

Optimal Replenishment Policies for Multiechelon Inventory Problems Under Advance Demand Information

Guillermo Gallego • Özalp Özer

Department of Industrial Engineering and Operations Research, Columbia University, New York, New York 10027

Department of Management Science and Engineering, Stanford University, Stanford, California

guillermo.gallego@columbia.edu • ozalp.ozel@stanford.edu

Customers and downstream supply chain partners often place, or can be induced to place, orders in advance of future requirements. We show how to optimally incorporate advance demand information into periodic-review, multiechelon, inventory systems in series. While the state space for series systems appears to be formidably large, we decompose the problem across locations, as in Clark and Scarf (1960), and reduce the state space at each location by using modified echelon inventory positions (that nets known requirements). We prove the optimality of state-dependent, echelon base-stock policies for finite and infinite horizon problems. We also show that myopic policies are optimal and very easy to compute when costs and demands are stationary. We provide managerial insights into the value of advance demand information through a numerical study.

(*Multiechelon; Stochastic Inventory System; Facilities-in-Series; Advance Demand Information*)

1. Introduction

There is a growing consensus that a portfolio of customers with different demand leadtimes, see Gressens and Brousseau (1999), can lead to higher, more regular revenues and better capacity utilization. Customers with positive demand leadtimes place orders in advance of their needs resulting in *advance demand information*. This gives rise to the problem of finding effective inventory control policies under advance demand information. Sellers may elicit advance demand information from buyers by providing incentives for early bookings. In this paper we take the incentive system as given and study the seller's operational problem of minimizing the discounted holding and penalty costs of managing a series system over a finite or infinite horizon, proving the optimality of state-dependent, echelon base-stock policies. We also show that myopic policies are optimal and very easy to compute when costs and demands are station-

ary. Our results can then be used to evaluate strategies to obtain advance demand information.

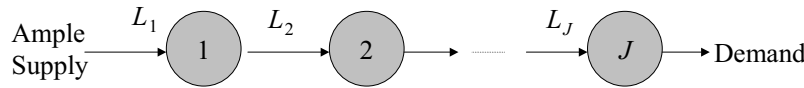
In particular, we analyze a periodic-review, single-item series system that incorporates advance demand information. Under advance demand information the demand seen during period t is of the form

$$D_t = (D_{t,t}, \dots, D_{t,t+N}) \quad (1)$$

where $D_{t,s}$ is the demand realized in period t for delivery in period $s \in \{t, \dots, t+N\}$, and where $N < \infty$ is the *information horizon* beyond which we do not collect advance demand information. Notice that $N = 0$ represents the classical case of no advance demand information. D_t becomes known with certainty at the end of period t .

As mentioned before, advance demand information may arise as buyers' response to price incentives. Buyers may be given a choice of demand leadtimes, $l \in \{0, \dots, N\}$, at decreasing unit costs c_l . In this case $D_{t,s}$ is the sum of the demands from buyers with demand

Figure 1 Serial System



leadtime $s - t \in \{0, 1, \dots, N\}$. Alternatively, the seller may offer buyers a menu of prices c_l with delivery leadtime $l \in \{0, \dots, N\}$. Each buyer needs to decide how many units to order at each price. This interesting problem was first studied by Fukuda (1964) for the case $N = 1$. For a stationary demand process, Fukuda's results can be interpreted as follows: If X_t is the demand process seen by a buyer, then $D_t = ((X_t - \Delta)^+, \min(\Delta, X_t))$ where Δ is a constant that increases with the price discount $c_0 - c_1$. Thus up to Δ units of the buyer's demand is purchased at c_1 and only the excess over Δ is purchased at c_0 . If there is more than one buyer, then the advance demand information should be aggregated across the buyers. Notice that in both cases the components of D_t may be dependent, but the vectors D_t are time independent if the buyer's demands are time independent.

The series system consists of J locations as illustrated in Figure 1. Location J satisfies the demands of one or more buyers with different demand leadtimes that collectively give rise to the advance demand information vector D_t . Each location satisfies its requirement from its immediate predecessor. Location 1 orders from an outside supplier with ample stock. Replenishment decisions are centralized and based on systemwide information. Orders from the outside supplier and shipments between locations arrive after fixed, location specific, leadtimes L_j . Unsatisfied demand at each location is fully backordered. A penalty cost is charged on backorders at location J only. Holding costs are charged on ending inventories at each location as well as on inventories in-transit to the downstream stage. The cost parameters and demands are allowed to be time dependent. We assume stationary cost and demands for the infinite horizon problem.

There is an extensive literature on *stationary*¹ series systems with independent demands without advance

demand information. This literature starts with the seminal paper of Clark and Scarf (1960). They show that the problem decomposes into single location problems and that echelon base-stock policies are optimal. This model was extended to the infinite horizon stationary case by Federgruen and Zipkin (1984). They show that base-stock policies are optimal for both the discounted and the average cost criterion. For a review of this literature we refer the reader to Federgruen (1993).

Chen and Song (2001) wrote the only paper, that we are aware of, to study infinite horizon average cost series system with a *nonstationary* demand process modulated by a finite state exogenous Markov chain. Markov modulated demand processes were first addressed by Song and Zipkin (1993) and Sethi and Cheng (1997) in the context of single location problems. Similar to the proofs of Chen and Zheng (1994), Chen and Song (2001) establish lower bounds on the cost of managing the average cost infinite horizon inventory problem in series. They construct feasible state-dependent policies that achieve these lower bounds, thereby proving their optimality. Later in the paper we discuss how a Markov modulated demand process can be combined with advance demand information for a finite horizon problem. The papers by Erkip et al. (1990), Song and Zipkin (1992, 1996), and Chen et al. (2000) also address nonstationary demand processes but for the one-warehouse multiretailer setting. Unlike Chen and Song (2001) and our paper, these papers focus on the performance of the system under a given policy or optimize within a given set of policies.

In contrast to multilocation, periodic-review inventory problems, a fair amount of research has been done on the analysis of single location inventory problems that incorporate the dynamic nature of demand updates. There are mainly four groups: those who make use of Bayesian updates (Scarf 1960 and Azuory 1985), those who incorporate time-series models (Johnson and Thompson 1975, Miller 1989,

¹ We refer to an inventory problem as stationary when the demand is stationary and the cost parameters are constant over time.

and Lovejoy 1990), those who are concerned with forecast revisions (Heath and Jackson 1994, Güllü 1996, Toktay and Wein 2001), and those who deal with state-dependent policies where the demand is governed by either an external process or updated by advance demand information (Song and Zipkin 1993, Sethi and Cheng 1997, Schwarz et al. 1998, Scheller-Wolf and Tayur 2001, Kapuscinski and Tayur 1998, Gallego and Özer 2001).

We end this review by noting that series systems are fundamental to the study of more general supply networks, including assembly systems as discussed in Rosling (1989) and distribution systems as in Federgruen (1993) and Özer (2003), who addresses a warehouse multiretailer distribution system with advance demand information.

In this paper we prove the optimality of state-dependent, echelon base-stock policies for both finite and infinite horizon α -discounted problems with advance demand information. To do this we show that the problem decomposes into single location periodic-review problems. We further reduce the dimension of each location's dynamic program by using, at each stage, the *modified inventory position*² concept introduced in Gallego and Özer (2001). To obtain the state-dependent, echelon base-stock levels for each location, we first solve the dynamic program for location J , which is a single location problem with advance demand information and linear ordering costs. We then use the optimal cost function for location J to obtain an *implicit* penalty cost function that measures the inability of location $J-1$ to respond to the requirements of location J . This implicit penalty cost function, together with location $J-1$ holding and ordering costs form the basis for the single period cost function of location $J-1$ and the single location dynamic-programming formulation. The solution to this DP provides an optimal echelon base-stock level for location $J-1$. We repeat this procedure and solve for locations $J-2, \dots, 2$ and end by solving the problem corresponding to location 1.

We also show that myopic base-stock policies are optimal for finite and infinite horizon problems when

demands and cost parameters are stationary. For the myopic result to hold with or without advance demand information, we require carefully set terminal conditions for finite horizon problems. See Veinott (1965) and the references therein for the optimality of myopic policies for single location problems. When the myopic policy is optimal, there is no need to incorporate advance demand information beyond each location's leadtime plus a review period.

The rest of the paper is organized as follows: In §2, we introduce the demand model. In §3, we analyze the single location model that is fundamental to the analysis of series systems. In §4, we carry out the analysis for the series system and provide an algorithm to compute the optimal policies. In §5, we establish the optimality of myopic policies for stationary problems. In §6, we extend the results to infinite horizon problems. In §7, we provide insights through a numerical study. We conclude in §8 and point out directions for future research.

2. The Demand Model

The demand model given by Equation (1) can incorporate several elements of randomness that are likely to arise in practice. In particular, the number of customers that arrive in a period, the order sizes, and the desired delivery dates can be all random. To see this, let N_t be the random number of customers that arrive during period t . Assume that customer j demands a random quantity $X_{t,j}$ to be delivered at a random fulfillment date $s_{t,j} \in \{t, \dots, t+N\}$. Then $D_{t,s} = \sum_{j=1}^{N_t} X_{t,j} I(s_{t,j} = s)$ where $I(\cdot)$ is an indicator function. Hariharan and Zipkin (1995) call $s_{t,j} - t$ the demand, or customer, leadtime and study a single location continuous-review model with constant customer leadtimes. Our model can be viewed as periodic-review generalization of Hariharan and Zipkin (1995) to the case of random demand leadtimes.

Given the advance demand information process $D_t = (D_{t,t}, \dots, D_{t,t+N})$, we can decompose the total demand $\sum_{r=s-N}^s D_{r,s}$ for period s at the beginning of period $t \in \{s-N, \dots, s\}$ into the *observed* and *unobserved* parts:

$$O_{t,s} = \sum_{r=s-N}^{t-1} D_{r,s}, \quad U_{t,s} = \sum_{r=t}^s D_{r,s}.$$

² The modified inventory position is the inventory on hand plus on order minus backorders minus the observed demand over the leadtime.

Notice that $O_{t,s}$ is known while $U_{t,s}$ is random at the beginning of period t . Notice also that $O_{t+1,s} = O_{t,s} + D_{t,s}$ and $U_{t+1,s} = U_{t,s} - D_{t,s}$ at the beginning of period $t+1$ after D_t is realized.

ASSUMPTION 1. *We assume that the advance demand information process D_t is time independent. The components of the demand vector $D_{t,t+i}$ and $D_{t,t+j}$ are, however, allowed to be dependent.*

This implies that $U_{t,s}$ is independent of $O_{t,s}$. However, the *total* demand for period s clearly depends on $O_{t,s}$. Indeed, the best estimate of the demand for period s at the beginning of period t is given by $O_{t,s} + \sum_{r=t}^s E[D_{r,s}]$. As discussed in the introduction, there are important cases where $U_{t,s}$ is independent of $O_{t,s}$. While it may be desirable to generalize the model to allow $U_{t,s}$ to depend on $O_{t,s}$, this requires an enlargement of the state space as discussed at the end of §3.

Notice that under this information structure at the beginning of period t , we know the N dimensional vector $\tilde{O}_t = (O_{t,t}, O_{t,t+1}, \dots, O_{t,t+N-1})$. We use the $\tilde{O}_t$ notation to emphasize that we are including all the components $O_{t,s}$, $s \in \{t, \dots, t+N\}$. Later, we will subsume the information contained in the first $L+1$ components of $\tilde{O}_t$ and use the notation O_t to refer to the resulting vector of reduced dimension. For convenience, we define $O_{t,s} \equiv 0$ for $s \geq t+N$. For easy reference we provide a glossary of notation in Appendix A.

2.1. MRP Serial Systems

Our model also bridges a gap between the classical stochastic inventory literature and MRP serial systems. The classical inventory system corresponds to the case $N=0$ where there is no advance demand information. On the other hand, MRP logic assumes future requirements are known with certainty. In practice, however, MRP is applied on a rolling horizon basis and forecasts for future requirements are updated in every period. Our model makes the forecast updates explicit. In particular, if we let $F_{t,s}$ be the forecast at the beginning of period t for period s demand, then

$$F_{t+1,s} = F_{t,s} + D_{t,s} - E_t[D_{t,s}]. \quad (2)$$

Notice that $D_{t,s} - E_t[D_{t,s}]$ is a mean zero random variable, so the forecast updates form a martingale. See Graves et al. (1986) and Heath and Jackson (1994) for an elegant exposition of the Martingale method of forecast evolution (MMFE). Graves et al. (1998) use the MMFE model to smooth production by making production updates a linear function of forecast updates while keeping safety stocks constant. Their paper provides a sensible heuristic for our model where the forecast updates may arise from factors other than advance demand information. The state-dependent base-stock policies obtained in this paper provide an optimal solution for a series MRP system where forecast updates arise from advance demand information, leadtimes are constant, and the objective is to minimize expected holding and backorder costs. To our knowledge, no other model has this capability.

2.2. Markov Modulated Demand and Advance Demand Information

The demand model can be extended to incorporate a fluctuating demand environment governed by an exogenous parameter, such as season and economic condition. In particular, it is possible to model the demand vector D_t as a state-dependent process that is governed by a discrete time-finite state Markov chain $\{\theta_t, t \geq 0\}$ as in Song and Zipkin (1993). In this case the distribution of the demand vector depends on the realization of θ_t . This generalization requires augmenting the state space of the dynamic programs to include the state of the Markov chain. The results and proofs for this case would follow the generalizations of the proofs in the paper.

We end this section by clarifying a few differences between MMD and ADI systems. For MMD systems, θ_t is an exogenous parameter that is independent of past demands and governs the distribution of D_{t+1} . For ADI systems, $\tilde{O}_t$ is an endogenous vector that depends on the demand history but does not influence D_{t+1} . One can envision a richer class of MMD systems with endogenous Markov chains and a richer class of ADI systems where $\tilde{O}_t$ influences future

demands. To our knowledge, such systems have not been studied in detail.

