility of serially correlated unobserved state variables. We return to this issue in Section 5 after we discuss second-stage estimation.

¹⁰Note that the distribution of actions given states $F(a|s)$ will still be observed. The identification problem arises in translating this distribution into policy functions. For example, if a firm has private information each period about both demand and marginal cost, we may not be able to consistently estimate its pricing policy as a function of these two shocks.

¹¹Although we will not focus on it, the possibility of unrealized states could have a substantive consequence for what the data reveals about parameters of firms' profit functions that are relevant only along "off equilibrium" paths. An example of this would be parameters that determine the payoffs in a punishment regime that is never triggered in equilibrium.

4. SECOND-STAGE ESTIMATION: RECOVERING THE STRUCTURAL PARAMETERS

The first stage generates estimates of the firms' equilibrium policy functions, the state transitions, and the value functions. In this section, we explain how to combine these estimates with the equilibrium conditions of the model to recover, or at least to bound, the structural parameters of the model.

To simplify this second-stage problem, we henceforth assume that the policy function and transition probabilities are parameterized by a finite parameter vector α and that this vector can be consistently estimated at the first stage.

ASSUMPTION S1: *The policy functions $\sigma_i(\mathbf{s}, v_i; \alpha)$ and transition probabilities $P(\mathbf{s}_{t+1}|\mathbf{s}_t, \mathbf{a}_t; \alpha)$ are parameterized by a finite parameter vector α , and there exists a consistent estimator $\hat{\alpha}_n$ with the property that $\sqrt{n}(\hat{\alpha}_n - \alpha_0) \rightarrow^d N(0, V_\alpha)$, where α_0 is the true parameter that generates the data.*

This assumption permits a nonparametric first stage with discrete action and state spaces or a parametric first stage with continuous actions and states. It rules out a nonparametric first stage with continuous actions or states. Hotz and Miller (1993) used a nonparametric kernel estimator to address the discrete action, continuous state case. Bajari, Chernozhukov, and Hong (2005) considered this same case using sieve estimation methods in the first stage. We are optimistic that our approach could be shown to work for a nonparametric first stage with continuous actions on a continuous state space, but we leave this for future research.

4.1. The Equilibrium Conditions

Our second-stage estimator makes use of the model's equilibrium conditions. Recall that the strategy profile σ is a Markov perfect equilibrium if and only if for all firms i , all states $\mathbf{s}$, and all alternative Markov policies σ'_i ,

$$(4) \quad V_i(\mathbf{s}; \sigma_i, \sigma_{-i}; \theta) \geq V_i(\mathbf{s}; \sigma'_i, \sigma_{-i}; \theta).$$

The equilibrium inequalities (4) define a set of parameters that *rationalize* the strategy profile σ in the sense that σ is a Markov perfect equilibrium of the game defined by P and θ . Let Θ_0 denote this set:

$$\Theta_0(\sigma, P) := \{\theta: \theta, \sigma, P \text{ satisfy (4) for all } \mathbf{s}, i, \sigma'_i\}.$$

The goal of second-stage estimation is to recover Θ_0 using our first-stage estimates of σ and P .

Depending on the model and its parameterization, the set Θ_0 may or may not be a singleton. This is the problem of *identification*, studied in the context of single-agent dynamic decision problems by Rust (1994) and Magnac

and Thesmar (2002) and in the context of dynamic games by Pesendorfer and Schmidt-Dengler (2003) and Bajari, Chernozhukov, and Hong (2005). These papers show that while point identification is often possible, it requires fairly strong assumptions, particularly in the multiagent case.¹² Even without it, however, knowledge of the set Θ_0 may convey useful information about the underlying parameters (e.g., Manski and Tamer (2002), Haile and Tamer (2003), Ciliberto and Tamer (2006)).

We therefore consider two alternative estimators. The first estimator requires that the model be identified; it yields standard point estimates of the parameters. The second estimator does not rely on identification; it yields consistent “bounds” estimates of the identified set of parameters.

Both estimators follow the same basic outline. We start by constructing empirical counterparts to all or a subset of the equilibrium inequalities (4) using the forward simulation procedure described in Section 3.3. We then look for the value(s) of θ that minimizes the (squared distance) violations of these inequalities.

4.2. Estimation of Identified Models

Before describing the estimators, it is useful to introduce a small amount of notation. Let $x \in \mathcal{X}$ index the equilibrium conditions, so that each x denotes a particular (i, s, σ'_i) combination. In a slight abuse of notation, define

$$g(x; \theta, \alpha) = V_i(s; \sigma_i, \sigma_{-i}; \theta, \alpha) - V_i(s; \sigma'_i, \sigma_{-i}; \theta, \alpha).$$

The dependence of $V_i(s; \sigma; \theta, \alpha)$ on α reflects the fact that σ and P are parameterized by α . The inequality defined by x is satisfied at θ, α if $g(x; \theta, \alpha) \geq 0$.

Define the function

$$Q(\theta, \alpha) := \int (\min\{g(x; \theta, \alpha), 0\})^2 dH(x),$$

where H is a distribution over the set $\mathcal{X}$ of inequalities.

The true parameter vector, θ_0 , satisfies

$$Q(\theta_0, \alpha_0) = 0 = \min_{\theta \in \Theta} Q(\theta, \alpha_0),$$

where $\Theta \subset \mathbb{R}^M$ contains θ_0 . Assuming that the model is identified and that H has sufficiently large support, $Q(\theta, \alpha_0) > 0$ for all $\theta \neq \theta_0$, so θ_0 uniquely

¹²These papers show that in a strategic setting, nonparametric identification requires a set of state variables that affect the profit functions of some but not all firms. This requirement is analogous to the standard exclusion restrictions needed for identification in a linear simultaneous equations model. Of course, it may also be the case that the model is overidentified. Although we do not pursue it here, our model suggests a natural specification test for overidentified models based on testing whether there is any parameter vector θ that satisfies the equilibrium conditions.

minimizes $Q(\theta, \alpha_0)$. Therefore, we propose to estimate θ by minimizing the sample analog of $Q(\theta, \alpha_0)$.

To do this, let $\{X_k\}_{k=1, \dots, n_I}$ be a set of inequalities from $\mathcal{X}$ chosen by the econometrician (representing independent and identically distributed draws from H). These might be selected in a variety of ways. One possibility is to draw firms and states at random and then consider alternative policies σ'_i that are slight perturbations of the estimated policy $\sigma_i(s, \nu_i; \hat{\alpha}_n)$, for example, $\sigma'_i(s, \nu_i) = \sigma_i(s, \nu_i; \hat{\alpha}_n) + \epsilon$. Another possibility is to focus on alternative policies that depart from the estimated policy at a single state. The particular method of selecting inequalities (the distribution H) will have implications for efficiency, but the only requirement for consistency is that H has sufficient support to yield identification.¹³

For each chosen inequality, X_k , the next step is to use the forward simulation procedure from Section 3.3 to construct analogs of each of the V_i terms. Let $\hat{g}(x; \theta, \hat{\alpha}_n)$ denote the empirical counterpart to $g(x; \theta, \alpha)$ computed by replacing the V_i terms with simulated estimates $\hat{V}_i$. Let n_s denote the number of simulation draws. We can then define

$$(5) \quad Q_n(\theta, \alpha) := \frac{1}{n_I} \sum_{k=1}^{n_I} (\min\{\hat{g}(X_k; \theta, \alpha), 0\})^2.$$

Our estimator minimizes this objective function at $\alpha = \hat{\alpha}_n$. That is,

$$\hat{\theta} := \arg \min_{\theta \in \Theta} Q_n(\theta, \hat{\alpha}_n).$$

We now specify sufficient conditions for this estimator to be consistent and asymptotically normal.

ASSUMPTION S2:

- (i) *The inequalities $X_1, \dots, X_{n_I} \in \mathcal{X}$ are independent and identically distributed draws from H .*
- (ii) *For each X_k , each $\hat{V}_i$ is computed using independent draws and satisfies $\mathbb{E}\hat{V}_i = V_i < \infty$. In addition, with probability 1, $\hat{V}$ is twice differentiable in θ and α , and three times differentiable in θ .*
- (iii) *As $n \rightarrow \infty$, both $n_s, n_I \rightarrow \infty$ and $n/n_s^2 \rightarrow 0$.*
- (iv) *The set $\Theta \subset \mathbb{R}^M$ is compact and $\theta_0 = \arg \min_{\theta \in \Theta} Q(\theta, \alpha_0)$.*
- (v) *There exists a full-rank matrix B_0 such that, for θ near θ_0 ,*

$$\frac{\partial}{\partial \theta} Q_n(\theta, \hat{\alpha}_n) = \frac{\partial}{\partial \theta} Q_n(\theta_0, \hat{\alpha}_n) + (B_0 + o_p(1))(\theta - \theta_0).$$

¹³Note that in some cases the space of feasible alternative policies may contain a large number of alternatives that provide little or no information about the parameters. In such cases it may be desirable to choose alternatives in a manner that places little or no weight on these policies. Holmes (2006) provided a nice illustration and solution to this problem.

