
Deep RL Dual Sourcing Inventory Management with Supply and Capacity Risk Awareness

Feng Liu
liudef@amazon.com

Ying Liu
liuyingy@amazon.com

Carson Eisenach
ceisen@amazon.com

Abstract

1 In this work, we investigate data-driven solutions aimed towards mitigating the
2 supply risk and optimizing cross-product constraints/resources for the inventory
3 management problem. We particularly focus on dual sourcing methodologies by
4 lead time for inventory planning and control. We investigate how to efficiently
5 employ Deep Reinforcement Learning (RL) models for learning and forecasting
6 the real-world supply chain processes under a range of assumptions. Moreover, we
7 also introduce a capacity control mechanism, designed to forecast shadow prices
8 given the capacity constraints of the inventory network. We highlight that instead
9 of directly modeling the complex physical constraints into the RL optimization
10 problem and solving the problem as a whole, our approach breaks down those
11 supply chain processes into different deep learning (DL) modules, leading to
12 improved performance on large real-world retail datasets. We also outline open
13 problems for future research to further investigate the efficacy of such models.

14 1 Introduction

15 Modern inventory management systems (IMS) in retail supply chain (such as Walmart and Amazon)
16 typically deploy multi-sourcing buying strategies, which employ a combined system with a just-
17 in-time (JIT) ordering strategy plus other specialized strategies that strive to achieve a balance
18 between supply shortage and inventory health across all products without compromising the inventory
19 management’s contribution to the overall retail service.

20 In real world supply chains, order quantities are subject to several post-processors, including modifi-
21 cation to meet vendor constraints such as minimum order and batch size constraints. Second, the
22 supply may be unreliable and vendors may only partially fill orders that they receive. This may occur
23 for multiple reasons, including that the vendor itself is out of stock. In the literature the proportion
24 of the original order quantity retailer ultimately receives is referred to as the yield or fill rate. In the
25 current state, there is no existing representation model to learn those external processes.

26 On the other hand, capacity limitations, both in terms of network inbound capacity and storage
27 capacity, require the supply chain system to assess which inventory to purchase and when to have it
28 arrive through all of its buying channels in order to have an effective and efficient capacity controller
29 (CC). For example, for JIT buys, traditional CC mechanisms typically have the ability to simulate
30 and compute weekly capacity costs to meet the available capacity for short forward-looking periods,
31 e.g. one to three months. However, as modern sourcing systems have changed to reflect increasing
32 complexity in global supply chain and the availability of direct-from-manufacturer cost discounts,
33 inventory decisions are made further into the future to cover planning horizons of longer periods, e.g.
34 more than six months.

35 This work is also motivated by prior works ([Madeka et al., 2022](#); [Andaz et al., 2024](#); [Eisenach
36 et al., 2024](#)) which established the viability of Deep RL for single-sourcing inventory planning
37 problem. But we forward this line of research by targeting for a more complex variant of the
38 aforementioned problem, wherein multi-sourcing channels of vendors are available and the primary

trade-off considered was to strike a balance for costs vs supply risks for different vendors. On the other hand, we aim to investigate effects of discount rates on long lead orders, stochastic quantity over time arrival profiles.

Here, we emphasize that a critical issue encountered in solving the multi-sourcing inventory problem via traditional methods such as dynamic programming is the unknown dynamics of a variety of underlying processes associated with inventory control. For instance, customer demand is not deterministic and exhibits volatility which are influenced by seasonality, external sourcing processes, etc. To incorporate those complex supply chain processes, we attempt to investigate whether it is possible to efficiently scale the training of decision policies from the vast amount of available historical data, without the requirement to build separate models for the various state variables with unknown stochastic behaviors.

We organize rest of the paper as summarized next. In Section 2, we mathematically formulate the dual-sourcing inventory management problem and thereby describe our dual sourcing RL methodology as a solution for the problem. In Section 3, we present our main results from numerical experiments. Additionally, in Section 4, we extend our dual sourcing RL baseline model by introducing a capacity mechanism control strategy and provide evaluation results. Finally, we conclude the paper with discussion on future work in Section 5.