3. Single Location Systems

We start by considering the dynamics of the single location, periodic-review inventory system. This problem is addressed in our earlier paper (Gallego and Özer 2001). Here we summarize the results that are necessary to study the series system. At the beginning of the period, the order placed a leadtime earlier, L periods ago, arrives. Based on the available information, the inventory manager decides the order quantity $z_t \geq 0$ and incurs a linear cost of $c_t z_t$. Demands over the period are satisfied, giving priority to existing backorders (if any). Excess demand is backlogged. Inventory holding and penalty costs are charged to the inventory level at the end of the period. The objective is to minimize the expected discounted cost of managing the system over a finite or infinite horizon.

The information available to the manager at the beginning of period t is given by I_t , the inventory on hand; B_t , the number of backorders; z_s for $s \in \{t-L+1, \dots, t-1\}$, the pipeline inventory and $\tilde{O}_t = (O_{t,t}, \dots, O_{t,t+N-1})$ the vector of observed demands. Let $\tilde{x}_t \equiv I_t + \sum_{s=t-L+1}^{t-1} z_s - B_t$, the inventory position before the ordering decision is made.

Let $\tilde{y}_t = \tilde{x}_t + z_t$. The net inventory at the end of period $t+L$ is given by

$$\tilde{y}_t - \sum_{s=t}^{t+L} O_{t,s} - \sum_{s=t}^{t+L} U_{t,s}.$$

The single period cost charged to period t is the discounted expected holding and penalty cost at the end of period $t+L$, and it is given by

$$G_t \left(\tilde{y}_t - \sum_{s=t}^{t+L} O_{t,s} \right) = \alpha^L E g_{t+L} \left(\tilde{y}_t - \sum_{s=t}^{t+L} O_{t,s} - \sum_{s=t}^{t+L} U_{t,s} \right),$$

where $\alpha \leq 1$ is the discount factor³ and $g_t(\cdot)$ is the holding and penalty cost function based on the net inventory at the end of period t . The expectation is taken with respect to the unobserved part of the lead-time demand.

³ We require $\alpha < 1$ for infinite horizon problems.

ASSUMPTION 2. We assume the cost function G_t is well defined, e.g., $g_t(x) \leq c(1 + |x|)$ and $ED_{t,s}^\rho < \infty$ for some $\rho > 1$ (regularity condition), that g_t is convex and $\lim_{|x| \rightarrow \infty} g_t(x) = \infty$ (coercive) for all t . This assumption is satisfied, for example, when holding and penalty costs are linear and imply that G_t is convex and coercive for all t .

ASSUMPTION 3. We assume that inventory (respectively backlogs) at the end of the planning horizon are salvaged (respectively purchased) at a unit rate c_{T+1} .

The dynamic programming formulation of this problem is given by

$$\tilde{V}_t(\tilde{x}_t, \tilde{O}_t) = -c_t \tilde{x}_t + \min_{y \geq \tilde{x}_t} \tilde{H}_t(y, \tilde{O}_t), \quad (3)$$

$$\begin{aligned} \tilde{H}_t(y, \tilde{O}_t) = & c_t y + G_t \left(y - \sum_{s=t}^{t+L} O_{t,s} \right) \\ & + \alpha E \tilde{V}_{t+1}(\tilde{x}_{t+1}, \tilde{O}_{t+1}) \end{aligned} \quad (4)$$

$$\tilde{V}_{T+1}(\tilde{x}_{T+1}, \cdot) \equiv -c_{T+1} \tilde{x}_{T+1}.$$

We propose the reduced state space (x_t, O_t) where

x_t : modified inventory position before the ordering decision, $= \tilde{x}_t - \sum_{s=t}^{t+L} O_{t,s}$,

O_t : vector of observed demands beyond the lead-time, that is beyond period $t+L$, $= (O_{t,t+L+1}, \dots, O_{t,t+N-1})$.

The dimension of the reduced state space is $(N-L-1)^+ + 1$. In terms of the reduced state space, the dynamic program is given by

$$V_t(x_t, O_t) = -c_t x_t + \min_{y \geq x_t} H_t(y, O_t), \quad (5)$$

$$H_t(y, O_t) = c_t y + G_t(y) + \alpha E V_{t+1}(x_{t+1}, O_{t+1}), \quad (6)$$

$$V_{T+1}(x_{T+1}, \cdot) \equiv -c_{T+1} x_{T+1}.$$

We define

$\tilde{y}_t(\tilde{O}_t)$: smallest minimizer of $\tilde{H}_t(\cdot, \tilde{O}_t)$ defined in Equation (4),

$y_t(O_t)$: smallest minimizer of $H_t(\cdot, O_t)$ defined in Equation (6),

y_t^m : smallest minimizer of $(c_t - \alpha c_{t+1})y + G_t(y)$,

y^m : smallest minimizer of $(1 - \alpha)cy + G(y)$.

All of the above minimizers exist under Assumption 2. Notice that the definition of y^m assumes stationary demand and costs. Both y_t^m and y^m ignore

observed demand information and the cost incurred in upcoming periods. We refer to policies that order up to y_t^m or y^m as myopic.

The following Lemma shows that the reduced state space comes at no loss of optimality. Theorem 1 summarizes some of the results in Gallego and Özer (2001). We refer the reader to this paper and to the Appendix for the rest of the proofs.

LEMMA 1. *For any given vector $\tilde{O}_t$, (i) $\tilde{V}_t(\tilde{x}_t, \tilde{O}_t) = V_t(x_t, O_t)$, (ii) $\tilde{H}_t(\tilde{y}_t, \tilde{O}_t) = H_t(y_t, O_t) + c_t \sum_{s=t}^{t+L} O_{t,s}$, where $\tilde{y}_t = y_t + \sum_{s=t}^{t+L} O_{t,s}$, (iii) $\tilde{y}_t(\tilde{O}_t) = y_t(O_t) + \sum_{s=t}^{t+L} O_{t,s}$.*

THEOREM 1. *If Assumption 2 holds, then for any given vector O_t we have*

(1) *Finite Horizon Problem:*

(i) $H_t(x, O_t)$ is convex.

(ii) *The state-dependent base-stock policy given by $y_t(O_t)$ is an optimal solution to the dynamic program given in Equation (5).*

(iii) $V_t(x, O_t) + c_t x$ is nondecreasing convex in x .

(2) *Infinite Horizon Problem Under Stationary Costs and Demands:*

(i) $\lim_{T \rightarrow \infty} V_t(\cdot, O_t)$ exists and converges uniformly to a convex function $V(\cdot, O_t)$. Furthermore, $\lim_{x \rightarrow \infty} V(x, O_t) = \infty$.

(ii) $\lim_{T \rightarrow \infty} H_t(\cdot, O_t)$ exists and converges uniformly to a convex function $H(\cdot, O_t)$. Let $y(O_t)$ be a minimizer.

(iii) $y(O_t) = y^m$ is an optimal policy for the infinite horizon problem.

(3) *Myopic Policies: If y_t^m is nondecreasing in t , then the myopic policy y_t^m is optimal for finite horizon problems. In particular, under stationary costs and demand distributions the myopic policy is optimal for finite as well as for infinite horizon problems.*

We remark that it is also possible to establish Theorem 1 in terms of the original state space of functional Equation (3). This allows us to compare the costs under the original and the reduced state space. This is also used for the decomposition proof of multilocation series system.

We now argue that the state space reduction discussed above is not plausible for a number of models

that make the unobserved part of the demand dependent on the observed part. Consider first the autoregressive model $D_{t,s} = a_{s-t} + \rho_{s-t} D_{t-1,s} + \epsilon_{ts}$. Here $O_{t,s}$ is not a sufficient statistic to compute the distribution of $U_{t,s}$ because it hides the value of $D_{t-1,s}$. Therefore a larger state space is required. Consider now a parsimonious model where $U_{t,s}$ depends directly on the observed part $O_{t,s} = \sum_{z=s-N}^{s-1} D_{z,s}$. This imposes strong distributional assumptions on $D_{r,s}$ for $r \in \{t, \dots, s\}$ because $U_{t,s} = \sum_{r=t}^s D_{r,s}$. Finally, as discussed in §2.2, modeling D_t as a state-dependent Markov chain would make the unobserved part dependent on the observed part. Note that in this case the state dependency of the policy would be due to both O_t and the state of the Markov chain.

4. Series Systems

We analyze a series multilocation, periodic-review inventory system under advance demand information. At the beginning of each period, after receiving previously placed orders and/or scheduled shipments, the inventory manager decides how much to order, $z_1 \geq 0$, from the outside supplier and how much to ship, $z_j \geq 0$, to location $j \in \{2, \dots, J\}$ from location $j-1$ in light of the advance demand information. A linear ordering/shipping cost $\sum_{j=1}^J c_{jt} z_j$ is charged to period t . This cost should be interpreted with care. c_{1t} is the acquisition cost from the external supplier, while c_{jt} , $j \geq 2$, are echelon, or value added, costs from shipping or transforming units from one stage to the next. We assume that an order from the outside supplier placed at the beginning of period t arrives at location 1 at the beginning of period $t + L_1$. Similarly, a shipment to location $j \in \{2, \dots, J\}$ arrives L_j periods later.

At the beginning of period t , the inventory manager knows the number of on-hand inventory I_{jt} at each location j , the backorders B_t at location J , and trans-shipments $\bar{z}_{jt} = (z_{j1}, \dots, z_{jL'_j})$ to each location j , where z_{js} is the shipment dispatched from location $j-1$ to location j at the beginning of period $t-s$ and $L'_j \equiv L_j - 1$. We assume without loss of generality that the on-hand inventory includes shipments placed a leadtime earlier and received at the beginning of the period. Under this convention, $\bar{z}_{jt}$ is

irrelevant whenever $L_j \in \{0, 1\}$. In addition to these variables, the inventory manager also knows $\tilde{O}_t = (O_{t,t}, \dots, O_{t,t+N-1})$.

Let us define

$$\hat{x}_{jt} = \sum_{i \geq j} I_{it} + \sum_{i > j} \sum_{s=t}^{t+L'_i} z_{is} - B_t,$$

echelon net inventory at location j ,

$$\tilde{x}_{jt} = \hat{x}_{jt} + \sum_{s=t}^{t+L'_j} z_{js},$$

echelon inventory position at location j ,

$$x_{jt} = \tilde{x}_{jt} - \sum_{s=t}^{t+L_j} O_{t,s}$$

modified echelon inventory position at location j ,

for each location $j \in \{1, \dots, J\}$. Notice that $\hat{x}_{jt} = I_{jt} - B_t$ and $\hat{x}_{jt} = \tilde{x}_{j+1,t} + I_{jt}$ for $j \in \{1, \dots, J-1\}$.

The holding and penalty costs are based on the inventory levels at the end of the period. We assume that the echelon holding costs are strictly positive. This is consistent with the idea that value is added as the item goes through the locations. The holding and penalty cost for period t is given by

$$\begin{aligned} & \sum_{j=1}^{J-1} h'_{jt} [\hat{x}_{j,t+1} - \hat{x}_{j+1,t+1}] + h'_{Jt} [\hat{x}_{J,t+1}]^+ + p_t [\hat{x}_{J,t+1}]^- \\ &= \sum_{j=1}^J h_{jt} \hat{x}_{j,t+1} + [p_t + h'_{Jt}] [\hat{x}_{J,t+1}]^- \end{aligned}$$

where $[x]^+ = \max(x, 0)$ and $[x]^- = \max(-x, 0)$, p_t is the penalty cost at location J , h'_{jt} is the local inventory holding cost and $h_{jt} = h'_{jt} - h'_{j-1,t}$ ($h'_{0t} = 0$) is the echelon holding cost at location j . Notice that shipments in transit to location $j+1$ are charged at location j 's holding cost rate.

The updates, at the end of period t , after observing the demand vector D_t , are given by

$$\tilde{O}_{t+1} = (O_{t+1,t+1}, \dots, O_{t+1,t+N}),$$

$$\tilde{z}_{j,t+1} = (z_{j,t+1}, \dots, z_{j,t+L'_j-1}),$$

$$\hat{x}_{j,t+1} = \hat{x}_{jt} + z_{jL'_j} - O_{t,t} - D_{t,t},$$

$$\tilde{x}_{j,t+1} = \tilde{x}_{jt} + z_j - O_{t,t} - D_{t,t},$$

$$x_{j,t+1} = x_{jt} + z_j - \sum_{s=t}^{t+L_j+1} D_{t,s} - O_{t,t+L_j+1}.$$

ASSUMPTION 4. We assume that costs continue to accrue up to period T . We assume a linear terminal condition of the form $-\sum_{j=1}^J c_{j,T+1} \tilde{x}_{j,T+1}$. The economic interpretation is that of a salvage value if $\tilde{x}_{j,T+1}$ is positive, and an acquisition cost if $\tilde{x}_{j,T+1}$ is negative. All costs after time $T+1$ are assumed to be zero. We take advance information up to the period T only, hence, $\tilde{O}_{T+1} = 0$.

Notice that locations $j < J$ do not need to wait until the beginning of period $T+1$ to salvage remaining inventories. Our formulation, which salvages remaining inventory at the beginning of period $T+1$, is without loss of generality provided that the unit costs $c_{j,T+1}$ are appropriately discounted, e.g., $\alpha^{L_j+\dots+L_1} c_{j,T+1} = c_{j,T+1-L_j-\dots-L_1}$. Notice also that at the beginning of period $T+1$, the expressions $\hat{x}$, $\tilde{x}$, and x are equivalent because there is no pipeline inventory and there is no observed demand over future periods.

To our knowledge all other papers in the literature charge zero terminal costs for the multiechelon systems. Linear terminal costs are more realistic and allow us to show that myopic policies are optimal for finite horizon problems with stationary costs and demand distributions, a result that fails to hold under zero terminal costs.