Part (i) was discussed above. Part (ii) says that the value function simulator is unbiased and uses independent draws for each inequality (the same draws may be used for the two value functions in a given inequality).¹⁴ It also makes some standard smoothness assumptions that would be satisfied in most applications (and the linear-in-parameters case, in particular, leads automatically to smoothness in θ). Part (iii) requires that as n goes to infinity, the number of inequalities sampled and the number of simulation draws per inequality also go to infinity, the latter at a rate faster than $\sqrt{n}$. This guarantees that the simulation error adds nothing to the asymptotic variance of the estimator. Part (iv) is the identification assumption. Finally, part (v) requires that the objective function satisfy a second-order expansion. It will typically be the case that

$$B_0 = B(\theta_0) \quad \text{where} \quad B(\theta) \equiv \frac{\partial^2}{\partial \theta \partial \theta'} Q(\theta, \alpha_0).$$

This assumes that the asymptotic objective function is twice differentiable on a neighborhood of θ_0 , which would typically be the case for an identified model.¹⁵ In addition to differentiability of the limit function, (v) would also typically require that an equicontinuity condition be satisfied.

PROPOSITION 1: *Under Assumptions S1 and S2,*

$$\sqrt{n}(\hat{\theta} - \theta_0) \xrightarrow{d} N(0, B_0^{-1} \Lambda_0 V_\alpha \Lambda_0' B_0^{-1}),$$

where

$$\Lambda_0 \equiv \frac{\partial^2}{\partial \theta \partial \alpha'} Q(\theta, \alpha) \quad \text{evaluated at} \quad \theta = \theta_0, \alpha = \alpha_0.$$

Because the simulation does not contribute anything to the asymptotic distribution of the estimator, the standard errors are determined by the first-stage sampling error adjusted for its effect on the second-stage estimates. Note that efficiency of the estimator depends on the distribution used to sample over inequalities (H) through the expectation operators in Λ_0 and B_0 . Deriving an expression for Λ_0 in terms of model primitives is difficult because of the complex way in which the first-stage parameters enter into the W terms. Therefore, in practice we believe it will typically be easiest to use subsampling or the bootstrap to estimate standard errors. This is the approach we follow in our Monte Carlo experiments.

¹⁴In fact, using the same draws for the two value functions within a given inequality would typically be desirable because the two simulated value functions would then be positively correlated, reducing the simulation variance in the inequality as a whole.

¹⁵For this to be true, it is important that the distribution H sample inequalities smoothly. If condition (v) is violated, then the estimator would not be asymptotically normal.

One point worth mentioning is that in a finite sample, the resulting parameter estimate $\hat{\theta}$ may not fully rationalize the estimated policies as an equilibrium of the model. Clearly if all inequalities are satisfied, the estimated policies are an equilibrium of the model defined by $\hat{\theta}$. This will be the case asymptotically. More generally, if the inequalities are nearly satisfied, the estimated policies will be “close” to being an equilibrium, but provided the simulation error is small, the violations will imply that at least one player has a profitable deviation given $\hat{\theta}$ and the estimated policies of the other players. This is typical of two-stage estimators, and one reason why Aguirregabiria and Mira (2007) suggested an iterative extension of these methods.

4.3. Bounds Estimation

Our minimum distance estimator extends naturally to models that may only be set identified. To develop this extension, we drop our earlier identification requirements, Assumption S2 parts (iv) and (v). We also drop part (i) and the assumption that $n_I \rightarrow \infty$, and instead consider estimation using a fixed set of n_I inequalities.¹⁶ Let

$$\Theta_{n_I}(\sigma, P) := \{\theta : \theta, \sigma, P \text{ satisfy (4) for all } x_k, k = 1, \dots, n_I\}.$$

If $x_1, \dots, x_{n_I}$ represent the full set of equilibrium conditions, then $\Theta_{n_I} = \Theta_0$; more generally, $\Theta_0 \subseteq \Theta_{n_I}$.

Our estimator and proof of consistency closely resemble those of Manski and Tamer (2002) and Haile and Tamer (2003). The distribution theory for the estimator originated in Chernozhukov, Hong, and Tamer (2007). Below we apply an extension of their method described by Romano and Shaikh (2006a).

Specifically, we consider the same squared distance function $Q_n(\theta, \hat{\alpha}_n)$ defined in (5), with the understanding that the relevant set of equilibrium conditions is fixed rather than sampled at random. Following Manski and Tamer (2002), we define our estimate of the set Θ_{n_I} to encompass all values of θ that come within a specified distance of minimizing $Q_n(\theta, \hat{\alpha}_n)$. That is,

$$(6) \quad \hat{\Theta}_n \equiv \left\{ \theta : Q_n(\theta, \hat{\alpha}_n) \leq \min_{\theta' \in \Theta} Q_n(\theta', \hat{\alpha}_n) + \mu_n \right\}$$

for some $\mu_n > 0$, where $\mu_n \rightarrow 0$.

Define the distance $\rho(\Theta, \Theta')$ between two sets $\Theta, \Theta' \subset \mathbb{R}^L$ as $\rho(\Theta, \Theta') \equiv \sup_{\theta \in \Theta} \inf_{\theta' \in \Theta'} |\theta - \theta'|$. Note that ρ is not symmetric; $\rho(\Theta, \Theta')$ measures the greatest distance from a point in Θ to the set Θ' .

¹⁶Note that it would also be possible to consider $n_I \rightarrow \infty$ here and simply drop the identification condition.

PROPOSITION 2: Under Assumptions S1 and S2(ii) and (iii),

$$\rho(\widehat{\Theta}_n, \Theta_{n_I}) \xrightarrow{P} 0.$$

Furthermore, if $\sup_{\theta \in \Theta} |Q_n(\theta, \hat{\alpha}_n) - Q(\theta, \alpha_0)|/\mu_n \rightarrow 0$, then

$$\rho(\Theta_{n_I}, \widehat{\Theta}_n) \xrightarrow{P} 0.$$

The first part of the proposition says that for large n , every point in $\widehat{\Theta}_n$ is close to a point in Θ_{n_I} . This is guaranteed under conditions similar to those required for standard estimators and holds even if $\mu_n = 0$ for all n . The second part says that for large n , every point in $\widetilde{\Theta}_0$ is close to a point in $\widehat{\Theta}_n$. Essentially this means that the entire set $\widetilde{\Theta}_0$ is eventually captured by the estimated set. This second part requires that μ_n go to zero slowly enough, with the required rate provided above.¹⁷

The primary difficulty in estimation is computation of the set $\widehat{\Theta}_n$. Manski and Tamer (2002) and Haile and Tamer (2003) used simulated annealing to sample the objective function and then constructed one-dimensional bounds on the parameters. Such an approach could also be used here in the general case. If the system of inequalities is linear in the parameters, then computation of $\widehat{\Theta}_n$ is much easier because it is a linear programming problem. One set of techniques in the operations research literature, dating back to Motzkin et al. (1953), uses the fact that Θ_n is a convex polyhedron to describe it in terms of its vertices. Alternatively, Bajari and Benkard (2004) suggested a Gibbs sampling procedure that produces simulation draws that are uniformly distributed over the set $\widehat{\Theta}$. These simulation draws can be used to estimate bounds on the parameters, θ . The latter method is particularly efficient at handling large numbers of inequalities, making it ideally suited to this problem.

Our model is similar to one considered by Chernozhukov, Hong, and Tamer (2007, Example 1), so their subsampling algorithm (Algorithm 2.1) could be used to construct confidence regions. We describe instead an extension of their procedure posed by Romano and Shaikh (2006a). Romano and Shaikh's procedure produces a set estimator $\widehat{\Theta}_{0.95}$ such that for all $\theta \in \widehat{\Theta}_n$,

$$\liminf_{n \rightarrow \infty} \Pr(\theta \in \widehat{\Theta}_{0.95}) \geq 0.95.$$

Note that since $\widehat{\Theta}_{0.95}$ is constructed such that the above condition is satisfied for all $\theta \in \widehat{\Theta}_n$, the condition is guaranteed to be satisfied at the true parameter value θ_0 .

¹⁷Note that because $\widetilde{\Theta}_0$ is a convex polyhedron in the linear-in-parameters case, consistency could be shown for that case even for $\mu_n = 0$.

In the context of our model, Romano and Shaikh's procedure can be described as follows:

- Step 1. Begin with a set Θ that is large enough that it is known with certainty that $\theta_0 \in \Theta$.
- Step 2. Construct B subsamples of size n_b and compute $Q_{n,b}(\theta, \hat{\alpha}_n)$ for each subsample, $b \in \{1, \dots, B\}$, for each $\theta \in \Theta$.
- Step 3. For each $\theta \in \Theta$, compute a critical value, $\hat{c}_n(0.95, \theta)$, such that

$$\hat{c}_n(0.95, \theta) = \inf \left\{ c : \frac{1}{B} \sum_{b=1}^B \mathbf{1}\{n_b Q_{n,b}(\theta, \hat{\alpha}_n) \leq c\} \geq 0.95 \right\}.$$

- Step 4. Compute $\hat{\Theta}_{0.95} = \{\theta : n \cdot Q_n(\theta, \hat{\alpha}_n) \leq \hat{c}_n(0.95, \theta)\}$.

