

Assemble-To-Order and Project Management under Throughput time Uncertainty: More Newsvendor Equations

Ton de Kok

Assemble-To-Order and Project Management under Throughput time Uncertainty: More
Newsvendor Equations

Problem Description

In High-Tech industry customer specific capital goods are manufactured. The implication of this customer specificity is that customers wait for some weeks to months before the product is delivered by the manufacturer. During this *customer lead time* the product is assembled from modules and parts, which are either taken from available stock or ordered from suppliers. The manufacturing process consists of several phases, some in parallel and some in sequence. Each phase produces a subassembly and the final phase produces the final product. In that way the network of phases is convergent. During the production of a subassembly several tests are used and sometimes rework is needed. The consequence is that the throughput time of each phase is different for each customer order. In many situations these differences cannot be predicted upfront from differences in the subassembly specifications. This implies that these differences should be considered as realizations of some underlying chance mechanism. In order to ensure timely delivery of the final product to the customer, the throughput time uncertainty must be absorbed by slack in the total planned throughput time of the product assembly from modules and parts to final product. The problem to be solved is to determine where and how much slack time is needed to ensure timely delivery of the product to the customer at minimum cost.

Key methodology

We solve the above described problem by quantitative modelling. We assume that each phase throughput time is a random variable with independent realizations. The phase throughput times are mutually independent. We assume that phase holding costs are charged from the moment the phase starts until the product is delivered to the customer. This cost structure

follows naturally from a cash flow perspective. At the start of each phase materials are released to be assembled in addition to the subassembly that is the input for this phase. During the phase additional labor and material costs are added and assuming these costs are uniformly added over time, these costs can be captured as a linear cost added at the start of the phase as well. These linear costs are only compensated at the moment the customer pays upon delivery of the product. Penalty costs are charged from the planned delivery time until the actual delivery of the product.

The slack time of each phase is determined by its *planned lead time*. Given the planned lead times of all phases and the precedence relationships between the phases, which constitute the convergent network, we can compute the planned start and completion times. The planned start times of each phase determine the release policy of the phase: if all preceding phases are completed at or before the planned start time of the phase under consideration, then this phase starts at its planned start time, otherwise it starts as soon as possible.

This *planned lead time problem* has been studied by several authors (e.g. Song et al. (2003)).

The contribution of this paper is that we conjecture a set of optimality equations, from which the optimal policy can be determined for any convergent structure. A formal proof of this conjecture is provided for the case studied by Song et al. (2003), which concerns a system consisting of N parallel phases that feed a single assembly phase. The proof is based on computing the marginal cost difference between the optimal policy and a set of carefully chosen policies that are perturbations of the optimal policy. Using the Kuhn-Tucker conditions on the cost function as a function of the planned lead times, we find the set of optimality equations.

Main results

These optimality equations can be written as generalized Newsvendor equations, each associated with a single phase. These equations generalize the ones derived in Dong et al.

(1994) from the serial planned lead times problem to the convergent problem. The optimality equations are mutually dependent, which implies that the convergent case is fundamentally more complicated than the serial case, which can be solved recursively phase by phase, solving single variable equations. The key observation that yields an efficient fixed point algorithm is that on a path from an arbitrary phase to the final phase, the influence of all other phases, assuming their planned lead times are optimal, can be modelled as an exogenous delay on the start of each phase on the path. Given these exogenous delays, the generalized Newsvendor equations can be recursively solved. Assuming initial exogenous delays equal to 0, we can iteratively solve the Newsvendor equations and recompute the exogenous delays. Our numerical experiments show, as stated above, fast convergence. Our comparison against generic non-linear optimization methods show that the solutions found are optimal or, in some cases, outperform the solutions found by the generic method, due to non-convexity of the cost function.

Managerial implications

The *generalized Newsvendor equations* conjecture implies that it is now possible to find optimal planned lead times for realistic situations. This has enabled us to identify that most of the slack time is positioned at the final stage, while other slack time is positioned to protect the critical path, which is primarily cost-based. This is partly explained by the fact that in realistic cases costs of phases are correlated with the length of the throughput times. Furthermore we briefly discuss potential application of our model to project networks with stochastic activity throughput times.

References

Song, J.S, Yano, C. and Lerssrisuriya, P., 2000, Contract Assembly: Dealing with Combined Supply Lead Time and Demand Quantity Uncertainty, MSOM, Vol. 2, No. 3, 287–296.