To simplify the exposition we formulate the dynamic programming recursion for a series system with two locations. We discuss how to extend the results to a series system with more than two locations at the end of this section and in Appendix B.

4.1. Series System with Two Locations

The problem here is to manage two locations in series up to period $T > L_1 + L_2$. The last order for location 1 is dispatched at the beginning of period $T - L_1 - L_2$ and arrives at location 1 at the beginning of period $T - L_2$. The last shipment to location 2 is initiated at the beginning of period $T - L_2$ and arrives at location 2 at the beginning of period T .

We define the initial state space by $(\hat{x}_{1t}, \tilde{z}_{1t}, \hat{x}_{2t}, \tilde{z}_{2t}, \tilde{O}_t)$. Later we show how to reduce the dimension

of the state space. The optimal cost-to-go starting from period t is given by

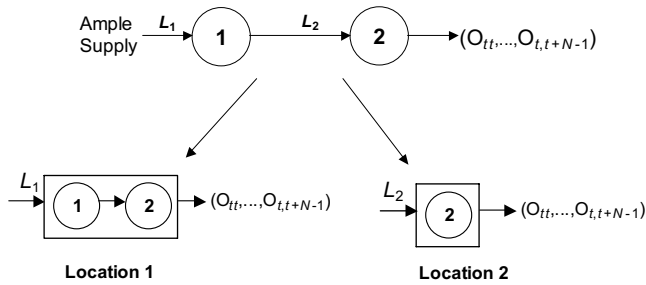
$$\begin{aligned} & \hat{J}_t(\hat{x}_{1t}, \bar{z}_{1t}, \hat{x}_{2t}, \bar{z}_{2t}, \tilde{O}_t) \\ &= \min_{(z_1, z_2) \in \mathcal{A}} \{c_{1t}z_1 + c_{2t}z_2 + h_{1t}E\hat{x}_{1,t+1} + Eg_t(\hat{x}_{2,t+1}) \\ & \quad + \alpha E\hat{J}_{t+1}(\hat{x}_{1,t+1}, \bar{z}_{1,t+1}, \hat{x}_{2,t+1}, \bar{z}_{2,t+1}, \tilde{O}_{t+1})\} \quad (7) \end{aligned}$$

where $\hat{J}_{T+1}(\hat{x}_1, \cdot, \hat{x}_2, \cdot, \cdot) \equiv -c_{1,T+1}\hat{x}_1 - c_{2,T+1}\hat{x}_2$, $g_t(x) = h_{2t}[x] + (p_t + h'_{2t})[x]^-$, and $\mathcal{A} = \{(z_1, z_2) \in Z^2: z_1 \geq 0, z_2 \geq 0, \text{ and } \tilde{x}_{2t} + z_2 \leq \hat{x}_{1t}\}$. The shipment constraint $\tilde{x}_{2t} + z_2 \leq \hat{x}_{1t}$ is equivalent to $z_2 \leq I_{1t}$, meaning that shipments to location 2 are bounded by the inventory on hand at location 1. Notice that g_t is convex and coercive for all $t \leq T$.

In the remainder of this section we show that the above dynamic program can be decomposed into two simpler problems. To do this, we first show that the state space of Equation (7) reduces to $(\hat{x}_{1t}, \bar{z}_{1t}, \tilde{x}_{2t}, \tilde{O}_t)$ by subsuming pipeline inventory, that is, $\tilde{x}_{2t} = \hat{x}_{2t} + \sum_{s=t}^{L_2} z_{2s}$. This yields an essentially equivalent dynamic program with a reduced state space and cost-to-go $J_t(\hat{x}_{1t}, \bar{z}_{1t}, \tilde{x}_{2t}, \tilde{O}_t)$. We then show that this dynamic program can be decomposed into two simpler dynamic programs plus a cost function that is independent of the decision variables z_1 and z_2 :

$$\begin{aligned} & J_t(\hat{x}_{1t}, \bar{z}_{1t}, \tilde{x}_{2t}, \tilde{O}_t) \\ &= \tilde{V}_t^1(\tilde{x}_{1t}, \tilde{O}_t) + \tilde{V}_t^2(\tilde{x}_{2t}, \tilde{O}_t) + \mathcal{F}_t(\hat{x}_{1t}, \bar{z}_{1t}, \tilde{O}_t). \end{aligned}$$

We further reduce the dimension of each of these programs as in Lemma 1. We conclude by establishing the optimality of *state-dependent, echelon base-stock policies* and describe how to compute them.



A shipment of z_2 initiated at the beginning of period t arrives at location 2 at the beginning of

period $t + L_2$. Consequently, the inventory on hand at location 2 at the end of period $t + L_2$ is given by

$$\tilde{x}_{2t} + z_2 - \left(\sum_{s=t}^{t+L_2} O_{t,s} + \sum_{s=t}^{t+L_2} U_{t,s} \right),$$

where the term in parenthesis is the leadtime demand for location 2, that is, the demand during periods $\{t, \dots, t + L_2\}$. We charge to period t the cost

$$h_{1t}E\hat{x}_{1,t+1} + G_t \left(\tilde{x}_{2t} + z_2 - \sum_{s=t}^{t+L_2} O_{t,s} \right),$$

where $G_t(x) = \alpha^{L_2} Eg_{t+L_2}(x - \sum_{s=t}^{t+L_2} U_{t,s})$. This cost accounting scheme is standard in the inventory literature. It assists in collapsing the state space, and in reducing the problem to

$$\begin{aligned} & J_t(\hat{x}_{1t}, \bar{z}_{1t}, \tilde{x}_{2t}, \tilde{O}_t) \\ &= \min_{(z_1, y) \in \mathcal{A}} \left\{ c_{1t}z_1 + c_{2t}(y - \tilde{x}_{2t}) + h_{1t}E\hat{x}_{1,t+1} \right. \\ & \quad \left. + G_t \left(y - \sum_{s=t}^{t+L_2} O_{t,s} \right) \right. \\ & \quad \left. + \alpha E J_{t+1}(\hat{x}_{1,t+1}, \bar{z}_{1,t+1}, \tilde{x}_{2,t+1}, \tilde{O}_{t+1}) \right\} \end{aligned}$$

where $J_{T+1}(\hat{x}_1, \cdot, \tilde{x}_2, \cdot) \equiv -c_{1,T+1}\hat{x}_1 - c_{2,T+1}\tilde{x}_2$ and $\mathcal{A} = \{(z_1, y) \in Z^2: z_1 \geq 0 \text{ and } \tilde{x}_{2t} \leq y \leq \hat{x}_{1t}\}$. The justification for this formulation is given by the following lemma.

LEMMA 2. $\hat{J}_t(\hat{x}_{1t}, \bar{z}_{1t}, \hat{x}_{2t}, \bar{z}_{2t}, \tilde{O}_t) = J_t(\hat{x}_{1t}, \bar{z}_{1t}, \tilde{x}_{2t}, \tilde{O}_t) + \mathcal{G}_t$.

The term $\mathcal{G}_t$, described in Appendix C, is a constant. It is independent of the period t decision variables, z_1 and y , so it can be dropped for optimization purposes. Next, we decompose the series system into two single location problems.

THEOREM 2. $J_t(\hat{x}_{1t}, \bar{z}_{1t}, \tilde{x}_{2t}, \tilde{O}_t) = \tilde{V}_t^1(\tilde{x}_{1t}, \tilde{O}_t) + \tilde{V}_t^2(\tilde{x}_{2t}, \tilde{O}_t) + \mathcal{F}_t$ where

$$\tilde{V}_t^2(\tilde{x}_{2t}, \tilde{O}_t) = -c_{2t}\tilde{x}_{2t} + \min_{\tilde{x}_{2t} \leq y} \tilde{H}_t^2(y, \tilde{O}_t), \quad (8)$$

$$\tilde{H}_t^2(y, \tilde{O}_t) = c_{2t}y + G_t \left(y - \sum_{s=t}^{t+L_2} O_{t,s} \right)$$

$$+ \alpha E \tilde{V}_{t+1}^2(\tilde{x}_{2,t+1}, \tilde{O}_{t+1}),$$

$$\tilde{V}_{T+1}^2(\tilde{x}_2, \cdot) \equiv -c_{2,T+1}\tilde{x}_2$$

and

$$\tilde{V}_t^1(\tilde{x}_{1t}, \tilde{O}_t) = -c_{1t}\tilde{x}_{1t} + \min_{y \geq \tilde{x}_{1t}} \left\{ c_{1t}y + \tilde{C}_t \left(y - \sum_{s=t}^{t+L_1} O_{t,s}, \tilde{O}_t \right) \right. \\ \left. + \alpha E \tilde{V}_{t+1}^1(\tilde{x}_{1,t+1}, \tilde{O}_{t+1}) \right\}, \quad (9)$$

$$\tilde{C}_t(y, \tilde{O}_t) = \alpha^{L_1} E \left\{ h_{1,t+L_1} \left(y - \sum_{s=t}^{t+L_1} U_{t,s} \right) + \alpha \tilde{I}P_{t+L_1+1} \right. \\ \left. \cdot \left(y - \sum_{s=t}^{t+L_1} U_{t,s}, \tilde{O}_{t+L_1+1} \right) \right\},$$

$$\tilde{I}P_t(y, \tilde{O}_t) = \tilde{H}_t^2(\min\{\tilde{y}_{2t}(\tilde{O}_t), y\}, \tilde{O}_t) - \tilde{H}_t^2(\tilde{y}_{2t}(\tilde{O}_t), \tilde{O}_t), \\ \tilde{V}_{T+1}^1(\tilde{x}_1, \cdot) \equiv -c_{1,T+1}\tilde{x}_1,$$

$\tilde{y}_{2t}(\tilde{O}_t)$ is the smallest minimizer of $\tilde{H}_t^2(\cdot, \tilde{O}_t)$ and $\mathcal{F}_t$ is a constant.

Here $\tilde{I}P_t$ denotes the implicit penalty cost function as in Clark and Scarf (1960). It measures the inability of location 1 to match the requirements of location 2 and depends on the state of the advance demand information in addition to the echelon inventory positions. So far we showed how to decompose Equation (7) into two smaller dimensional dynamic programs. We further reduce the state space of the dynamic program $\tilde{V}_t^2(\tilde{x}_{2t}, \tilde{O}_t)$ by subsuming the observed part of the state-specific leadtime demand⁴ and introducing the modified inventory position concept as in Lemma 1. It is not clear, however, whether a state-space reduction for location 2's dynamic program will preclude us from reducing the state space of location 1's dynamic program. Recall that the implicit penalty cost function is based on the location 2's cost function $\tilde{V}_t^2$, hence the state space of $\tilde{I}P_t$ depends on that of $\tilde{V}_t^2$. We show next why this is not a problem.

As in Lemma 1, the dynamic program for location 2 in Equation (8) is equivalent to

$$V_t^2(x_{2t}, O_t^2) = -c_{2t}x_{2t} + \min_{x_{2t} \leq y} H_t^2(y, O_t^2), \quad (10) \\ H_t^2(y, O_t^2) = c_{2t}y + G_t(y) + \alpha E V_{t+1}^2(x_{2,t+1}, O_{t+1}^2), \\ V_{T+1}^2(x, \cdot) \equiv -c_{2,T+1}x, \\ y_{2t}(O_t^2) : \text{smallest minimizer of } H_t^2(\cdot, O_t^2).$$

⁴ the first L_2 components of vector $\tilde{O}_t$

where $O_t^2 \equiv (O_{t,t+L_2+1}, \dots, O_{t,t+N-1})$. Because of Assumption 4, no demand information is recorded for periods after T , so $x_{2t} = \tilde{x}_{2t}$ for $t = T + 1$.

We argue that the implicit penalty cost function can be written in terms of the reduced state space for location 2. Let $IP_t(x, O_t^2) = H_t^2(\min\{x, y_{2t}(O_t^2)\}, O_t^2) - H_t^2(y_{2t}(O_t^2), O_t^2)$. We use this implicit penalty cost function and the next lemma to reduce the state space of location 1's dynamic program.

LEMMA 3. $IP_t(x, O_t) = \tilde{I}P_t(\tilde{x}, \tilde{O}_t)$ where $x = \tilde{x} - \sum_{s=t}^{t+L_2} O_{t,s}$ and $O_t = (O_{t,t+L_1+1}, \dots, O_{t,t+N-1})$ for all fixed $L \in \{0, 1, \dots, N-1\}$.

We can now assert that $\tilde{I}P_t(\tilde{x}, \tilde{O}_t) = IP_t(\tilde{x} - \sum_{s=t}^{t+L_2} O_{t,s}, O_t^2)$. In other words, we can obtain $\tilde{I}P_t(\tilde{x}_t, \tilde{O}_t)$, to be used in the location 1's new dynamic program, from $V_t^2(x, O_t^2)$. Hence, the DP for location 1 is independent of the state-space reduction for location 2. For the evaluation of $\tilde{C}_t(y, \tilde{O}_t)$, notice that $O_t^1 \equiv (O_{t,t+L_1+1}, \dots, O_{t,t+N-1})$ contains all the relevant information about $\tilde{O}_{t+L_1+1}$ that is available at the beginning of time t . We apply Lemma 1 to reduce the DP for location 1 to

$$V_t^1(x_{1t}, O_t^1) = -c_{1t}x_{1t} + \min_{y \geq x_{1t}} H_t^1(y, O_t^1), \quad (11) \\ H_t^1(y, O_t^1) = c_{1t}y + C_t(y, O_t^1) + \alpha E V_{t+1}^1(x_{1,t+1}, O_{t+1}^1), \\ C_t(y, O_t^1) = \alpha^{L_1} E \left\{ h_{1,t+L_1} \left(y - \sum_{s=t}^{t+L_1} U_{t,s} \right) + \alpha \tilde{I}P_{t+L_1+1} \right. \\ \left. \cdot \left(y - \sum_{s=t}^{t+L_1} U_{t,s}, \tilde{O}_{t+L_1+1} \right) \right\}, \\ V_{T+1}^1(x, \cdot) \equiv -c_{1,T+1}x, \\ y_{1t}(O_t^1) : \text{smallest minimizer of } H_t^1(y, O_t^1).$$

LEMMA 4. $C_t(\cdot, O_t^1)$ is convex and $\lim_{|x| \rightarrow \infty} C_t(x, O_t^1) = \infty$.