Romano and Shaikh (2006a) showed that this procedure leads to a consistent estimate of the set $\hat{\Theta}_{0.95}$.¹⁸

5. SERIALLY CORRELATED UNOBSERVED STATE VARIABLES

A limitation to our approach is the assumption that all commonly known states, $\mathbf{s}$, are observed. This leaves the *independent* private shocks as the only source of unobserved variation over time. This section discusses alternative methods for incorporating serially correlated unobserved state variables.

In many cases, unobserved state variables can be recovered in first-stage estimation. Examples include the estimation of unobserved product characteristics in demand models (e.g., Berry (1994), Berry, Levinsohn, and Pakes (1995)) and unobserved productivity shocks in production functions (e.g., Olley and Pakes (1996)). In such cases, once the values of the unobserved state variables are recovered, it is “as if” they were observable in the first place (subject to caveats about estimation error) and it is, therefore, simple to allow for quite general forms of temporal dependence. We believe that this is an important and potentially useful case due to the widespread use of the methods listed above.

With the advent of long panels, another possibility is to use panel data methods. Suppose that data are available on several markets over time and that the

¹⁸Imbens and Manski (2004) argued convincingly that the above pointwise confidence regions are logically preferred to confidence regions for the identified set as a whole. However, Romano and Shaikh (2006b) also provided an estimator for a confidence region for the identified set as a whole. In that case, it is necessary to follow a “step down” procedure that estimates $\hat{\Theta}_{0.95}$ iteratively. Starting with $\hat{\Theta}_{0.95}^0 = \Theta$, at each iteration i , a cutoff value $\hat{c}_n^i(0.95)$ is computed such that

$$\hat{c}_n^i(0.95) = \inf \left\{ c : \frac{1}{B} \sum_{b=1}^B \mathbf{1}\left\{n_b \sup_{\theta \in \hat{\Theta}_{0.95}^{(i-1)}} Q_{n,b}(\theta, \hat{\alpha}_n) \leq c\right\} \geq 0.95 \right\}.$$

Then a new estimate is computed: $\hat{\Theta}_{0.95}^i = \{\theta : n \cdot Q_n(\theta, \hat{\alpha}_n) < \hat{c}_n^i(0.95)\}$. The procedure is terminated when the set obtained from iteration i is identical to that obtained in iteration $i - 1$. This procedure will yield larger confidence regions than the procedure described in the text.

unobserved state variables are market specific. If a long panel is available, it may be sometimes possible to estimate the policy functions separately for each market. This would result in different value function estimates for each market. However, in the second-stage estimation, inequalities from all markets could be pooled together. Note that such an approach would also accommodate different equilibria in different markets. Of course, the main difficulty would be finding data sets large enough to make this approach feasible.

A related parametric approach could be used in the event that there were not enough data to obtain precise estimates for the market specific policy functions. If we were willing to assume a parametric form for the policy functions, then it would be possible to use standard panel data methods to estimate them. In that case it may even be possible to account for time-varying serially correlated unobservables.

Aguirregabiria and Mira (2007) suggested a fourth method. They noted that if the joint distribution between the observed and unobserved states can be recovered, this distribution can be used to control for the effect of the unobserved states in estimating the model. In our context, it seems that learning this joint distribution requires computing an equilibrium to the model. For the models we have in mind, this may not be computationally feasible.

Arcidiacono and Miller (2007) suggested a fifth method. In a dynamic discrete choice framework, they show that the EM algorithm can be used to maximize a likelihood function that integrates out over the unobserved states in the second stage. This novel approach even allows for time-varying serially correlated unobserved states. Because their method requires a likelihood function in the second stage, however, it would be difficult to apply directly to our model in general.

6. APPLYING THE ESTIMATION APPROACH

In this section we apply our general algorithm to two specific examples: a version of Rust's (1987) machine replacement model involving a single agent who makes a discrete choice each period, and a version of the Ericson–Pakes (1995) dynamic oligopoly model involving multiple firms that make both discrete (entry/exit) and continuous (price/investment) decisions in each period.

6.1. *Example: Dynamic Discrete Choice*

We first consider a simple version of Rust's (1987) machine replacement model as introduced in Section 2.1. In this single-agent dynamic discrete choice context, our approach is very similar to that of Hotz, Miller, Sanders, and Smith (1994), the difference being the second-stage procedure. The primary advantage of our second stage is that it extends easily to more general models such as the dynamic oligopoly example below, but the single-agent model provides a useful starting example.

Recall that in the model of Section 2.1, the manager's profit function,

$$\pi(a, s, \nu; \theta) = a[-R - \nu(1)] + (1 - a)[- \mu s - \nu(0)],$$

is linear in the parameters $\theta = (R, \mu)$. For our Monte Carlo exercise, we set $\mu = 1$, $R = 4$, and $\beta = 0.9$. We assume that the age of the machine takes values $s \in \{1, 2, 3, 4, 5\}$. Finally, we specify that $\nu(0)$ and $\nu(1)$ are distributed independently $N(0, 1)$ and this distribution is treated as known in the estimation. Using this parameterized model, we simulated data sets of (relatively small) sizes $n = 50, 100, 200$, and 400 and applied our estimation routine.

The first-stage estimates for this model are the average replacement probabilities in the simulated data for the five states. Recall that the optimal policy in this model has a cutoff feature. Let $\sigma(s, \nu)$ denote the optimal policy, let $V(s; \sigma)$ denote the value function, and define $v(1, s) = -R + \beta V(1; \sigma)$ and $v(0, s) = -\mu s + \beta V(\min\{s + 1, 5\}; \sigma)$ to be the choice-specific value functions. Then

$$(7) \quad \sigma(s, \nu) = 1 \quad \Leftrightarrow \quad v(1, s) + \nu(1) \geq v(0, s) + \nu(0).$$

Therefore, $\Pr(a = 1|s) = \Phi([v(1, s) - v(0, s)]/\sqrt{2})$, so we can invert the estimated choice probabilities and recover the difference in the choice-specific value functions and hence the policy rule at each state.

To simulate the value function, observe that it is linear in the parameters $\theta = (R, \mu)$. In particular,

$$\begin{aligned} V(s; \sigma; \theta) &= W^1(s; \sigma) \cdot R + W^2(s; \sigma) \cdot \mu + W^3(s; \sigma) \\ &= \mathbb{E} \left[- \sum_{t=0}^{\infty} \beta^t \sigma(s_t, \nu_t) \middle| s_0 = s \right] \cdot R \\ &\quad + \mathbb{E} \left[- \sum_{t=0}^{\infty} \beta^t (1 - \sigma(s_t, \nu_t)) s_t \middle| s_0 = s \right] \cdot \mu \\ &\quad + \mathbb{E} \left[\sum_{t=0}^{\infty} \beta^t \nu(\sigma(s_t, \nu_t)) \middle| s_0 = s \right]. \end{aligned}$$

The first term is the expected present cost of replacement decisions, the second term is the expected present cost of maintenance, and the third term is the expected present value of the realized preference shocks, where $\nu(\sigma(s_t, \nu_t))$ is the preference shock corresponding to the action that is actually chosen at time t . Note that for any given policy σ , each $W^m(s|\sigma)$ can be computed by simulation once and $V(s|\sigma, \theta)$ can be obtained immediately for a given θ .

To perform the second-stage estimation, we constructed alternative state-dependent cutoff rule policies by drawing cutoff points from a normal distribution with standard deviation 0.5 centered at the cutoff points estimated in

TABLE I
DYNAMIC DISCRETE CHOICE MONTE CARLO^a

	Mean	SE (Real)	5% (Real)	95% (Real)	SE (Subsampling)
<i>n</i> = 400, <i>n_I</i> = 200, <i>n_s</i> = 1000					
μ	1.00	0.14	0.79	1.24	0.10
<i>R</i>	4.02	0.53	3.24	4.96	0.39
<i>n</i> = 200, <i>n_I</i> = 200, <i>n_s</i> = 500					
μ	0.99	0.18	0.72	1.37	0.17
<i>R</i>	4.00	0.78	2.94	5.95	0.86
<i>n</i> = 100, <i>n_I</i> = 200, <i>n_s</i> = 250					
μ	0.94	0.32	0.47	1.48	0.35
<i>R</i>	3.75	1.26	1.92	5.70	1.15
<i>n</i> = 50, <i>n_I</i> = 200, <i>n_s</i> = 150					
μ	0.89	0.54	0.11	2.03	0.47
<i>R</i>	3.57	2.35	0.60	8.16	2.27

^a500 Monte Carlo runs; 25 subsamples of size *n*/2.

the first stage. We computed the estimator and standard errors 500 times for various values of *n_s* and *n_I*. The standard error estimates were obtained using 25 subsamples of size *n*/2.

The results, shown in Table I, show several things. First, the estimators are close to unbiased for all but the smallest sample sizes. Second, the subsampled standard errors are close to, though generally slightly smaller than, the true standard errors. The differences are most likely due to the small sample sizes used, such that the asymptotic approximations are imperfect. Finally, despite the small size of the data sets, the estimates are relatively precise. For *n* = 400, the *t* statistics are on the order of 6–8; for *n* = 100, they fall to about 3–4. Even for *n* = 50, *t* statistics are on the order of 2. These results are encouraging, particularly as many real-world data sets may be far above these sizes.

6.2. Example: Dynamic Oligopoly

We now apply the estimation approach to the dynamic oligopoly model outlined in Section 2.2. We first flesh out a few remaining details of the model and then describe estimation.

6.2.1. The dynamic oligopoly model and equilibrium

We start by positing specific functional forms for the demand system, the state transitions, and the cost functions. We assume a logit demand system for the goods. There is a mass *M* of consumers and consumer *r* derives utility *U_{ri}* from good *i*, where

$$U_{ri} = \gamma_0 \ln(z_i) + \gamma_1 \ln(y_r - p_i) + \varepsilon_{ri}.$$

Here z_i is the current quality of firm i 's product, p_i is the product's price, y_r is income, γ_0 and γ_1 are parameters, and ε_{rj} is an independent and identically distributed logit error term.¹⁹ For simplicity, all consumers have the same income $y_r = y$. We also assume firms have constant marginal costs of production equal to

$$\text{mc}(q_i; \mu) = \mu.$$

Each period, firms choose investment levels $I_{it} \in \mathbb{R}_+$ to increase their product quality the next period. We model the evolution of product quality following Pakes and McGuire (1994). Firm i 's investment is successful with probability

$$\rho I_{it} / (1 + \rho I_{it}),$$

where ρ is a parameter. If firm i 's investment is successful, its quality z_i increases by 1 and otherwise remains unchanged. There is also an outside good, whose quality moves up with probability δ each period. Firm i 's cost of investment is²⁰

$$C(I_i) = \xi \cdot I_i.$$

The scrap value realized on exiting, ϕ , is fixed and equal for all firms. Each period, the potential entrant draws a private entry cost ν_{et} from a uniform distribution on $[\nu^L, \nu^H]$.

In a Markov perfect equilibrium, each incumbent firm i sets its quantity to maximize its static profits, and uses an optimal investment strategy $I_i(\mathbf{s})$ and exit strategy $\chi_i(\mathbf{s})$. Each potential entrant follows a strategy $\chi_e(\mathbf{s}, \nu_e)$ that calls for it to enter if the expected profit from doing so exceeds its entry cost. The state variable $\mathbf{s}_t = (N_t, z_{1t}, \dots, z_{N_t}, z_{\text{out},t})$ includes number of incumbent firms and the current product qualities. We restrict attention to equilibria where the incumbents use symmetric strategies.

The model parameters are $\gamma_0, \gamma_1, \mu, \xi, \phi, \nu^L, \nu^H, \rho, \delta, \beta$, and y . As suggested above, the demand parameters γ_0, γ_1 and the marginal cost parameter μ can be estimated using static methods, so we do not include them in

¹⁹This simple demand specification does not include natural features such as an unobserved product characteristic. In principle this is a straightforward extension as long as one has an instrument such as an observed cost shifter to estimate the demand system. We do not consider it here because adding cost shifters to the model would substantially complicate the equilibrium computation needed to generate the Monte Carlo data.

²⁰We consider a very simple deterministic investment cost. As far as the estimation is concerned, it is straightforward to include a stochastic shock to the marginal cost of investment, but again this would substantially complicate the equilibrium computation required to generate the simulated data, so we are unable to consider it. For a detailed discussion of how our approach can be used to estimate the model when there is a stochastic shock to the marginal cost of investment, see Akerberg, Benkard, Berry, and Pakes (2007).

TABLE II
DYNAMIC OLIGOPOLY MONTE CARLO PARAMETERS

Parameter	Value	Parameter	Value
Demand		Discount factor β	0.925
γ_1	1.5	Investment cost ξ	1
γ_0	0.1	Marginal cost μ	3
M	5	Entry cost distribution	
y	6	ν^l	7
Investment evolution		ν^h	11
δ	0.7	Scrap value ϕ	6
ρ	1.25		

the vector of dynamic parameters θ . The remaining payoff parameters are $\theta = (\xi, \phi, \nu^L, \nu^H)$. The state transition parameters are ρ and δ . We assume the discount factor β and consumer income y are known to the econometrician.

For the Monte Carlo experiments, the parameters were set at the values shown in Table II. Computing an equilibrium for this model, which is required to generate the Monte Carlo data, is not trivial. To keep computation manageable, we considered a model in which a maximum of three firms could be active in each period. We then generated data sets of varying numbers of periods, from 100 to 400. Typically, real-world data sets would have more firms and fewer periods. Unfortunately, computing an equilibrium in which there are a large number of incumbent firms is prohibitive. Instead we chose a longer period length to make the number of firm-year observations comparable to available data sets.

6.2.2. Estimating the dynamic oligopoly model

The first stage requires estimation of the state transition probabilities and the policy functions; we also estimate the demand and marginal cost parameters at this stage. For the state transitions, we use the observed investment levels and product qualities to estimate the transition parameters ρ and δ by maximum likelihood. For the demand system, we similarly use maximum likelihood, and the observed quantity, price, and product quality data to estimate γ_0 and γ_1 , the parameters of the logit system. We then recover the marginal cost parameter μ from the static markup formula for optimal pricing. In general, the demand parameters are recovered very precisely and the investment evolution parameters are recovered somewhat less precisely.

Estimating the investment and exit policies is straightforward because the incumbent firms do not receive private investment or exit shocks in the simple numerical example we consider. Because not every possible state $\mathbf{s}$ is observed in the small data sets we consider, we use local linear regressions with a normal kernel to estimate both the investment policy $I(\mathbf{s})$ and the exit policy $\chi(\mathbf{s})$.

Because the state space is discrete, these estimates achieve parametric rates of convergence even though some smoothing is employed.

To estimate the equilibrium entry policy $\chi_e(\mathbf{s}, \nu_e)$, observe that it has a cutoff form,

$$(8) \quad \chi_e(\mathbf{s}, \nu_e) = 1 \quad \Leftrightarrow \quad \nu_e \leq \beta \mathbb{E}[V_e(\mathbf{s}_{t+1}; \boldsymbol{\sigma}) \mid \boldsymbol{\sigma}, \mathbf{s}_t = \mathbf{s}, \chi_e = 1],$$

where $V_e(\mathbf{s}_{t+1}; \boldsymbol{\sigma})$ is the value function that the entrant will have next period, as an incumbent, given the strategies $\boldsymbol{\sigma} = (I, \chi, \chi_e)$.

There is a variety of ways to handle entry in the estimation process. If the entry cost distribution G_e is known, the entry policy can be estimated using a Hotz–Miller approach: estimate the probability of entry at each state and invert the estimated probability to recover the value of entry. If only the parametric form of G_e is known, the same can be done conditional on parameter values. Because the value function of an entrant is the same as that of an incumbent, however, this is unnecessary. Instead, we simply estimate the probability of entry at each state, using local linear regression, and use this to calculate the incumbent value function. Two dynamic parameters, the investment cost ξ and the scrap value ϕ , enter the value function but not the entry cost distribution.

Given a strategy profile $\boldsymbol{\sigma} = (I, \chi, \chi_e)$, the incumbent value function is

$$\begin{aligned} V_i(\mathbf{s}; \boldsymbol{\sigma}) &= W^1(\mathbf{s}; \boldsymbol{\sigma}) + W^2(\mathbf{s}; \boldsymbol{\sigma}) \cdot \xi + W^3(\mathbf{s}; \boldsymbol{\sigma}) \cdot \phi \\ &= \mathbb{E} \left[\sum_{t=0}^{\infty} \beta^t \tilde{\pi}_i(\mathbf{s}_t) \mid \mathbf{s}_0 = \mathbf{s} \right] - \mathbb{E} \left[\sum_{t=0}^{\infty} \beta^t I_i(\mathbf{s}_t) \mid \mathbf{s}_0 = \mathbf{s} \right] \cdot \xi \\ &\quad + \mathbb{E} \left[\sum_{t=0}^{\infty} \beta^t \chi_i(\mathbf{s}_t) \mid \mathbf{s}_0 = \mathbf{s} \right] \cdot \phi. \end{aligned}$$

The first term $\tilde{\pi}_i(\mathbf{s}_t)$ is the static profit of incumbent i given that the current state is $\mathbf{s}_t$. This number is computed directly from the first-stage estimates of marginal cost and demand by calculating the static Bertrand–Nash pricing equilibrium. The second term is the expected present value of future investment costs. The third term is the expected present value of the scrap value on exit. Note that ξ and ϕ , the two parameters to be estimated, factor out linearly.

To apply the minimum distance estimator, we constructed alternative investment and exit policies by drawing a mean-zero normal error and adding it to the estimated first-stage investment and exit policies.²¹ We used $n_s = 2,000$

²¹So, for instance, if the estimated investment policy specifies an investment $I_i(\mathbf{s})$ in state $\mathbf{s}$, an alternative policy might be to invest $I_i(\mathbf{s}) + 0.01$ in state $\mathbf{s}$. The errors on the alternative investment policies had standard deviation 0.3 and on the alternative exit policies had standard deviation 0.5.

simulated paths, each having length at most 80 periods (shorter if, as was typical, exit occurred prior to the last period), to compute the present value W^1, W^2, W^3 terms for these alternative policies.