To summarize, the problem of finding an optimal policy reduces to solving two simpler, single location, dynamic programs. The state space for these programs is of dimension $1 + (N - L_j - 1)^+$ for $j = 1, 2$. We are now in a position to present a simple algorithm to compute an optimal policy. The *first step* of this algorithm involves solving the single location dynamic program for location 2, given by Equation (10). Notice that this recursion has the same form as Equation (5).

In particular, the functions $G_t(\cdot)$ satisfy the conditions of Theorem 1. The *second step* involves solving the single location dynamic program for location 1, given by (11). This requires the implicit penalty cost function that is now available from the solution to the dynamic program for the second location. Lemma 4 and Theorem 1 imply the optimality of state-dependent base-stock policy for location 1. We summarize our results in the following theorem.

THEOREM 3. *An echelon state-dependent base-stock policy is optimal. In particular, an optimal base-stock level for location j at time t is given by $y_{jt}(O_t^j)$ for $j = 1, 2$.*

Note that the decomposition of the series systems into stages should be done in a way that preserves advance demand information for use at all the decomposed stages. Notice that the implicit penalty cost depends on the advance demand information. On the other hand, if the inventory manager is only concerned with incorporating advance demand information up to the minimum of the leadtimes plus one review period, that is, $N \leq \min(L_1, L_2) + 1$, then the state space for each location is one-dimensional. To see this, notice that for $N \leq L_2 + 1$ the vector O_t^2 disappears, so the state space, the implicit penalty cost function, and C_t are all of dimension one. This result and $N \leq L_1 + 1$ imply that the DP for location 1 is also of dimension 1. As a consequence, *all* the classical results and algorithms for series systems apply after modifying the inventory position, at each stage, to subsume the observed part of the leadtime demand.

4.2. Series System with $J > 2$ Locations

The proof of this case is an extension of the results we obtained for two locations in series. We apply Lemma 2 and Theorem 2 recursively to decompose the system into single location problems. Optimality of base-stock policies then follows directly from Lemma 4 and Theorem 1. The outline of the proof is in Appendix B.

5. Myopic Policies

From Theorem 1 we know that myopic policies are optimal for single location problems whenever the sequence of myopic base-stock levels is nondecreasing. In particular, the myopic policy is optimal for

finite and infinite horizon problems with stationary costs and demand distributions. In this section we show that myopic policies are also optimal for two-location series system for finite horizon problems with stationary costs and demand distributions. We refer to Veinott (1965) and the references therein for finite horizon single location models. In this section we show that the myopic policy is optimal for a *finite horizon multiechelon inventory problem in series* with and without advance demand information. We also show that for a nonstationary problem the optimal base-stock level is bounded.

ASSUMPTION 5. *We assume that $h_{1t} = h_1$, $c_{1t} = c_1$ for all $t \leq T - L_1 - L_2$, and $h_{2t} = h_2$, $c_{2t} = c_2$ and $p_t = p$ for $t \leq T - L_2$.*

Next, we define the myopic cost functions

$$\mathcal{L}^2(y) = (1 - \alpha)c_2y + G(y), \quad (12)$$

$$IP^m(x) = \mathcal{L}^2(\min\{y_2^m, x\}) - \mathcal{L}^2(y_2^m),$$

$$C^m(x) = \alpha^{L_1} E \left[h_1 \left(x - \sum_{s=t}^{t+L_1} U_{t,s} \right) + \alpha IP^m \left(x - \sum_{s=t}^{t+L_1} U_{t,s} \right) \right],$$

$$\mathcal{L}^1(y) = (1 - \alpha)c_1y + C^m(y),$$

where y_i^m is the smallest minimizer of $\mathcal{L}^i(y)$ for $i = \{1, 2\}$. Notice that $\mathcal{L}^i$ is convex and $\lim_{|y| \rightarrow \infty} \mathcal{L}^i(y) = \infty$, so y_i^m is finite. We refer to the policy that orders up to the base-stock level y_i^m as the myopic policy for location i . We refer to the policy that orders up to y_1^m for location 1 and up to y_2^m for location 2 as the myopic policy.

THEOREM 4. *For stationary problems, the myopic policy is optimal for finite horizon problems.*

This result implies that for stationary problems information beyond the leadtimes do not affect the base-stock levels and the optimal costs. In other words, at time t knowing that there is some demand to come in any periods after $t + L_2$ will not change the order-up-to level for location 2. In turn, this implies that the implicit penalty costs are stationary and unaffected by the observation beyond location 2's leadtime. Therefore the problem corresponding to the first echelon, that is location 1's problem, is independent of the observations beyond its own leadtime.

Notice that we do not need to solve dynamic programs for each location anymore to obtain the optimal base-stock levels. A straightforward numerical integration would provide us the myopic cost functions $\mathcal{L}^i$ and their minimizers. We have an easy corollary.

COROLLARY 1. $IP_t(x, O_t^2) = IP^m(x)$ for all $t \leq T - L_1 - L_2$.

In words, under stationary cost parameters $IP_t(x, O_t^2)$ is independent of the advance demand information beyond the leadtime, and hence it is stationary and given by the myopic implicit penalty cost function $IP^m(x)$. We use this corollary to establish the optimality of myopic policies for the infinite horizon problem.

Next, we compare the myopic policy with the optimal policy for nonstationary demand and cost structure. In this case the myopic policy is defined as before; the only difference is that cost functions are time dependent; that is, y_{jt}^m is the smallest minimizer of $\mathcal{L}_t^j(y)$.

PROPOSITION 1. For nonstationary problems and any vector O_t^j and t , we have $y_{jt}(O_t^j) \leq y_{jt}^m$ for $j = 1, 2$.

6. Infinite Horizon Discounted Cost Model

For the infinite horizon problem, we continue to assume that all cost parameters and the demand are stationary. At the beginning of period t the initial state space is given by $(\hat{x}_{1t}, \bar{z}_{1t}, \tilde{x}_{2t}, \tilde{O}_t)$. Assume that we are given a policy $Y_\infty = ((y_{1t}, \tilde{y}_{2t}), (\tilde{y}_{1t+1}, \tilde{y}_{2t+1}), \dots)$. After implementing this policy at the beginning of each period $s \in \{t, t+1, \dots\}$, the *echelon net inventory* at location 1 is y_{1s} and the *echelon inventory position* at location 2 is $\tilde{y}_{2s}$. The expected discounted cost of managing this inventory under this policy for an infinite horizon is given by

$$\begin{aligned} J(\hat{x}_{1t}, \bar{z}_{1t}, \tilde{x}_{2t}, \tilde{O}_t | Y_\infty) \\ = \lim_{T \rightarrow \infty} E \sum_{s=t}^T \alpha \left\{ c_1(y_{1s} - \hat{x}_{1s}) + h_1 \hat{x}_{1,s+1} + c_2(\tilde{y}_{2s} - \tilde{x}_{2s}) \right. \\ \left. + G\left(y_{2s} - \sum_{r=t}^{t+L_2} O_{s,r}\right) \right\}. \end{aligned}$$

Our aim is to determine a policy that minimizes the above cost function. Notice that the limit exists and may be $+\infty$ because the one-period cost function is nonnegative. Let

$$J(\hat{x}_{1t}, \bar{z}_{1t}, \tilde{x}_{2t}, \tilde{O}_t) = \inf_{Y_\infty \in \Pi} J(\hat{x}_{1t}, \bar{z}_{1t}, \tilde{x}_{2t}, \tilde{O}_t | Y_\infty)$$

denote the cost of an optimal policy where Π is the class of infinite horizon measurable policies. If there exists an optimal policy, it achieves the infimum above for every state, and we refer to this policy as the one that solves the infinite horizon discounted cost problem. The optimality equation, also known as Bellman's equation, for this problem is given by

$$\begin{aligned} J(\hat{x}_{1t}, \bar{z}_{1t}, \tilde{x}_{2t}, \tilde{O}_t) \\ = \min_{(z_1, y) \in \mathcal{A}} \left\{ c_1 z_1 + c_2(y - \tilde{x}_{2t}) + h_1 E \hat{x}_{1,t+1} \right. \\ \left. + G\left(y - \sum_{s=t}^{t+L_2} O_{t,s}\right) \right. \\ \left. + \alpha E J(\hat{x}_{1,t+1}, \bar{z}_{1,t+1}, \tilde{x}_{2,t+1}, O_{t+1}) \right\}. \quad (13) \end{aligned}$$

Let $\pi^* \in \Pi$ be a stationary policy. Under positivity conditions, which are satisfied in our problem, π^* is an optimal policy if and only if it satisfies the right-hand side of (13); see Proposition 1.3 in Bertsekas (1995, p. 143). We show that this policy orders from the outside supplier according to the policy that solves the infinite horizon version of the location 1 problem and ships to the second location according to the policy that solves the infinite horizon version of the location 2 problem. To prove this we consider the finite horizon stationary case and the limit as the horizon grows to infinity. The finite horizon dynamic program for each location depends on time to go $T - t$. We consider the limit of the functions J_t and V_t^i as $T \rightarrow \infty$ and show that the decomposition for the finite horizon problems extends to the infinite horizon problems as well.

LEMMA 5. For every fixed vector $\tilde{O}_t$

- (1) $\lim_{T \rightarrow \infty} \tilde{V}_t^2(\tilde{x}_{2t}, \tilde{O}_t)$ exists and converges uniformly to a limit function $\tilde{V}^2(\tilde{x}_{2t}, \tilde{O}_t)$;
- (2) $\lim_{T \rightarrow \infty} \tilde{V}_t^1(\hat{x}_{1t}, \tilde{O}_t)$ exists and converges uniformly to a limit function $\tilde{V}^1(\hat{x}_{1t}, \tilde{O}_t)$; and

(3) $\lim_{T \rightarrow \infty} J_t(\hat{x}_{1t}, \tilde{z}_{1t}, \tilde{x}_{2t}, \tilde{O}_t)$ exists and converges uniformly to a limit function $J(\hat{x}_{1t}, \tilde{z}_{1t}, \tilde{x}_{2t}, \tilde{O}_t)$ that is equal to $\tilde{V}^1(\tilde{x}_{1t}, \tilde{O}_t) + \tilde{V}^2(\tilde{x}_{2t}, \tilde{O}_t) + \mathcal{F}_t$. The limit value is also the optimal value for the infinite horizon problem.

The two-location infinite horizon problem decomposes into two infinite horizon, single location problems. Next, we apply the state space reduction as suggested by Lemma 1 to obtain the limiting value of Equation (11). Let

$$V^2(x_{2t}, O_t^2) = -c_2 x_{2t} + \min_{y \geq x_{2t}} \{c_2 y + G(y) + \alpha E V^2(x_{t+1}, O_{t+1}^2)\}.$$

Notice that this is the optimality equation for the infinite horizon location 2 problem. Consequently, due to Theorem 1, an optimal policy for the above problem is given by y_2^m , the smallest minimizer of $\mathcal{L}^2(y) = (1 - \alpha)c_2 y + G(y)$. Similarly let

$$V^1(x_{1t}, O_t^1) = -c_1 x_{1t} + \min_{y \geq x_{1t}} \{c_1 y + C^m(y) + \alpha E V^1(x_{1,t+1}, O_{t+1}^1)\} \quad (14)$$

where $C^m(y) = \alpha^{L_1} E[h_1(y - \sum_{s=t}^{t+L_1} U_{t,s}) + \alpha IP^m(x - \sum_{s=t}^{t+L_1} U_{t,s})]$. Convexity of IP^m implies the convexity of $C^m(\cdot)$ and the $\lim_{|x| \rightarrow \infty} C^m(x) = \infty$. Hence, Theorem 1 suffices to show the optimality of base-stock policy with base-stock level y_1^m for Equation (14). Thus because of Theorem 1, Theorem 3, and Lemma 5, the myopic policy π^* , where the location 1 and 2 use the myopic base-stock level y_1^m and y_2^m , respectively satisfies the right-hand side of Equation (13); hence it is optimal. We summarize our result in the following Theorem.

THEOREM 5. *A base-stock policy is optimal for infinite horizon series system where base-stock levels are given by the myopic base-stock levels y_i^m , $i = 1, 2$.*

7. Numerical Study

For stationary costs and demand, the myopic policy is optimal. Hence we only need to search for the minimizer of the myopic cost functions $\mathcal{L}^i$ to obtain the optimal echelon base-stock levels. For the non-stationary problems, however, we need to first solve location 2's dynamic program, calculate the implicit penalty costs, and then solve location 1's dynamic program. To do this we use a backward induction

algorithm. The idea is to solve the dynamic program starting from the last period, which is a single period problem, by evaluating the cost for each instance of the state space, choosing an action that minimizes the cost, and repeating these steps until the first stage is reached.

```
FOR  $i = 2$  to 1 do
  Initialize  $V_{T+1}^i(x, O^i) \equiv c_{iT+1}x$  for all  $x$  and  $O^i$ 
  FOR  $t = T$  to 1 do
    Given  $O^i$  evaluate  $H_t^i(y, O^i)$  for all  $y$  and choose
    a minimizer  $y_t^i(O^i)$ 
    If  $i \neq 1$ , evaluate the implicit penalty cost
    function  $IP_t^i(y, O^i)$ 
    Evaluate  $V_t^i(x, O^i)$ 
  end FOR
end FOR
```

The expectations that appear in the dynamic programs are straight forward numerical integrations. The convolutions of random variables are easier to calculate when the demand process is chosen from a family of regenerative distributions.⁵ Recall that the dimension of the state space for location j is $1 + (N - L_j - 1)^+$ for $j = \{1, 2\}$. The computational burden is mainly because of the increase in the dimension of the state space. We limit our numerical study to problems where $L_1 = L_2$ and $N = L_1 + 2$. The state space for these problems is of dimension 2 for both locations. They are general enough to capture the main ideas.