We experimented with two alternatives for estimating the distribution of entry costs. One approach we considered was to first apply the minimum distance estimator to estimate ξ and ϕ . Here we exploit the fact that G_e does not enter the incumbent value function. We then observe that given a state $\mathbf{s}$, the probability of entry equals

$$(9) \quad \Pr(\chi_e(\mathbf{s}, \nu_e) = 1) = G_e(\beta \mathbb{E}[V(\mathbf{s}'; \boldsymbol{\sigma}) \mid \boldsymbol{\sigma}, \mathbf{s}, \chi_e = 1]).$$

Using our estimates of ξ and ϕ , we computed the incumbent value function $V(\mathbf{s}' \mid \boldsymbol{\sigma})$ at different values of $\mathbf{s}$, and used this to compute the argument of G_e at $|S|$ points. Because the left-hand side at these points is just the observed probability of entry at state $\mathbf{s}$, we have identified $G_e(\cdot)$ on a discrete grid. We then used local linear regression to infer the remaining points.

Our other approach to estimating entry costs uses prior knowledge that G_e has a uniform distribution. Using this knowledge, we estimated the highest and lowest possible entry costs, ν_L and ν_H , along with the investment cost and scrap value parameters, in the second stage of estimation. To do this, we include optimality conditions for entry along with those for exit and investment in applying the minimum distance estimator. This approach has a potential benefit in that it uses information in the entry decisions to help estimate the investment and exit parameters.

To implement the Monte Carlo experiment, we generated 500 data sets and computed our estimator for each. For each data set, we used subsampling based on 20 subsamples of size $n/2$, to compute standard errors. The computational burden of the second-stage estimates was about a third that of the first-stage estimates. Together, each round of estimation took under a minute with little difference across sample sizes.

The results are shown in Tables III and IV. For small sample sizes, there is a slight bias in the estimates of the exit value, which is reduced with more observations. We found that this bias goes away entirely if the true first-stage functions are used instead of the estimated first stage, so we conclude that the second-stage bias is generated primarily by sampling error in the first-stage local linear estimates. The investment cost parameter estimates are essentially unbiased even for the smallest sample size ($n = 100$). The difference between the two most likely reflects the difference in the number of first-stage observations. For the $n = 100$ case, there were typically about 250 investment observations, but only about 25 entry or exit observations. Thus, for any given sample size, the investment policy function is substantially better estimated than the entry and exit policy functions.

Similar to the last example, the subsampled standard errors are, on average, slightly smaller than the true standard errors. This is again likely due to the

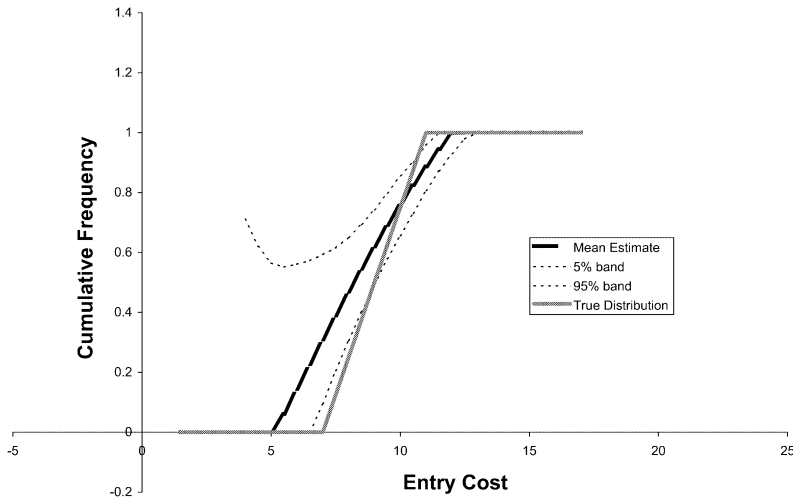
TABLE III
DYNAMIC OLIGOPOLY WITH NONPARAMETRIC ENTRY DISTRIBUTION

	Mean	SE (Real)	5% (Real)	95% (Real)	SE (Subsampling)
$n = 400, n_I = 500$					
ξ	1.01	0.05	0.91	1.10	0.03
ϕ	5.38	0.43	4.70	6.06	0.39
$n = 200, n_I = 500$					
ξ	1.01	0.08	0.89	1.14	0.05
ϕ	5.32	0.56	4.45	6.33	0.53
$n = 100, n_I = 300$					
ξ	1.01	0.10	0.84	1.17	0.06
ϕ	5.30	0.72	4.15	6.48	0.72

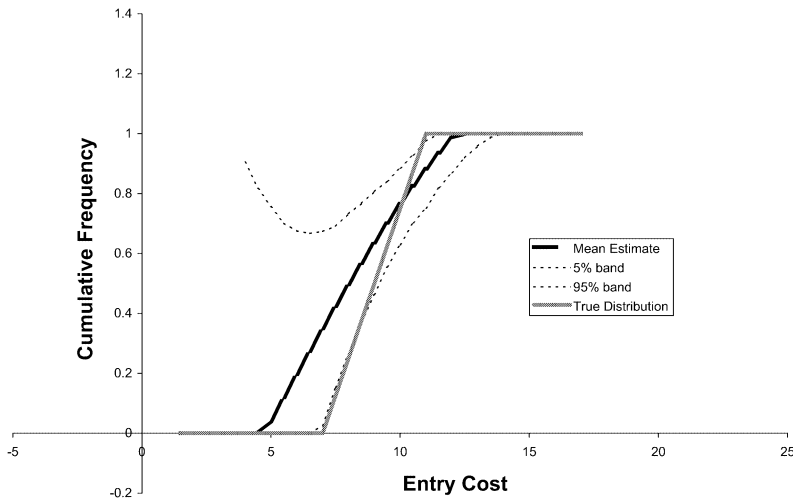
small sample sizes used. Overall the estimates are surprisingly precise given the small sample sizes. For $n = 400$, t statistics are on the order of 12–20. For the $n = 100$ case, in which there are very few observations of entry and/or exit, t statistics still average between 7 and 10. For these small data sets, the first-stage estimates are not very accurate pointwise. Thus, the second-stage estimation algorithm must be averaging across points in such a way as to come up with precise parameter estimates anyway. Because real-world data sets often contain more observations than this last case, we believe that these results support the method's potential in applications.

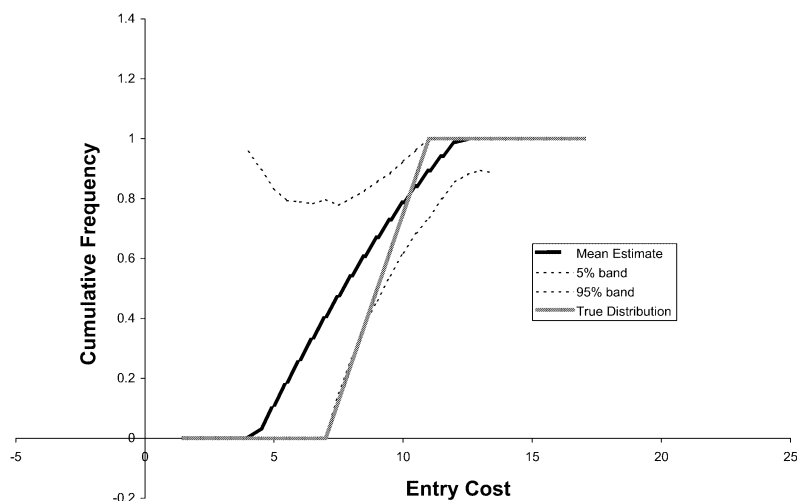
TABLE IV
DYNAMIC OLIGOPOLY WITH PARAMETRIC ENTRY DISTRIBUTION

	Mean	SE (Real)	5% (Real)	95% (Real)	SE (Subsampling)
$n = 400, n_I = 500$					
ξ	1.01	0.06	0.92	1.10	0.04
ϕ	5.38	0.42	4.68	6.03	0.41
ν^l	6.21	1.00	4.22	7.38	0.26
ν^h	11.2	0.67	10.2	12.4	0.30
$n = 200, n_I = 500$					
ξ	1.01	0.07	0.89	1.13	0.05
ϕ	5.28	0.66	4.18	6.48	0.53
ν^l	6.20	1.16	3.73	7.69	0.34
ν^h	11.2	0.88	9.99	12.9	0.40
$n = 100, n_I = 300$					
ξ	1.01	0.10	0.84	1.17	0.06
ϕ	5.43	0.81	4.26	6.74	0.75
ν^l	6.38	1.42	3.65	8.43	0.51
ν^h	11.4	1.14	9.70	13.3	0.58

FIGURE 1.—Entry cost distribution for $n = 400$.