We assume $L_1 = L_2 = 1$ and $N = 3$, and also that all cost parameters are stationary. The demand vector is given by $D_t = (D_{t,t}, D_{t,t+1}, D_{t,t+2}, D_{t,t+3})$. Observe that the vectors O_t^1 and O_t^2 are scalars and equal to $D_{t-1,t+2}$. For our computational study we model the arrival of customers by a Poisson process with mean Λ_t . An arriving customer requests her demand to be fulfilled $l \in 0, \dots, N$ periods later with probability p_{tl} where $\sum_{l=1}^N p_{tl} = 1$. Hence the demand $D_{t,t+l}$ is Poisson with mean $\lambda_{tl} = p_{tl}\Lambda_t$. One interpretation is that with probability p_{tl} the demand leadtime will be l periods. For stationary problems another interpretation is that $p_l * 100\%$ of the demand for any period s is placed during period $s - l$.

⁵ A class of distribution is said to be regenerative if the superposition of distributions from this class belongs to the same class. Poisson is a regenerative distribution.

Table 1 Optimal Base-Stock Levels and Resulting Cost for $h_1 = 1, h_2 = 3, c_1 = 10, c_2 = 30$

$p = 19, 20$ periods						$p = 19, 20$ periods					
No.	$(\lambda_0, \lambda_1, \lambda_2, \lambda_3)$	y_{1t}	y_{2t}	Cost	∇ %	No.	$(\lambda_0, \lambda_1, \lambda_2, \lambda_3)$	y_{1t}	y_{2t}	Cost	∇ %
1	(4, 0, 0, 0)	20	10	2,806		5	(3, 0, 0, 0)	15	8	2,142	
2	(0, 4, 0, 0)	11	6	2,405	14.3	6	(0, 3, 0, 0)	8	4	1,830	14.5
3	(0, 0, 4, 0)	0	0	1,936	31.0	7	(0, 0, 3, 0)	0	0	1,452	32.2
4	(0, 0, 0, 4)	0	0	1,780	36.4	8	(0, 0, 0, 3)	0	0	1,338	37.5
$p = 19, 40$ periods						$p = 9, 40$ periods					
No.	$(\lambda_0, \lambda_1, \lambda_2, \lambda_3)$	y_{1t}	y_{2t}	Cost	∇ %	No.	$(\lambda_0, \lambda_1, \lambda_2, \lambda_3)$	y_{1t}	y_{2t}	Cost	∇ %
9	(2, 0, 0, 0)	11	6	1,978		13	(3, 0, 0, 0)	13	7	2,672	
10	(0, 2, 0, 0)	6	3	1,698	14.2	14	(0, 3, 0, 0)	7	4	2,372	11.2
11	(0, 0, 2, 0)	0	0	1,322	33.2	15	(0, 0, 3, 0)	0	0	1,983	25.8
12	(0, 0, 0, 2)	0	0	1,246	37.0	16	(0, 0, 0, 3)	0	0	1,869	30.0
$p = 19, 20$ periods						$p = 99, 40$ periods					
No.	$(\lambda_0, \lambda_1, \lambda_2, \lambda_3)$	y_{1t}	y_{2t}	Cost	∇ %	No.	$(\lambda_0, \lambda_1, \lambda_2, \lambda_3)$	y_{1t}	y_{2t}	Cost	∇ %
17	(3, 0, 0, 0)	15	8	2,142		21	(3, 0, 0, 0)	19	10	3,801	
18	(2.4, 0.6, 0, 0)	14	7	2,081	2.8	22	(2.4, 0.6, 0, 0)	17	10	3,625	4.6
19	(2.4, 0, 0.6, 0)	14	7	2,033	5.1	23	(2.4, 0, 0.6, 0)	16	9	3,498	8.0
20	(2.4, 0, 0, 0.6)	12	7	1,964	8.3	24	(2.4, 0, 0, 0.6)	16	9	3,364	11.5
$p = 99, 20$ periods						$p = 99, 40$ periods					
No.	$(\lambda_0, \lambda_1, \lambda_2, \lambda_3)$	y_{1t}	y_{2t}	Cost	∇ %	No.	$(\lambda_0, \lambda_1, \lambda_2, \lambda_3)$	y_{1t}	y_{2t}	Cost	∇ %
25	(3, 0, 0, 0)	18	10	3,439		31	(1, 0, 0, 0)	8	5	1,436	
26	(2, 1, 0, 0)	15	8	3,122	9.2	32	(0.5, 0.5, 0, 0)	7	4	1,265	11.9
27	(1, 1, 1, 0)	10	6	2,559	25.6	33	(0, 0.5, 0.5, 0)	3	2	934	34.9
28	(0, 1, 1, 1)	4	3	1,899	44.8	34	(0, 0.5, 0, 0.5)	3	2	872	39.2
29	(0, 0, 1, 2)	0	0	1,367	60.3	35	(0, 0, 0.5, 0.5)	0	0	603	58.0
30	(0, 0, 0, 3)	0	0	1,338	61.1	36	(0, 0, 0, 1)	0	0	596	58.2

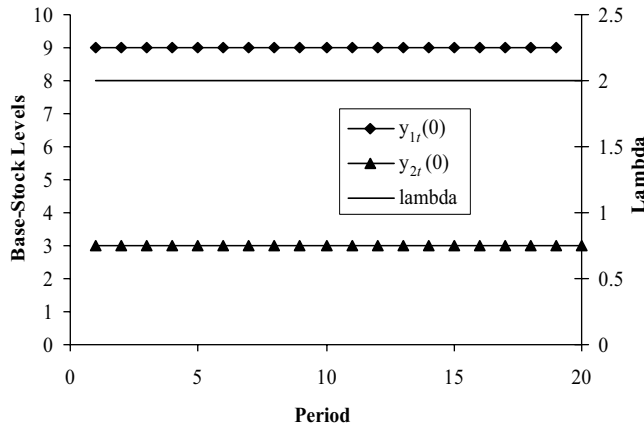
To portray the necessary insights, and also to illustrate our theoretical results numerically, we present a subset of problem instances with the parameters $h_1, h_2 = 1, 2, 3, 6, c_1, c_2 = 10, 20, 30, p = 9, 19, 99, \alpha = 0.95, T = 20, 40$. We start by analyzing the benefits of using advance demand information for stationary problems.

Note that under our numerical study the demand at any period s is given by $D_{s-3,s} + D_{s-2,s} + D_{s-1,s} + D_{s,s}$. In our first set of experiments (as summarized in Table 1), we address the stationary demand case; that is, $\Lambda_t = \Lambda$ and $p_{it} = p_i$. In this case, the mean value of demand for any period s is given by $\Lambda = \sum_{l=0}^3 \lambda_l$. We fix Λ and change the value of p_i to study different advance demand information scenarios. Consider the following two extreme scenarios. Under the first scenario assume that the inventory manager is incapable of obtaining advance demand information. We model this first case by setting $p_0 = 1$ and $p_1 = p_2 = p_3 \equiv 0$. On the other extreme, assume that the inventory manager implements an aggressive strategy to convince all customers to place orders three periods in advance. This

scenario is modeled by $p_0 = p_1 = p_2 \equiv 0, p_3 = 1$. All the other possibilities lie between these two extreme scenarios.

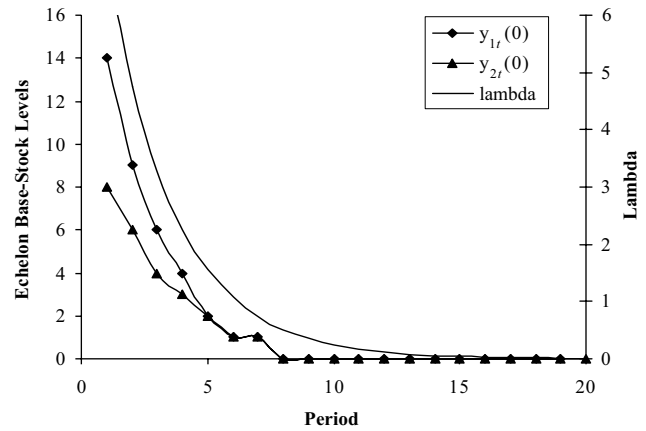
We divide Table 1 into eight groups reflecting different parameter set; Λ , penalty cost p and the horizon length T . For each such set we explore the impact of having more advance demand information. Notice the reduction in echelon base-stock levels as more of the demand is known in advance. In particular, the base-stock levels drop to zero when the manager obtains all the demand information two periods in advance. This illustrates how advance demand information allows a fundamental bridge between build-to-stock and build-to-order systems. Our model can be used to assess the cost and benefit of such a shift in series production systems. Even in the case of moderate advance demand information, modeled in Experiments 17 through 24, the cost reductions range from 4% to 11%. It is also worth noticing that having advance demand information is more desirable for a system with high penalty costs.

Figure 2 Stationary Demand



Note. $h_1 = 1$; $h_2 = 3$; $p = 19$; $c_1 = 10$; $c_2 = 30$.

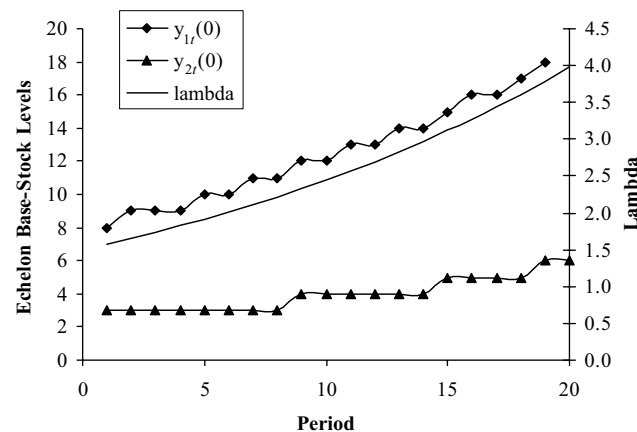
Figure 3(b) Ramp-Down Demand



Note. $h_1 = 1$; $h_2 = 1$; $p = 9$; $c_1 = 10$; $c_2 = 10$.

Next we turn our attention to the nonstationary demand and its impact on the optimal echelon base-stock levels. We assume for the next set of experiments that $p_{ti} = p_{ij}$ for all i, j . Figures 2 and 3 exhibit optimal echelon base-stock levels at each period. Our benchmark is Figure 2. It addresses the case for stationary demand process where the mean is $\Lambda_t = 2$ for all $t = \{1, \dots, 20\}$. Figure 3(a) depicts the result for a ramp-up demand case where the average number of customers arriving, Λ_t , is increasing over time. Figure 3(b) is for a ramp-down demand case. In Table 2, as a complement to these figures, we tabulate the optimal echelon base-stock levels at Period 1 as

Figure 3(a) Ramp-Up Demand



Note. $h_1 = 1$; $h_2 = 3$; $p = 19$; $c_1 = 10$; $c_2 = 30$.

a function of the observed demand beyond the location's leadtime and the myopic base-stock levels. We observe that myopic base-stock levels are optimal for stationary and ramp-up demand. This observation verifies both the optimality of echelon base-stock policies and the optimality of myopic policy for stationary problems. For these cases information beyond each location's leadtime does not affect the echelon base-stock levels. Hence, information beyond the leadtime has no operational value for both the stationary and the ramp-up demand process.

8. Conclusion

In this paper we establish the optimal control policy for a series system that incorporates advance demand information. We show the optimality of state-dependent, echelon base-stock policies. We also

Table 2 Is Myopic Policy Optimal for $t = 1$?

$D_{t-1, t+1}$	0	1	2	3	4	5	y_i^m
Stationary Demand							
$y_{1t}(D_{t-1, t+1})$	9	9	9	9	9	9	9
$y_{2t}(D_{t-1, t+1})$	3	3	3	3	3	3	3
Ramp-Up Demand							
$y_{1t}(D_{t-1, t+1})$	8	8	8	8	8	8	8
$y_{2t}(D_{t-1, t+1})$	3	3	3	3	3	3	3
Ramp-Down Demand							
$y_{1t}(D_{t-1, t+1})$	14	14	15	15	15	16	13
$y_{2t}(D_{t-1, t+1})$	8	8	8	8	8	8	8

prove that myopic policies are optimal for stationary problems. Hence the computational effort to solve this problem dramatically reduces for the stationary case. For problems where the information horizon is smaller than $\min\{L_1, \dots, L_J\} + 1$, the decomposed dynamic programs for each location reduce to a single dimensional problem. Hence, existing software for classical series systems can be used to solve our problem for this special case provided that the inventory positions are redefined as suggested in the paper. If the manager, however, wants to incorporate additional information beyond the $\min\{L_1, \dots, L_J\} + 1$, then he has to deal with state-dependent policies. We propose an algorithm that recursively solves the dynamic program for this more general case. We provide a numerical study to shed some light on the value of advance demand information for a series system. We illustrate how this information provides a bridge between build-to-stock and build-to-order production systems. Our results provide a tool to assess the cost and benefit of advance demand information and can be used to design an effective production system.

The scheme developed in this paper assumes that the control is centralized. The policy can be decentralized. The problem in this case is a pull system where each echelon orders based on their local base-stock levels. In this case, echelon inventory managers need to know the advance demand information seen only by the last stage. Another interesting future research direction is to analyze and develop incentive systems to share the benefits of advance demand information so as to induce the party down stream in the supply chain to share this information with the other parties located upstream in the chain.

Appendix A. Glossary of Notation

Notation for Advance Demand Information Model:

- N : length of information horizon;
- $D_{t,s}$: orders placed for period $s \in t, \dots, t+N$ during period t ;
- $O_{t,s}$: observed (known) part of period s demand at the beginning of period t ;
- $U_{t,s}$: unobserved (unknown) part of period s demand at the beginning of period t ;
- $\tilde{O}_t = (O_{t,t}, O_{t,t+1}, \dots, O_{t,t+N-1})$, vector of observed demands;
- $O_t = (O_{t,t+L+1}, \dots, O_{t,t+N-1})$, vector of observed demands beyond the leadtime;

$$\begin{aligned} O_t^1 &= (O_{t,t+L_1+1}, \dots, O_{t,t+N-1}); \\ O_t^2 &= (O_{t,t+L_2+1}, \dots, O_{t,t+N-1}). \end{aligned}$$

Notation for the Multilocation Problems:

- T : terminal period for the finite horizon problem;
- J : total number of locations in the series system and also denotes the last location;
- L_j : leadtime for a shipment dispatched from location $j-1$ to location j ;
- $L'_j = L_j - 1$;
- p_t : penalty cost (at location J) per unit;
- h'_{jt} : local inventory holding cost at location j per unit;
- h_{jt} : echelon holding cost at location j per unit;
 $= h'_{jt} - h'_{j-1,t}$ for $j \in \{2, \dots, J\}$ and $h_{1t} = h'_{1t}$;
- c_{jt} : shipment cost per unit in period t ;
- I_{jt} : inventory on hand at location j at the beginning of period t ;
- B_t : backorders at location J at period t ;
- $\hat{x}_{jt}$: echelon *net* inventory at location j ;
- $\tilde{x}_{jt} = \hat{x}_{jt} + \sum_{s=t}^{t+L'_j} z_{js}$, echelon inventory *position* at location j ;
- $x_{jt} = \tilde{x}_{jt} - \sum_{s=t}^{t+L_j} O_{t,s}$, *modified* echelon inventory position at location j ;
- $\tilde{z}_{jt} = (z_{j1}, \dots, z_{jL'_j})$ trans-shipments to location j ;
- z_{js} : shipment dispatched from location $j-1$ to location j at the beginning of period $t-s$;
- z_1 : orders from outside supplier to be received at location 1;
- z_j : shipments to location $j > 2$ from location $j-1$.

Optimal Dynamic Programming Cost Functions:

Our convention is to drop the tilde signs to emphasize that the observed part of the leadtime demand has been subsumed and we operate with a reduced state space. At period t

- $\hat{J}_t(\hat{x}_{1t}, \tilde{z}_{1t}, \hat{x}_{2t}, \tilde{z}_{2t}, \tilde{O}_t)$: before decomposition and before subsuming the pipeline inventory;
- $J_t(\hat{x}_{1t}, \tilde{z}_{1t}, \hat{x}_{2t}, \tilde{O}_t)$: before decomposition after subsuming $\tilde{z}_{2t}$;
- $\tilde{V}_t^2(\hat{x}_{2t}, \tilde{O}_t)$: DP for location 2 after the decomposition and before subsuming the observed part of the leadtime demand;
- $V_t^2(x_{2t}, O_t^2)$: after subsuming the observed part of the leadtime demand;
- $\tilde{V}_t^1(\hat{x}_{1t}, \tilde{z}_{1t}, \tilde{O}_t)$: for location 1 after decomposition before subsuming $\tilde{z}_{1t}$;
- $\tilde{V}_t^1(\hat{x}_{1t}, \tilde{O}_t)$: after subsuming the pipe-line inventory;
- $V_t^1(x_{1t}, O_t^1)$: after subsuming the observed part of the leadtime demand.

Appendix B. Extension to $J > 2$ Case

The optimal cost-to-go function for this more general case, similar to Equation (7), is given by

$$\begin{aligned} &\hat{J}_t(\hat{x}_{1t}, \tilde{z}_{1t}, \dots, \hat{x}_{Jt}, \tilde{z}_{Jt}, \tilde{O}_t) \\ &= \min_{(z_1, \dots, z_J) \in \mathcal{A}^J} \left\{ \sum_{i=1}^J c_{it} z_i + \sum_{i=1}^{J-1} h_{it} E \hat{x}_{i,t+1} + E g_t(\hat{x}_{J,t+1}) \right. \\ &\quad \left. + \alpha E \hat{J}_{t+1}(\hat{x}_{1,t+1}, \tilde{z}_{1,t+1}, \dots, \hat{x}_{J,t+1}, \tilde{z}_{J,t+1}, O_{t+1}) \right\} \quad (15) \end{aligned}$$

where $\hat{J}_{T'+\sum_{i=1}^J L_{i+1}}(\hat{x}_1, \dots, \hat{x}_J, \dots, \cdot) \equiv -\sum_{i=1}^J c_i \hat{x}_i$, $g_t(x) = h_{jt}[x] + (p_{jt} + h'_{jt})[x]$, and $\mathcal{A}' = \{(z_1, \dots, z_J): z_i \geq 0, \forall i \text{ and } \tilde{x}_{it} + z_i \leq \hat{x}_{i-1,t}, \forall i > 1\}$. Notice that shipment constraints $\tilde{x}_{it} + z_i \leq \hat{x}_{i-1,t}$ are equivalent to $z_i \leq L_{i-1,t}$, meaning that shipments to location i are bounded by the inventory on hand at location $i-1$. Notice also that g_t is convex and $\lim_{|x| \rightarrow \infty} g_t(x) = \infty$.

We first apply Lemma 2 to reduce the state space and embed $\tilde{z}_{jt}$ into the net inventory of the last location. Hence, the cost-to-go function after this transformation is given by $J_t(\hat{x}_{1t}, \tilde{z}_{1t}, \dots, \tilde{x}_{Jt}, \tilde{O}_t)$. Next, using the arguments in Theorem 2, we decompose the cost function.

$$J_t(\hat{x}_{1t}, \tilde{z}_{1t}, \dots, \tilde{x}_{Jt}, \tilde{O}_t) = \hat{J}_t(\hat{x}_{1t}, \tilde{z}_{1t}, \dots, \hat{x}_{j-1,t}, \tilde{z}_{j-1,t}, \tilde{O}_t) + \tilde{V}_t^j(\tilde{x}_{jt}, \tilde{O}_t).$$

The first cost function is the cost-to-go function for a $J-1$ location series system, and the second cost function is the single location problem for location J . The location J problem is similar to Equation (8). The implicit penalty cost function to be used in $\hat{J}_t$ is based on the cost of not being able to satisfy the requirements of location J . Hence,

$$\begin{aligned} & \hat{J}_t(\hat{x}_{1t}, \tilde{z}_{1t}, \dots, \hat{x}_{j-1,t}, \tilde{z}_{j-1,t}, \tilde{O}_t) \\ &= \min_{(z_1, \dots, z_{j-1}) \in \mathcal{A}'} \left\{ \sum_{i=1}^{j-1} c_{it} z_i + \sum_{i=1}^{j-2} h_{it} E \hat{x}_{i,t+1} + \tilde{P}_t(\hat{x}_{j-1,t+1}, \tilde{O}_t) \right. \\ & \quad \left. + \alpha E \hat{J}_{t+1}(\hat{x}_{1,t+1}, \tilde{z}_{1,t+1}, \dots, \hat{x}_{j-1,t+1}, \tilde{z}_{j-1,t+1}, \tilde{O}_{t+1}) \right\}. \end{aligned} \quad (16)$$

From this point on we use Theorem 2 recursively to decompose the entire system into single location problems. The arguments in the proof of Theorem 2 enable us to embed the pipeline inventory $\tilde{z}_{jt}$ into the net inventory of that location and to decompose the j location series system into a single location problem and a series system with $j-1$ locations. Once the J location series system is decomposed into single location problems, we apply Lemma 1 to further collapse the state space of these single location problems. Theorem 1 suffices to conclude the optimality of state-dependent base-stock policy for each problem. In summary,

Apply Lemma 2 to Equation (15): $\hat{J}_t(\hat{x}_{1t}, \tilde{z}_{1t}, \dots, \hat{x}_{Jt}, \tilde{O}_t) \rightarrow J_t(\hat{x}_{1t}, \tilde{z}_{1t}, \dots, \tilde{x}_{Jt}, \tilde{O}_t);$

FOR $i = J$ to 1 do

Apply Theorem 2 and decompose $J_t(\hat{x}_{1t}, \tilde{z}_{1t}, \dots, \tilde{x}_{it}, \tilde{O}_t)$ into two subproblems:

$$\hat{J}_t(\hat{x}_{1t}, \tilde{z}_{1t}, \dots, \hat{x}_{i-1,t}, \tilde{z}_{i-1,t}, \tilde{O}_t) + \tilde{V}_t^i(\tilde{x}_{it}, \tilde{O}_t)$$

Next subsume $\tilde{z}_{i-1,t}$:

$$\hat{J}_t(\hat{x}_{1t}, \tilde{z}_{1t}, \dots, \hat{x}_{i-1,t}, \tilde{z}_{i-1,t}, \tilde{O}_t) \rightarrow J_t(\hat{x}_{1t}, \tilde{z}_{1t}, \dots, \tilde{x}_{i-1,t}, \tilde{O}_t)$$

end FOR;

Apply Lemma 1 to all single location problems: $\tilde{V}_t^i(\tilde{x}_{it}, \tilde{O}_t) \rightarrow V_t^i(x_{it}, O_t^i).$

We remark that the algorithm described in the numerical section can be used to solve systems with several locations by changing the first loop to “FOR $i = J$ to 1 do.”

Appendix C. The Proofs

PROOF OF THEOREM 2. We first show that

$$J_t(\hat{x}_{1t}, \tilde{z}_{1t}, \tilde{x}_{2t}, \tilde{O}_t) = \hat{V}_t^1(\hat{x}_{1t}, \tilde{z}_{1t}, \tilde{O}_t) + \tilde{V}_t^2(\tilde{x}_{2t}, \tilde{O}_t). \quad (17)$$

Second, we subsume the location 1 pipeline inventory into its echelon net inventory and conclude our proof by showing that $\hat{V}_t^1(\hat{x}_{1t}, \tilde{z}_{1t}, \tilde{O}_t) = \hat{V}_t^1(\hat{x}_{1t}, \tilde{O}_t) + \mathcal{F}_t$. Let us first define the DP:

$$\begin{aligned} \hat{V}_t^1(\hat{x}_{1t}, \tilde{z}_{1t}, \tilde{O}_t) &= -c_{1t} \hat{x}_{1t} + \min_{\hat{x}_{1t} \leq y} \{c_{1t} y + h_{1t} E \hat{x}_{1,t+1} + \tilde{P}_t(\hat{x}_{1t}, \tilde{O}_t) \\ & \quad + \alpha E \hat{V}_{t+1}^1(\hat{x}_{1,t+1}, \tilde{z}_{1,t+1}, \tilde{O}_{t+1})\}, \end{aligned}$$

where $\hat{V}_{T+1}^1(\hat{x}_1, \cdot, \cdot) \equiv -c_{1,T+1} \hat{x}_1$. Notice that our first claim Equation (17) is true by definition for $t = T+1$. Now suppose it holds for $t+1$. Then

$$\begin{aligned} & J_t(\hat{x}_{1t}, \tilde{z}_{1t}, \tilde{x}_{2t}, \tilde{O}_t) \\ &= \min_{(z_1, y) \in \mathcal{A}} \left\{ c_{1t} z_1 + h_{1t} E \hat{x}_{1,t+1} + \alpha E \hat{V}_{t+1}^1(\hat{x}_{1,t+1}, \tilde{z}_{1,t+1}, \tilde{O}_{t+1}) \right. \\ & \quad \left. + c_{2t}(y - \tilde{x}_{2t}) + G_t \left(y - \sum_{s=t}^{t+L_2} O_{t,s} \right) + \alpha E \tilde{V}_{t+1}^2(\tilde{x}_{2,t+1}, \tilde{O}_{t+1}) \right\}. \end{aligned} \quad (18)$$

There are two decision variables to consider, z_1 and y . The ordering decision z_1 has an impact on the echelon net inventory at location 1. These decision variables are linked only by the constraint $\tilde{x}_{2t} \leq y \leq \hat{x}_{1t}$. We first fix z_1 and optimize over y . This results in

$$-c_{2t} \tilde{x}_{2t} + \min_{\tilde{x}_{2t} \leq y \leq \hat{x}_{1t}} \tilde{H}_t^2(y, \tilde{O}_t).$$

This is a capacitated version of Problem (8) where the decision variable y is not restricted to be less than $\hat{x}_{1,t}$. Notice that $\tilde{y}_{2t}(\tilde{O}_t)$ exists because $G_t(\cdot)$ satisfies Assumption 2. If $\hat{x}_{1t} \geq \tilde{y}_{2t}(\tilde{O}_t)$, then the upper bound induces no penalty. On the other hand, if $\hat{x}_{1t} < \tilde{y}_{2t}(\tilde{O}_t)$, then it is optimal to set $y = \hat{x}_{1t}$ and the implicit penalty cost is $\tilde{H}_t^2(\hat{x}_{1t}, \tilde{O}_t) - \tilde{H}_t^2(\tilde{y}_{2t}(\tilde{O}_t), \tilde{O}_t)$. Consequently,

$$\begin{aligned} -c_{2t} \tilde{x}_{2t} + \min_{\tilde{x}_{2t} \leq y \leq \hat{x}_{1t}} \tilde{H}_t^2(y, \tilde{O}_t) &= \tilde{V}_t^2(\tilde{x}_{2t}, \tilde{O}_t) + \tilde{H}_t^2(\min\{\tilde{y}_{2t}(\tilde{O}_t), \hat{x}_{1t}\}, \tilde{O}_t) \\ & \quad - \tilde{H}_t^2(\tilde{y}_{2t}(\tilde{O}_t), \tilde{O}_t) \\ &= \tilde{V}_t^2(\tilde{x}_{2t}, \tilde{O}_t) + \tilde{P}_t(\hat{x}_{1t}, \tilde{O}_t). \end{aligned}$$