Our nonparametric estimate of entry costs is displayed in Figures 1–3. The figures show that for all three sample sizes the entry cost distribution is recovered quite well. Table IV reports our entry cost estimates based on the second, parametric approach. Again the results are relatively accurate. It appears that one can recover the distribution of entry costs with reasonably sized data sets. Perhaps surprisingly, joint estimation did not improve the estimates of the in-

FIGURE 2.—Entry cost distribution for $n = 200$.

FIGURE 3.—Entry cost distribution for $n = 100$.

vestment and scrap parameters. This suggests that in this model entry behavior contains little information about these parameters.

6.2.3. *An alternative second-stage estimator*

The results above suggest that the inequality estimator may exhibit bias in small samples. This bias arises because the second-stage objective function is nonlinear in the first-stage estimates. Sampling error in the first stage, therefore, can bias the second-stage objective and through this, estimates of θ .

In a similar context, Pakes, Ostrovsky, and Berry (2007) suggested that one possible way to alleviate finite sample bias is to use a method of moments estimator based on highly aggregated moment conditions. The idea is that aggregation helps to average out the first-stage estimation error, reducing (but not eliminating) bias in the second-stage estimates. In this section, we explore how this approach applies in our setting.

The most obvious way to construct aggregated second-stage moment conditions in our model is to match the observed investment, exit, and entry rates to those predicted by the model when the model is evaluated using the estimated value functions. For each θ , we construct these moment conditions as follows:

- Step 1. Estimate the first stage (policy functions, transition probabilities, and demand and cost parameters).
- Step 2. For each state observed in the data and each state that is reachable in one period from a state observed in the data, construct simulated estimates of the value function $\hat{V}(s)$.
- Step 3. For each state observed in the data, use the value function estimates from Step 2 to compute the predicted entry probability, the pre-

dicted exit policy for each incumbent, and the predicted investment for each incumbent.

- Step 4. Construct moment conditions using the average over all states of the difference between the observed investment, entry, and exit at each state and those predicted by the model.
- Step 5. Estimate θ by minimizing a quadratic form in these moment conditions.

Note that, similarly to above, if V is linear in θ , then the simulations in Step 2 can be performed in advance once, and then Step 2 involves only multiplying together two vectors for each observed $\mathbf{s}$, for each θ .

As described in Step 4 above, this method provides only three aggregate moment conditions (entry, exit, and investment), allowing us to identify at most three parameters. In our Monte Carlo's experiments we have four parameters and thus require additional moment conditions. We thus add additional moment conditions based on the following criteria:

- (i) the covariances between exit and investment, and the firm's own state (two moment conditions);
- (ii) the covariances between entry, exit, and investment, and the sum of rival firms' states (three moment conditions).

This provides eight moment conditions in all. Note that additional moment conditions could also be constructed by interacting other functions of $\mathbf{s}$ with entry, exit, or investment. We use an optimal weight matrix when there are enough observations to compute these weights.

Results from this estimator are shown in Table V. Even at the smallest sample sizes, the bias in the aggregate moments estimator is much smaller than for the inequality estimator, particularly for the cost of investment parameter and for the lower bound on the entry cost. Standard errors are larger in some cases and smaller in others, so there is no obvious gain or loss in efficiency between the two estimators.

These results suggest that the aggregate moments estimator may have desirable finite sample properties compared to the inequality estimator. The aggregate moments approach does have some drawbacks in that it requires solving for the optimal investment policy, which may not always be as straightforward as it is in our model. A more realistic model also might allow for private investment costs, so one would have to compute the investment policy as a function of this shock and then simulate the moment conditions. In our particular example, investment has an analytic solution and there are no private investment costs, so these issues do not arise. Of course, the inequality estimator also has the advantage of applying easily to set-identified models, but in cases where the model is point identified and aggregate moments are easy to compute, we think the aggregate moments approach should be considered as well.

TABLE V
DYNAMIC OLIGOPOLY MODEL: METHOD OF MOMENTS ESTIMATES

	Mean	SE (Real)	5% (Real)	95% (Real)	SE (Subsampling)
$n = 400, n_I = 500$					
ξ	1.01	0.03	0.96	1.06	0.03
ϕ	6.09	0.18	5.86	6.36	0.05
ν^l	7.04	1.01	5.42	8.35	0.90
ν^h	11.1	0.65	9.97	12.0	0.57
$n = 200, n_I = 500$					
ξ	1.01	0.05	0.92	1.08	0.04
ϕ	6.06	0.24	5.74	6.38	0.10
ν^l	6.89	1.98	3.52	9.16	1.35
ν^h	11.0	1.54	9.21	12.5	0.91
$n = 100, n_I = 300$					
ξ	1.01	0.06	0.90	1.10	0.04
ϕ	6.02	0.34	5.43	6.48	0.15
ν^l	6.86	2.34	2.61	9.16	1.44
ν^h	11.1	1.83	8.73	13.58	1.09

7. CONCLUSION

This paper describes two new estimators applicable to a large class of dynamic environments. The estimators exploit the assumption that observed behavior is consistent with Markov perfect equilibrium. In that case, agents' beliefs can be recovered from observations of equilibrium play. Once agents' beliefs are known, the structural parameters can be solved for using the optimality conditions for equilibrium.

The biggest advantage of the approach is that it avoids the need for equilibrium computation. Avoiding equilibrium computation solves two problems. First, computing an equilibrium even once, for even the simplest of empirical dynamic oligopoly models, can be computationally prohibitive. In contrast, in our Monte Carlo experiments we found that the overall computational burden of the new estimators tends to be no more than that of many commonly used static estimation methods. Second, because equilibrium beliefs are obtained from the data, there is no need for the researcher to make assumptions about which of many potential equilibria is being played. These two benefits do come at some cost, because in avoiding equilibrium computations, some efficiency is compromised. However, the Monte Carlo experiments show that the approach still works quite well for fairly small data sets.

Both estimators are also conceptually straightforward and relatively easy to program in standard statistical packages. Given that there is currently a variety of options for estimating single-agent dynamic models, we expect that the estimators will be most useful for estimating dynamic games. In particular, we

hope our approach will facilitate future empirical work on dynamic oligopoly, such as the recent work of Ryan (2006).

Dept. of Economics, University of Minnesota, 271 19th Avenue South, Minneapolis, MN 55455, U.S.A., and NBER; bajari@econ.umn.edu,

Graduate School of Business, Stanford University, 518 Memorial Way, Stanford, CA 94305-5015, U.S.A., and NBER; lanierb@stanford.edu,

and

Dept. of Economics, Stanford University, 579 Serra Mall, Stanford, CA 94305-6072, U.S.A., and NBER; jdlevin@stanford.edu.

Manuscript received July, 2005; final revision received November, 2006.

APPENDIX: PROOFS OF PROPOSITIONS 1 AND 2

To shorten notation, let $h(y) = (\min\{y, 0\})^2$. The sample objective function is then

$$Q_n(\theta, \alpha) = \frac{1}{n_I} \sum_{k=1}^{n_I} h(\widehat{g}(X_k; \theta, \alpha)),$$

while the asymptotic objective function is given by

$$Q(\theta, \alpha_0) = \mathbb{E}_X h(g(X_k; \theta, \alpha_0)),$$

where X_k has distribution H on $\mathcal{X}$.

We will use the following properties of the above functions:

1. h has a continuous first derivative, its second derivative is continuous everywhere except at zero, and its third and higher derivatives are zero everywhere.
2. g and its derivatives with respect to α are smooth in θ .
3. $Q(\theta, \alpha_0)$, $Q_n(\theta, \alpha)$, and $(\partial Q_n(\theta, \alpha))/\partial \alpha$ are smooth in θ .

We will do the asymptotics in the number of first-stage observations, n . The reason for this is that the number of inequalities sampled, n_I , and the number of simulation draws per inequality, n_s , are both under the researcher's control.

LEMMA A1: $\sup_{\theta \in \Theta} |Q_n(\theta, \widehat{\alpha}_n) - Q(\theta, \alpha_0)| = o_p(1)$.

PROOF: Because Q_n and Q are continuous in θ , it suffices to show that $Q_n(\theta, \widehat{\alpha}_n) - Q(\theta, \alpha_0)$ is $o_p(1)$ for each θ :

$$\begin{aligned} & Q_n(\theta, \widehat{\alpha}_n) - Q(\theta, \alpha_0) \\ &= \frac{1}{n_I} \sum_{k=1}^{n_I} \{h(\widehat{g}(X_k; \theta, \widehat{\alpha}_n)) - h(\widehat{g}(X_k; \theta, \alpha_0))\} \end{aligned}$$

$$\begin{aligned}
& + \frac{1}{n_I} \sum_{k=1}^{n_I} \{h(\widehat{g}(X_k; \theta, \alpha_0)) - \mathbb{E}[h(\widehat{g}(X_k; \theta, \alpha_0))]\} \\
& + \frac{1}{n_I} \sum_{k=1}^{n_I} \{\mathbb{E}[h(\widehat{g}(X_k; \theta, \alpha_0))] - h(g(X_k; \theta, \alpha_0))\} \\
& + \frac{1}{n_I} \sum_{k=1}^{n_I} h(g(X_k; \theta, \alpha_0)) - \mathbb{E}_X h(g(X_k; \theta, \alpha_0)).
\end{aligned}$$

The first term, equal to $Q_n(\theta, \widehat{\alpha}_n) - Q_n(\theta, \alpha_0)$, represents the effect of using a preliminary estimator for α rather than the true value. By a mean-value expansion (Assumption S2(ii)),

$$Q_n(\theta, \widehat{\alpha}_n) - Q_n(\theta, \alpha_0) = \frac{\partial Q_n(\theta, \alpha_n^*)}{\partial \alpha} (\widehat{\alpha}_n - \alpha),$$

where $\widehat{\alpha}_n^* \rightarrow \alpha_0$ lies between $\widehat{\alpha}_n$ and α_0 . The derivative term satisfies a weak law of large numbers (WLLN), so this term is $o_p(1)$ by Assumption S1.