Substituting into Equation (18) we obtain

$$\begin{aligned} J_t(\hat{x}_{1t}, \tilde{z}_{1t}, \tilde{x}_{2t}, \tilde{O}_t) &= \min_{z_1 \geq 0} \{c_{1t} z_1 + h_{1t} E \hat{x}_{1,t+1} + \alpha E \hat{V}_{t+1}^1(\hat{x}_{1,t+1}, \tilde{z}_{1,t+1}, \tilde{O}_{t+1}) \\ & \quad + \tilde{V}_t^2(\tilde{x}_{2t}, \tilde{O}_t) + \tilde{P}_t(\hat{x}_{1t}, \tilde{O}_t)\} \\ &= \min_{z_1 \geq 0} \{c_{1t} z_1 + h_{1t} E \hat{x}_{1,t+1} + \tilde{P}_t(\hat{x}_{1t}, \tilde{O}_t) \\ & \quad + \alpha E \hat{V}_{t+1}^1(\hat{x}_{1,t+1}, \tilde{z}_{1,t+1}, \tilde{O}_{t+1})\} + \tilde{V}_t^2(\tilde{x}_{2t}, \tilde{O}_t) \\ &= \hat{V}_t^1(\hat{x}_{1t}, \tilde{z}_{1t}, \tilde{O}_t) + \tilde{V}_t^2(\tilde{x}_{2t}, \tilde{O}_t), \end{aligned}$$

completing the induction argument and proving our first claim Equation (17). Now let us prove our second claim. Notice that if we initiate an order for z_1 units at the beginning of period t , it arrives at location 1 at the beginning of period $t + L_1$. Consequently, the echelon net inventory at location 1 at the end of period $t + L_1$ is given by

$$\hat{x}_{1t} + \sum_{l=1}^{L_1'} z_{1l} + z_1 - \left(\sum_{s=t}^{t+L_1} O_{t,s} + \sum_{s=t}^{t+L_1} U_{t,s} \right) - \tilde{x}_{1t} + z_1 - \sum_{s=t}^{t+L_1} O_{t,s} - \sum_{s=t}^{t+L_1} U_{t,s},$$

where $\sum_{s=t}^{t+L_1} O_{t,s} + \sum_{s=t}^{t+L_1} U_{t,s}$ is the leadtime demand for location 1, that is, the demand during periods $\{t, \dots, t + L_1\}$. The echelon net inventory at the end of period $t + L_1$ impacts the holding cost charged to period $t + L_1$ and the implicit penalty cost charged to period $t + L_1 + 1$. Thus, the location 1 problem reduces to $\tilde{V}_t^1$, defined as in Equation (9). The difference between $\hat{V}_t^1$ and $\tilde{V}_t^1$ consists of the unaccounted holding costs over periods $\{t, \dots, t + L_1\}$ and unaccounted implicit penalty costs over periods $\{t, \dots, t + L_1\}$:

$$\mathcal{F}_t(\hat{x}_{1t}, \tilde{x}_{1t}, \tilde{O}_t) = \sum_{s=t}^{t+L_1'} \alpha^{s-t} h_{1s} E[\hat{x}_{1,s+1}] + \sum_{s=t}^{t+L_1} \alpha^{s-t} E[\tilde{IP}_s(\hat{x}_{1s}, \tilde{O}_s)].$$

Notice that the $\hat{x}_{1s}$ for $s \in \{t, \dots, t + L_1\}$ are independent from the decision z_1 and z_2 made at time t . Hence for optimization purposes we can eliminate this constant from further consideration. $\square$

PROOF OF THEOREM 4. Before we establish the optimality of the myopic policy we should understand what happens at time $T - L_2 + 1$. At the beginning of period $T - L_2 + 1$, the inventory on hand at location 1 is never shipped to location 2. Hence, it can and should be salvaged at this time and the corresponding cost is charged to period $T - L_1 - L_2 + 1$. Notice that this avoids further location 1 holding costs. However, we have formulated the problem with salvage at time $T + 1$. This was done for convenience, and is equivalent to our intended formulation provided that: (i) we do not charge holding costs to location 1 after period $T - L_2$, and (ii) we adjust the salvage value accordingly; that is, $c_{1,T+1} = c_1/\alpha^{L_2}$. See also the discussion after Assumption 4.

According to Theorem 3a state-dependent base-stock policy is optimal. Recall that we do not release shipments to location 2 after period $T - L_2$. Our first task is to show that an optimal policy for location 2 is to order up to y_2^m for all $t \leq T - L_2$. To do this, it is important to write the cost function for period $T - L_2 + 1$. Because holding and penalty costs are shifted by L_2 units and because holding and penalty costs are zero for $t > T$, the only nonzero costs are those related to the terminal condition at time $T + 1$. Consequently from Equation (10)

$$\begin{aligned} V_{T-L_2+1}^2(x, O_{T-L_2+1}^2) &= -c_2 \left(x - E \sum_{s=T-L_2+1}^{T+1} U_{T-L_2+1,s} \right), \\ H_{T-L_2}^2(y, O_{T-L_2}^2) &= c_2 y + G(y) + \alpha E V_{T-L_2+1}^2(y - W, O_{T-L_2+1}^2) \\ &= \mathcal{L}^2(y) + \alpha c_2 E \left\{ \sum_{s=T-L_2+1}^{T+1} U_{T-L_2+1,s} + W \right\}, \end{aligned}$$

where $W = \sum_{s=T-L_2}^{T+1} D_{T-L_2,s} + O_{T-L_2,T+1}$ is a nonnegative random variable and the last term is a constant. Because the optimal base-stock level for period $T - L_2$ is obtained by minimizing $H_{T-L_2}^2(y, O_{T-L_2}^2)$, it follows that $y_{T-L_2}(O_{T-L_2}^2) = y_2^m$. Assume for an induction argument that the myopic policy is optimal for $t + 1$. Then we have for all y

$$\begin{aligned} V_{t+1}^2(y, O_{t+1}^2) &= -c_2 y + H_{t+1}^2(\max(y, y_2^m), O_{t+1}^2), \\ H_t^2(y, O_t^2) &= c_2 y + G(y) + \alpha E V_{t+1}^2(y - X, O_{t+1}^2) \\ &= \mathcal{L}^2(y) + \alpha E \{ c_2 X + H_{t+1}^2(\max(y - X, y_2^m), O_{t+1}^2) \}, \end{aligned}$$

where $X = \sum_{s=t}^{t+L_2+1} D_{t,s} + O_{t,t+L_2+1}$ is a nonnegative random variable. For any $y \leq y_2^m$, the term in brackets is independent of y , hence the minimizer of $\mathcal{L}^2(\cdot)$ is the minimizer of $H_t^2(\cdot)$. Notice that for $y > y_2^m$ both $\mathcal{L}^2(\cdot)$ and $H_{t+1}^2(\cdot, O_{t+1}^2)$ are nondecreasing and $E H_{t+1}^2(\max(y - X, y_2^m), O_{t+1}^2)$ is a convex combination of nondecreasing functions and constants, hence $H_t^2(\cdot, O_{t+1}^2)$ is nondecreasing in this region. The optimal base-stock level for period t is obtained by minimizing $H_t^2(\cdot, O_t^2)$, so it follows that $y_t(O_t^2) = y_2^m$, concluding the induction argument for the first stage.

Notice also that from the definition of the implicit penalty cost we have $\tilde{IP}_t(x, O_t^2) = H_t^2(\min\{y_2^m, x\}, O_t^2) - H_t^2(y_2^m, O_t^2) = \mathcal{L}^2(\min\{y_2^m, x\}) - \mathcal{L}^2(y_2^m) = IP^m(x)$ for $x \leq y_2^m$, otherwise it is equal to zero. From the definition of $C_t(y, O_t^1)$ given in Equation (11) we can also conclude that $C_t(y, O_t^1) = C^m(y)$. Next we establish the optimality of a myopic policy for location 1. We need to show that $y_{1t}(O_t^1) = y_1^m$ for all $t \leq T - L_1 - L_2$ because no orders from the outside supplier are allowed after time $T - L_1 - L_2$. Let us prove the optimality of the myopic policy first for $t = T - L_1 - L_2$. To do so it is important to write the cost function for period $T - L_1 - L_2 + 1$. Due to the salvage assumption and the fact that after time $T - L_2$ location 1 has no on-hand inventory,

$$\begin{aligned} V_{T-L_1-L_2+1}^1(x, O_{T-L_1-L_2+1}^1) &= -c_1 \left(x - E \sum_{s=T-L_1-L_2+1}^{T-L_2+1} U_{T-L_1-L_2+1,s} \right), \\ H_{T-L_1-L_2}^1(y, O_{T-L_1-L_2}^1) &= c_1 y + C^m(y) \\ &\quad + \alpha E V_{T-L_1-L_2+1}^1(y - Z, O_{T-L_1-L_2+1}^1) \\ &= \mathcal{L}^1(y) + \alpha c_1 E \left\{ \sum_{s=T-L_1-L_2+1}^{T-L_2+1} U_{T-L_1-L_2+1,s} + Z \right\}, \end{aligned}$$

where $Z = \sum_{s=T-L_1-L_2}^{T-L_2+1} D_{T-L_1-L_2,s} + O_{T-L_1-L_2,T-L_2+1}$ is a nonnegative random variable and the last term is a constant. Thus, $y_{T-L_1-L_2}^1(O_{T-L_1-L_2}^1) = y_1^m$ is optimal. This sets the stage for a similar induction argument (as described in location 2's DP) that establishes the optimality of the myopic policy for the first location. Combining our results, we have that the myopic policy is optimal for finite horizon stationary series systems. $\square$

PROOF OF LEMMA 1. It is based on induction. Notice that $\tilde{x}_{T+1} = x_{T+1}$ because by definition $O_{T+1} \equiv 0$, or in other words the manager does not allow customers to place orders beyond the terminal

period (i.e., the store is closed). Hence part (i) is true for $T + 1$. Assume it is also true for $t + 1$. Then we can write

$$\begin{aligned}\tilde{V}_t(\tilde{x}_t, \tilde{O}_t) &= \min_{\tilde{y}_t \geq \tilde{x}_t} \left\{ c_t(\tilde{y}_t - \tilde{x}_t) + G_t \left(\tilde{y}_t - \sum_{s=t}^{t+L} O_{t,s} \right) \right. \\ &\quad \left. + \alpha E \tilde{V}_{t+1}(\tilde{y}_t - O_{t,t} - D_{t,t}, \tilde{O}_{t+1}) \right\} \\ &= \min_{\tilde{y}_t \geq \tilde{x}_t} \left\{ c_t(\tilde{y}_t - \tilde{x}_t) + G_t \left(\tilde{y}_t - \sum_{s=t}^{t+L} O_{t,s} \right) \right. \\ &\quad \left. + \alpha EV_{t+1} \left(\tilde{y}_t - O_{t,t} - D_{t,t} - \sum_{s=t+1}^{t+L+1} O_{t+1,s}, O_{t+1} \right) \right\} \\ &= \min_{y_t \geq x_t} \left\{ c_t(y_t - x_t) + G_t(y_t) \right. \\ &\quad \left. + \alpha EV_{t+1} \left(y_t - \sum_{s=t}^{t+L+1} D_{t,s} - O_{t,t+L+1}, O_{t+1} \right) \right\} \\ &= V_t(x_t, O_t).\end{aligned}$$

The third inequality is obtained by substituting the update $O_{t+1,s}$ for $O_{t,s} + D_{t,s}$, x_t for $\tilde{x}_t - \sum_{s=t}^{t+L} O_{t,s}$ and y_t for $\tilde{y}_t - \sum_{s=t}^{t+L} O_{t,s}$ and rearranging the terms. This verifies the induction hypothesis and the proof of the first part. Next we show (i) $\Rightarrow$ (ii).

$$\begin{aligned}\tilde{H}_t(\tilde{y}_t, \tilde{O}_t) &= c_t \tilde{y}_t + G_t \left(\tilde{y}_t - \sum_{s=t}^{t+L} O_{t,s} \right) + \alpha E \tilde{V}_{t+1}(\tilde{x}_{t+1}, \tilde{O}_{t+1}) \\ &= c_t \tilde{y}_t + G_t \left(\tilde{y}_t - \sum_{s=t}^{t+L} O_{t,s} \right) + \alpha EV_{t+1}(x_{t+1}, O_{t+1}) \\ &= c_t \left(y_t + \sum_{s=t}^{t+L} O_{t,s} \right) + G_t(y_t) + \alpha EV_{t+1}(x_{t+1}, O_{t+1}) \\ &= H_t(y_t, O_t) + c_t \sum_{s=t}^{t+L} O_{t,s}.\end{aligned}$$

Next we show that (ii) $\Rightarrow$ (iii). Recall that $\tilde{y}_t(\tilde{O}_t)$ is the smallest minimizer of $\tilde{H}_t(\cdot, \tilde{O}_t)$. From (ii), and $y_t \equiv \tilde{y}_t - \sum_{s=t}^{t+L} O_{t,s}$, it follows that $\tilde{y}_t(\tilde{O}_t) - \sum_{s=t}^{t+L} O_{t,s}$ is the smallest minimizer of $H_t(\cdot, O_t)$. $\square$