The second term represents the simulation error (the expectation is the expected simulation outcome). Each element of the sum is independent by Assumption S2(ii) and has expectation zero by construction. As $n_s \rightarrow \infty$ by Assumption S2(iii), this term is also $o_p(1)$.

The third term represents the simulation bias. Doing a mean-value expansion of one term in the sum gives

$$\begin{aligned}
& h(\widehat{g}(X_k; \theta, \alpha_0)) - h(g(X_k; \theta, \alpha_0)) \\
& = h'(g^*(X_k; \theta, \alpha_0))(\widehat{g}(X_k; \theta, \alpha_0) - g(X_k; \theta, \alpha_0)),
\end{aligned}$$

where g^* lies between g and $\widehat{g}$. This term has expectation zero because $\widehat{g}$ is unbiased (Assumption S2(ii)). Because each term is independent, the sum also satisfies a WLLN. Thus this term is also $o_p(1)$.

The last term represents the fact that all the inequalities are sampled only asymptotically. It satisfies a WLLN as long as $Q(\theta, \alpha_0)$ exists and is finite, so the last term is $o_p(1)$ given that $n_I \rightarrow \infty$. Q.E.D.

PROOF OF PROPOSITION 1: For consistency, observe that by our identification assumption, if $|\hat{\theta} - \theta_0| > \delta > 0$, then there exists some $\varepsilon(\delta) > 0$ such that $Q(\hat{\theta}, \alpha_0) > \varepsilon$, so

$$\Pr(|\hat{\theta} - \theta_0| > \delta) \leq \Pr(Q(\hat{\theta}, \alpha_0) > \varepsilon(\delta)).$$

Now observe that

$$\begin{aligned}
Q(\hat{\theta}, \alpha_0) &= Q(\hat{\theta}, \alpha_0) - Q_n(\hat{\theta}, \widehat{\alpha}_n) + Q_n(\hat{\theta}, \widehat{\alpha}_n) \\
&\leq |Q(\hat{\theta}, \alpha_0) - Q_n(\hat{\theta}, \widehat{\alpha}_n)| + Q_n(\hat{\theta}, \widehat{\alpha}_n)
\end{aligned}$$

$$\begin{aligned}
&\leq \sup_{\theta \in \Theta} |Q(\theta, \alpha_0) - Q_n(\theta, \hat{\alpha}_n)| + Q_n(\hat{\theta}, \hat{\alpha}_n) \\
&= o_p(1) + Q_n(\hat{\theta}, \hat{\alpha}_n) \\
&\leq o_p(1) + Q_n(\theta_0, \hat{\alpha}_n) = o_p(1).
\end{aligned}$$

It follows that $\hat{\theta}_n \rightarrow^p \theta_0$.

We now consider the asymptotic distribution of $\hat{\theta}_n$. Differentiating $Q_n(\theta, \hat{\alpha}_n)$, evaluating at $\theta = \theta_0$, and pre-multiplying by $\sqrt{n}$ gives

$$\begin{aligned}
&\sqrt{n} \frac{\partial}{\partial \theta} Q_n(\theta_0; \hat{\alpha}_n) \\
&= \frac{\sqrt{n}}{n_I} \sum_{k=1}^{n_I} \left(\frac{\partial}{\partial \theta} h(\hat{g}(X_k; \theta_0, \hat{\alpha}_n)) - \frac{\partial}{\partial \theta} h(\hat{g}(X_k; \theta_0, \alpha_0)) \right) \\
&\quad + \frac{\sqrt{n}}{n_I} \sum_{k=1}^{n_I} \left(\frac{\partial}{\partial \theta} h(\hat{g}(X_k; \theta_0, \alpha_0)) - \mathbb{E} \frac{\partial}{\partial \theta} h(\hat{g}(X_k; \theta_0, \alpha_0)) \right) \\
&\quad + \frac{\sqrt{n}}{n_I} \sum_{k=1}^{n_I} \left(\mathbb{E} \frac{\partial}{\partial \theta} h(\hat{g}(X_k; \theta_0, \alpha_0)) - \frac{\partial}{\partial \theta} h(g(X_k; \theta_0, \alpha_0)) \right) \\
&\quad + \frac{\sqrt{n}}{n_I} \sum_{k=1}^{n_I} \frac{\partial}{\partial \theta} h(g(X_k; \theta_0, \alpha_0)).
\end{aligned}$$

These four terms are analogous to those in Lemma A1.

The first term is the first-stage sampling error term. Doing an element by element mean-value expansion of the first term gives the expression

$$\begin{aligned}
&\frac{\sqrt{n}}{n_I} \sum_{k=1}^{n_I} \left(\frac{\partial}{\partial \theta} h(\hat{g}(X_k; \theta_0, \hat{\alpha}_n)) - \frac{\partial}{\partial \theta} h(\hat{g}(X_k; \theta_0, \alpha_0)) \right) \\
&= \left(\frac{1}{n_I} \sum_{k=1}^{n_I} \frac{\partial^2 h(\hat{g}(X_k; \theta_0, \alpha_n^*))}{\partial \theta \partial \alpha'} \right) \sqrt{n}(\hat{\alpha}_n - \alpha_0).
\end{aligned}$$

By a WLLN, consistency of $\hat{\alpha}_n$, and consistency of $\hat{g}$, then as n_s and n_I go to ∞ with n ,

$$\begin{aligned}
&\left(\frac{1}{n_I} \sum_{k=1}^{n_I} \frac{\partial^2 h(\hat{g}(X_k; \theta_0, \alpha_n^*))}{\partial \theta \partial \alpha'} \right) \xrightarrow{p} \Lambda_0 \equiv \frac{\partial^2}{\partial \theta \partial \alpha'} \mathbb{E}_X h(g(X_k; \theta_0, \alpha_0)) \\
&= \frac{\partial^2}{\partial \theta \partial \alpha'} Q(\theta_0, \alpha_0).
\end{aligned}$$

Note that the derivative above can be brought outside the integral because the function inside is bounded. Thus, under these assumptions the first term has limiting distribution $N(0, \Lambda_0 V_\alpha \Lambda_0)$.

The second term is the simulation variance term. It is mean zero by construction and is the sum of independent terms (Assumptions S2(ii)). Therefore, a central limit theorem applies and the second term is asymptotically normal with rate $\sqrt{n_I}$ and variance matrix of the form S_{n_I}/n_s that disappears with n_s . As n_I and n_s go to infinity with n (Assumptions S2(iii)), this term contributes nothing to the asymptotic variance (this would be true even for fixed n_s).

The third term is the simulation bias term. Doing a second-order mean-value expansion of $h(\widehat{g})$ around g for one element in the sum for the third term gives

$$\begin{aligned} & \frac{\partial}{\partial \theta_l} h(\widehat{g}(X_k; \theta_0, \alpha_0)) - \frac{\partial}{\partial \theta_l} h(g(X_k; \theta_0, \alpha_0)) \\ &= \frac{\partial}{\partial \theta} \frac{\partial}{\partial \theta'} h(g(X_k; \theta_0, \alpha_0)) \times (\widehat{g}(X_k; \theta_0, \alpha_0) - g(X_k; \theta_0, \alpha_0)) \\ & \quad + \frac{1}{2} \frac{\partial^2}{\partial \theta \partial \theta'} \frac{\partial}{\partial \theta_l} h(g^*(X_k; \theta_0, \alpha_0)) (\widehat{g}(X_k; \theta_0, \alpha_0) - g(X_k; \theta_0, \alpha_0))^2, \end{aligned}$$

where g^* lies between $\widehat{g}$ and g . Note that the derivatives of h on the right-hand side exist with probability 1. Taking the expectations with respect to the simulation error gives

$$\begin{aligned} & \mathbb{E} \frac{\partial}{\partial \theta_l} \widehat{h}(X_k; \theta_0, \alpha_0) - \frac{\partial}{\partial \theta_l} h(X_k; \theta_0, \alpha_0) \\ &= \frac{1}{2} \frac{\partial^2}{\partial \theta \partial \theta'} \frac{\partial}{\partial \theta_l} h(g^*(X_k; \theta_0, \alpha_0)) \text{Var}(\widehat{g}(X_k; \theta_0, \alpha_0)). \end{aligned}$$

Since $g^* \rightarrow g$, this term goes to zero at rate n_s . Therefore, as long as n_s goes to infinity faster than $\sqrt{n}$, the third term contributes nothing to the asymptotic variance.

The last term is the standard term in the first-order expansion. Note that in this case it is always zero (since h is always zero at the true values of the parameters) and thus drops out completely. This is because of the lack of sampling error in the second stage.

Putting all four terms together gives

$$\sqrt{n} \frac{\partial}{\partial \theta} Q_n(\theta_0, \widehat{\alpha}_n) \xrightarrow{d} N(0, \Lambda_0 V_\alpha \Lambda_0).$$

Under Assumption S2(v), the asymptotic distribution of the estimator is given by

$$\sqrt{n}(\widehat{\theta} - \theta_0) \xrightarrow{d} N(0, B_0^{-1}(\Lambda_0 V_\alpha \Lambda_0') B_0^{-1}). \quad Q.E.D.$$

PROOF OF PROPOSITION 2: The first part follows from Lemma A1. The second part follows from Proposition 5b in Manski and Tamer (2002). *Q.E.D.*

REFERENCES

- ACKERBERG, D., C. L. BENKARD, S. T. BERRY, AND A. PAKES (2007): "Econometric Tools for Analyzing Market Outcomes," in *Handbook of Econometrics*, Vol. 6, ed. by J. J. Heckman and E. Leamer. Amsterdam: Elsevier. [1333,1356]
- AGUIRREGABIRIA, V., AND P. MIRA (2002): "Swapping the Nested Fixed Point Algorithm: A Class of Estimators for Discrete Markov Decision Models," *Econometrica*, 70, 1519–1543. [1333]
- (2007): "Sequential Estimation of Dynamic Discrete Games," *Econometrica*, 75, 1–53. [1333,1350,1353]
- ARCIDIACONO, P., AND R. A. MILLER (2007): "CCP Estimation of Dynamic Discrete Choice Models with Unobserved Heterogeneity," Mimeo, Duke University. [1353]
- BAJARI, P., AND C. L. BENKARD (2004): "Demand Estimation with Heterogeneous Consumers and Unobserved Product Characteristics: A Hedonic Approach," Working Paper W10278, NBER. [1351]
- BAJARI, P., V. CHERNOZHUKOV, AND H. HONG (2005): "Semiparametric Estimation of a Dynamic Game of Incomplete Information," Mimeo, Duke University. [1346,1347]
- BENKARD, C. L. (2004): "A Dynamic Analysis of the Market for Wide-Bodied Commercial Aircraft," *Review of Economic Studies*, 71, 581–611. [1331,1339]
- BERESTEANU, A., AND P. B. ELLICKSON (2006): "The Dynamics of Retail Oligopoly," Mimeo, Duke University. [1333]
- BERRY, S. T. (1994): "Estimating Discrete Choice Models of Product Differentiation," *Rand Journal of Economics*, 25, 242–262. [1352]
- BERRY, S. T., AND A. PAKES (2000): "Estimation from the Optimality Conditions for Dynamic Controls," Mimeo, Yale University. [1333]
- BERRY, S. T., J. LEVINSOHN, AND A. PAKES (1995): "Automobile Prices in Market Equilibrium," *Econometrica*, 63, 841–890. [1339,1352]
- BRESNAHAN, T. F., AND P. REISS (1991): "Empirical Models of Discrete Games," *Journal of Econometrics*, 81, 57–81. [1334]
- CHERNOZHUKOV, V., H. HONG, AND E. TAMER (2007): "Estimation and Confidence Regions for Parameter Sets in Econometric Models," *Econometrica*, 75, 1243–1284. [1350,1351]
- CHING, A. (2005): "Consumer Learning and Heterogeneity: Dynamics of Demand for Prescription Drugs after Patent Expiration," Mimeo, University of Toronto. [1332]
- CILIBERTO, F., AND E. TAMER (2006): "Market Structure and Multiple Equilibria in Airline Markets," Mimeo, Northwestern University. [1334,1347]
- DORASZELSKI, U., AND M. SATTERTHWAIT (2007): "Computable Markov-Perfect Industry Dynamics: Existence, Purification and Multiplicity," Mimeo, Harvard University. [1337]
- ERICSON, R., AND A. PAKES (1995): "Markov-Perfect Industry Dynamics: A Framework for Empirical Work," *Review of Economic Studies*, 62, 53–83. [1331,1334,1337,1353]
- GEWEKE, J., AND M. KEANE (2001): "Computationally Intensive Methods for Integration in Econometrics," in *Handbook of Econometrics*, Vol. 5, ed. by J. J. Heckman and E. Leamer. Amsterdam: Elsevier Science, 3463–3568. [1332]
- GOWRISANKARAN, G., AND R. TOWN (1997): "Dynamic Equilibrium in the Hospital Industry," *Journal of Economics and Management Strategy*, 6, 45–74. [1331]
- HAILE, P. A., AND E. TAMER (2003): "Inference with an Incomplete Model of English Auctions," *Journal of Political Economy*, 111, 1–51. [1334,1347,1350,1351]
- HOLMES, T. J. (2006): "The Diffusion of Wal-Mart and Economies of Density," Mimeo, University of Minnesota. [1348]

- HOTZ, V. J., AND R. A. MILLER (1993): "Conditional Choice Probabilities and the Estimation of Dynamic Models," *Review of Economic Studies*, 60, 497–529. [1332,1340,1341,1346]
- HOTZ, V. J., R. A. MILLER, S. SANDERS, AND J. SMITH (1994): "A Simulation Estimator for Dynamic Models of Discrete Choice," *Review of Economic Studies*, 61, 265–289. [1334,1340,1342,1353]
- IMAI, S., N. JAIN, AND A. CHING (2005): "Bayesian Estimation of Dynamic Discrete Choice Models," Mimeo, University of Toronto. [1332]
- IMBENS, G., AND C. MANSKI (2004): "Confidence Intervals for Partially Identified Parameters," *Econometrica*, 72, 1845–1857. [1352]
- JOFRE-BONET, M., AND M. PESENDORFER (2003): "Estimation of a Dynamic Auction Game," *Econometrica*, 71, 1443–1489. [1333]
- MAGNAC, T., AND D. THESMAR (2002): "Identifying Dynamic Discrete Decision Processes," *Econometrica*, 70, 801–816. [1347]
- MANSKI, C. F., AND E. TAMER (2002): "Inference in Regressions with Interval Data on a Regressor or Outcome," *Econometrica*, 70, 519–546. [1347,1350,1351,1369]
- MOTZKIN, T. S., H. RAIFFA, G. L. THOMPSON, AND R. M. THRALL (1953): "The Double Description Method," in *Contributions to Theory of Games*, Vol. 2, ed. by H. W. Kuhn and A. W. Tucker. Princeton, NJ: Princeton University Press. [1351]
- OLLEY, G. S., AND A. PAKES (1996): "The Dynamics of Productivity in the Telecommunications Equipment Industry," *Econometrica*, 64, 1263–1298. [1352]
- PAKES, A., AND P. MCGUIRE (1994): "Computing Markov-Perfect Nash Equilibria: Numerical Implications of a Dynamic Differentiated Product Model," *Rand Journal of Economics*, 25, 555–589. [1331,1337,1356]
- (2001): "Stochastic Approximation for Dynamic Models: Markov Perfect Equilibrium and the 'Curse' of Dimensionality," *Econometrica*, 69, 1261–1281. [1331]
- PAKES, A., M. OSTROVSKY, AND S. BERRY (2007): "Simple Estimators for the Parameters of Discrete Dynamic Games (with Entry/Exit Examples)," *Rand Journal of Economics*, forthcoming. [1333,1362]
- PESENDORFER, M., AND P. SCHMIDT-DENGLER (2003): "Identification and Estimation of Dynamic Games," Working Paper 9726, NBER. [1333,1334,1347]
- (2007): "Asymptotic Least Squares Estimators for Dynamic Games," Mimeo, London School of Economics. [1334]
- ROMANO, J. P., AND A. M. SHAIKH (2006a): "Inference for Identifiable Parameters in Partially Identified Econometric Models," Mimeo, University of Chicago. [1350-1352]
- (2006b): "Inference for Identified Set in Partially Identified Econometric Models," Mimeo, University of Chicago. [1352]
- RUST, J. (1987): "Optimal Replacement of GMC Bus Engines: An Empirical Model of Harold Zurcher," *Econometrica*, 55, 999–1033. [1332,1334,1337,1353]
- (1994): "Structural Estimation of Markov Decision Processes," in *Handbook of Econometrics*, Vol. 4, ed. by R. F. Engle and D. L. McFadden. Amsterdam: Elsevier Science, 3082–3139. [1340,1346]
- RYAN, S. (2006): "The Costs of Environmental Regulation in a Concentrated Industry," Mimeo, Massachusetts Institute of Technology. [1333,1365]
- RYAN, S., AND C. TUCKER (2007): "Heterogeneity and the Dynamics of Technology Adoption," Mimeo, MIT. [1333]
- SWEETING, A. (2006): "The Costs of Product Repositioning: The Case of Format Switching in the Commercial Radio Industry," Mimeo, Northwestern University. [1333]