PROOF OF LEMMA 2. In the statement of the Lemma we referred to $\mathcal{G}_t(\hat{x}_{2t}, \hat{z}_{2t}, \tilde{O}_t)$ as $\mathcal{G}_t$. The difference between $\hat{J}_t$ and J_t consists of the unaccounted costs over periods $\{t, \dots, t + L'_2\}$. That is, $\mathcal{G}_t(\hat{x}_{2t}, \hat{z}_{2t}, \tilde{O}_t) = \sum_{s=t}^{t+L'_2} \alpha^{s-t} E g_s(\hat{x}_{2t, s+1})$. We want to verify that these costs do not depend on the decisions made at time t . To see this, let $G_t^i(y) = \alpha^i E g_{t+i}(y - \sum_{s=t}^{t+i} U_{t,s})$. Recall that we do not incur any cost after the terminal period $T + 1$. Now we can write $\mathcal{G}_t(\hat{x}_{2t}, \hat{z}_{2t}, \tilde{O}_t) = G_t^0(\hat{x}_{2t} - O_{t,t}) + G_t^1(\hat{x}_{2t} + z_{2t, t'_2} - \sum_{s=t}^{t+1} O_{t,s}) + G_t^2(\hat{x}_{2t} + z_{2t, t'_2} + z_{2t, t'_2-1} - \sum_{s=t}^{t+2} O_{t,s}) + \dots + G_t^{L'_2}(\hat{x}_{2t} - \sum_{s=t}^{t+L'_2} O_{t,s})$, which is independent of the decisions made at time t , that is z_1 and z_2 . $\square$

PROOF OF LEMMA 3. $\tilde{IP}_t(\tilde{x}, \tilde{O}_t) = \tilde{H}_t^2(\min\{\tilde{y}_{2t}(\tilde{O}_t), \tilde{x}\}, \tilde{O}_t) - \tilde{H}_t^2(\tilde{y}_{2t}(\tilde{O}_t), \tilde{O}_t) = H_t^2(\min\{\tilde{y}_{2t}(\tilde{O}_t), \tilde{x}\} - \sum_{s=t}^{t+L} O_{t,s}, O_t) - H_t^2(\tilde{y}_{2t}(\tilde{O}_t) - \sum_{s=t}^{t+L} O_{t,s}, O_t) = H_t^2(\min\{y_{2t}(O_t), x\}, O_t) - H_t^2(y_{2t}(O_t), O_t) = IP_t(x, O_t)$. The second and third equalities follow from Lemma 1. $\square$

PROOF OF LEMMA 4. Notice that the implicit penalty cost inherits its convexity from $\tilde{H}_t^2$. In fact, $\tilde{IP}_t(\cdot, \tilde{O}_t)$ is nonincreasing convex. Thus, $\tilde{C}_t(y, \tilde{O}_t) = h_{1t}y + \alpha E \tilde{IP}_{t+1}(y, \tilde{O}_{t+1})$ is convex and $\lim_{|x| \rightarrow \infty} \tilde{C}_t(x, \tilde{O}_t) = \infty$ for all t . We can express $C_t(y, \tilde{O}_t) = \alpha^{L_1} E \tilde{C}_{t+L_1}(y - \sum_{s=t}^{t+L_1} U_{t,s}, \tilde{O}_{t+L_1})$. So, C_t inherits these properties from $\tilde{C}_t$. $\square$

PROOF OF PROPOSITION 1. We define $\nabla f(x, y) = f(x + 1, y) - f(x, y)$ as the first difference of function f . We first add and subtract $\alpha E c_{2, t+1} x_{2, t+1}$ to the right-hand side of Equation (10) and use the state updates to obtain $H_t^2(y, O_t^2) = (c_{2t} - \alpha c_{2, t+1})y + G_t(y) + \alpha E[V_{t+1}^2(x_{2, t+1}, O_{t+1}^2) - c_{2, t+1} x_{2, t+1}] + \alpha c_{2, t+1} (E \sum_{s=t}^{t+L_2+1} D_{t,s} + O_{t, t+L_2+1})$. From this we have $\nabla H_t^2(y, O_t^2) = \nabla \mathcal{L}_t^2(y) + \alpha \nabla[V_{t+1}^2(x_{2, t+1}, O_{t+1}^2) + c_{2, t+1} x_{2, t+1}]$. Notice that the second term is nonnegative due to Theorem 1, Part 1 (iii). Hence $\nabla H_t^2(y, O_t^2) \geq \nabla \mathcal{L}_t^2(y)$. This implies $y_{2t}(O_t^2) \leq y_{2t}^m$ because both $H_t^2(\cdot, O_t^2)$ and $\mathcal{L}_t^2(\cdot)$ are convex. To show that $y_{1t}(O_t^1) \leq y_{1t}^m$, it is enough to show that $\nabla H_t^1(y, O_t^1) \geq \nabla \mathcal{L}_t^1(y)$. Now $\nabla H_t^1(y, O_t^1) = (c_{1t} - \alpha c_{1, t+1}) + \nabla C_t(y, O_t^1) + \alpha E[V_{t+1}^1(x_{1, t+1}, O_{t+1}^1) + c_{1, t+1} x_{1, t+1}]$. The last term is positive because of Theorem 1, Part 1 (iii). On the other hand, from Equation (12) $\nabla \mathcal{L}_t^1(y) = (c_{1t} - \alpha c_{1, t+1}) + \alpha^{L_1} h_{1, t+L_1} + \alpha^{L_1+1} E \nabla IP^m(y - \sum_{s=t}^{t+L_1} U_{t,s})$. To complete the proof it is enough to show that $\nabla C_t(y, O_t^1) \geq \alpha^{L_1} h_{1, t+L_1} + \alpha^{L_1+1} E \nabla IP^m(y - \sum_{s=t}^{t+L_1} U_{t,s})$. However, Equation (11) we have $\nabla C_t(y, O_t^1) = \alpha^{L_1} h_{1, t+L_1} + \alpha^{L_1+1} E \nabla IP_{t+L_1+1}(y - \sum_{s=t}^{t+L_1} U_{t,s}, O_t^1)$. To conclude we need to show $\nabla IP_t(y, O_t) \geq \nabla IP^m(y)$. However, notice that $\nabla IP_t(y, O_t) = \nabla H_t^2(\min\{y, y_{2t}(O_t^2)\}, O_t^2) \geq \nabla \mathcal{L}_t^2(\min\{y, y_{2t}(O_t^2)\}) = \nabla IP^m(y)$ from Equation (12) and the inequality is from the first part of the proof. $\square$

PROOF OF LEMMA 5. Let $\tilde{V}^1(\tilde{x}_{1t}, \tilde{O}_t) = -c_{1t} \tilde{x}_{1t} + \min_{y \geq \tilde{x}_{1t}} \{c_{1t}y + C^m(\tilde{y} - \sum_{s=t}^{t+L_1} O_{t,s}) + \alpha E \tilde{V}^1(\tilde{x}_{1, t+1}, O_{t+1})\}$. Notice that C^m does not involve optimal cost function due to Corollary 1, and hence it is stationary. This result is important to establish the proof. Part 1 follows immediately from Theorem 1, Part 2, the proof of which is similar to that of Iglehart (1963). To prove the second part, recall first that for stationary problems the implicit penalty cost function is given by $IP^m(\cdot)$. Hence the dynamic program for $\tilde{V}_t^1$ has stationary single period costs that satisfy Assumption 2 and hence has similar structure to that of $\tilde{V}_t^2$. Now one can use similar arguments as in Theorem 1 to show that $\tilde{V}_t^1$ converges to $\tilde{V}^1(\tilde{x}_{1t}, \tilde{O}_t)$, concluding the proof of the second part. From Theorem 2 we know that the $J_t = V_t^1 + V_t^2 + \mathcal{F}_t$. Taking the limit of both sides and using Part 1 and 2 implies the first part of Lemma 5, Part 3. So far we have shown that the limit of the finite horizon problem converges and decomposes into two limiting functions. Next we want to show that this limit function, that is J , is actually the optimal value for the infinite horizon problem. Because the positivity assumption and the compactness of the set $U_t((\hat{x}_{1t}, \hat{z}_{1t}, \hat{x}_{2t}, \tilde{O}_t), \lambda)$ in Proposition 1.7, Bertsekas (1995, p. 148) is satisfied, we can conclude that J is the optimal value, concluding the proof of the lemma. $\square$

References

Azuory, K. S. 1985. Bayes solution to dynamic inventory models under unknown demand distribution. *Management Sci.* 31 1150–1160.

- Bertsekas, D. 1995. *Dynamic Programming and Optimal Control*, Vol II. Athena Scientific, Belmont, MA.
- Chen, F., J. Song. 2001. Optimal policies for multiechelon inventory problems with Markov modulated demand. *Oper. Res.* **49** 226–234.
- , Y.-S. Zheng. 1994. Lower bounds for multiechelon stochastic inventory systems. *Management Sci.* **40** 1426–1443.
- , J. K. Ryan, D. Simchi-Levi. 2000. The impact of exponential smoothing forecasts on the bullwhip effect. *Naval Res. Logist.* **47** 269–286.
- Clark, A., H. Scarf. 1960. Optimal policies for a multi-echelon inventory problem. *Management Sci.* **6** 475–490.
- Erkip, N., W. H. Hausman, S. Nahmias. 1990. Optimal centralized ordering policies in multiechelon inventory systems with correlated demands. *Management Sci.* **36** 381–392.
- Federgruen, A. 1993. Centralized planning models for multiechelon inventory systems under uncertainty, Chapter 3. S. C. Graves et al., eds. *Handbooks in OR & MS*, Vol. 4. Elsevier Science Publishers B.V., North-Holland, Amsterdam, The Netherlands.
- , P. Zipkin. 1984. Computational issues in an infinite horizon, multiechelon inventory model. *Oper. Res.* **32** 818–836.
- Fukuda, Y. 1964. Optimal policies for the inventory problem with negotiable leadtime. *Management Sci.* **10** 690–708.
- Gallego, G., Ö. Özer. 2001. Integrating replenishment decisions with advance demand information. *Management Sci.* **47** 1344–1360.
- , P. Zipkin. 1999. Stock positioning and performance estimation in series production-transportation systems. *Manufacturing & Service Oper. Management* **1** 77–87.
- Graves, S. C., D. B. Kletter, W. B. Hetzel. 1998. A dynamic model for requirement planning with application to supply chain optimization. *Oper. Res.* **46** 35–49.
- , H. C. Meal, S. Dasu, Y. Qiu. 1986. Two-stage production planning in a dynamic environment. S. Axsater, C. Schneeweiss, E. Silver, eds. *Multi-Stage Production Planning and Inventory Control*. Springer-Verlag, Berlin, Germany, 9–43.
- Gressens, B., C. Brousseau. 1999. The value propositions of dynamic pricing in business-to-business e-commerce. *CRM Project*, Vol. 1. http://www.crmproject.com/authors.asp?a_ID=158.
- Güllü, R. 1996. On the value of information in dynamic production/inventory problems under forecast evolution. *Naval Res. Logist.* **43** 289–303.
- Hariharan, R., P. Zipkin. 1995. Customer order information, lead times, and inventories. *Management Sci.* **41** 1599–1607.
- Heath, D., P. Jackson. 1994. Modeling the evolution of demand forecasts with application to safety stock analysis in production/distribution Systems. *IIE Trans.* **26** 17–30.
- Iglehart, D. L. 1963. Optimality of (s, S) policies in the infinite horizon dynamic inventory problem. *Management Sci.* **9** 259–267.
- Johnson, G., H. Thompson. 1975. Optimality of myopic inventory policies for certain depended processes. *Management Sci.* **21** 1303–1307.
- Kapuscinski R., S. Tayur. 1998. A capacitated production-inventory model with periodic demand. *Oper. Res.* **46** 899–911.
- Lovejoy, W. 1990. Myopic policies for some inventory models with uncertain demand distributions. *Management Sci.* **6** 724–738.
- Miller, L. M. 1989. Scarf's state reduction method, flexibility, and a dependent demand inventory model. *Oper. Res.* **34** 83–90.
- Özer, Ö. 2003. Replenishment strategies for distribution systems under advance demand information. *Management Sci.* **49** 255–272.
- Rosling, K. 1989. Optimal inventory policies for assembly systems under random demands. *Oper. Res.* **37** 565–579.
- Scarf, H. 1960. Bayes solutions of the statistical inventory problem. *Naval Logist. Res. Quart.* **7** 591–596.
- Scheller-Wolf, A. S., S. Tayur. 2001. A Markovian dual-source production-inventory model with order bands. Working paper, Carnegie Mellon University, Pittsburgh, PA.
- Schwarz, L. B., N. C. Petruzzi, K. Wee. 1998. The value of advance-order information and the implication for managing the supply chain: An information/control/buffer portfolio perspective. Working paper, Kranert Graduate School of Management, Purdue University, West Lafayette, IN.
- Sethi, P. S., F. Cheng. 1997. Optimality of (s, S) policies in inventory models with Markovian demand. *Oper. Res.* **45** 931–939.
- Song, J.-S., P. Zipkin. 1992. Evaluation of base-stock policies in multiechelon inventory systems with state dependent demands part I: State-in dependent policies. *Naval Res. Logist.* **39** 715–728.
- , ———. 1993. Inventory control in fluctuating demand environment. *Oper. Res.* **41** 351–370.
- , ———. 1996. Evaluation of base-stock policies in multiechelon inventory systems with state dependent demands part II: State-dependent policies. *Naval Res. Logist.* **43** 381–396.
- Toktay L. B., L. W. Wein. 2001. Analysis of a forecasting-production-inventory system with stationary demand. *Management Sci.* **47** 1268–1281.
- Veinott, A. 1965. Optimal policy for a multi-product, dynamic, non-stationary inventory problem. *Management Sci.* **12** 206–222.

The consulting Senior Editor for this manuscript was Garrett J. van Ryzin. This manuscript was received on February 3, 2003, and was with the author 35 days for 1 revision. The average review cycle time was 25 days.

Copyright 2003, by INFORMS, all rights reserved. Copyright of Manufacturing & Service Operations Management is the property of INFORMS: Institute for Operations Research and its content may not be copied or emailed to multiple sites or posted to a listserv without the copyright holder's express written permission. However, users may print, download, or email articles for individual use.