2 A Deep RL Approach for the Dual Sourcing Problem

2.1 Modeling the Dual Sourcing Problem as an Exo-IDP

In this section we model the dual sourcing problem as an *Exogenous Interactive Decision Process* (Exo-IDP) following the framework introduced in Madeka et al. (2022). In an Exo-IDP, the state describing the system dynamics decompose into (i) *Exogenous* processes, which evolve independently of the buy policy’s actions, and (ii) *Endogenous* processes that are impacted by the agent’s actions and the exogenous signals.

We consider the case of a retailer managing a set $\mathcal{A}$ of products for T time steps, where the objective is to maximize revenue by placing orders to both *long lead time* (LLT) and JIT sources. To succinctly describe our process, we focus on just one product $i \in \mathcal{A}$, though we note that decisions can be made jointly for every product.

State The price received at sale, costs incurred on JIT purchase, and holding costs are denoted as p_t^i , $c_t^{J,i}$, and h_t^i , respectively. For LLT orders, the retailer typically receives a discount on the cost of goods sold, so there is a different cost incurred on purchase $c_t^{L,i}$. Additionally, the demand for product i at time t is denoted as d_t^i . The aforementioned set of variables are completely exogenous, and therefore their evolution is independent of any policy’s interaction with the Exo-IDP.

Together these exogenous processes form the state as follows,

$$s_t^i \triangleq (d_t^i, p_t^i, c_t^{J,i}, c_t^{L,i}, \rho_t^{J,i}, \rho_t^{L,i}, M_t^{J,i}, M_t^{L,i}). \quad (1)$$

The history of the joint process up to time t is defined as

$$H_t := \{(k_0^i, s_1^i, \dots, s_t^i)\}_{i=1}^{|\mathcal{A}|},$$

where k_0^i is the initial inventory level. Product-level histories can be defined similarly.

Actions For product i , action a_t^i implies placing orders via JIT, LLT channels at time t . In other words, agent’s interaction with the Exo-IDP is only via placing orders. More precisely,

$$a_t^i \triangleq (q_t^{J,i}, q_t^{L,i}). \quad (2)$$

So we have $a_t^i \in \mathbb{R}^2$. For a class of policies parameterized by θ , we can defined the actions as

$$a_t^i = \pi_{\theta,t}^i(H_t). \quad (3)$$

We define the set of these policies as $\Pi \triangleq \{\pi_t^i(\cdot; \theta) | \theta \in \Theta, i \in \mathbb{A}, t \in [0, T-1]\}$.

External Sourcing Processes After orders are created and submitted to vendors, they can arrive in multiple shipments over time, and the total arriving quantity may not necessarily sum up to the order quantity placed. Any constraints the vendor imposes on the retailer's orders $\mathbf{M}_t^{J,i} \in \mathbb{R}^{d_v}$ – such as minimum order quantities and batch sizes – are exogenous to the ordering decisions. We define an order quantity post-processor on the JIT source $f_p^J : \mathbb{R}_{\geq 0} \times \mathbb{R}^{d_v} \rightarrow \mathbb{R}_{\geq 0}$ that may modify the order quantity. The final order quantity submitted to the vendor is denoted as $\tilde{q}_t^{J,i} := f_p^J(q_t^{J,i}, \mathbf{M}_t^{J,i})$. Similarly, the final LLT order quantity is defined as $\tilde{q}_t^{L,i} := f_p^L(q_t^{L,i}, \mathbf{M}_t^{L,i})$.

At every time t , the vendor has allocated a supply $U_t^{J,i}$ that denotes the maximum number of units it can send (regardless the amount we order), which will arrive over from the current week up to L_1 weeks in the future according to an exogenous *arrival shares* process $(\rho_{t,0}^{J,i}, \dots, \rho_{t,L_1}^{J,i})$ where $\sum_l \rho_{t,l}^{J,i} = 1$ and $\rho_{t,l}^{J,i} \geq 0$ for all i, t and l . The arrival quantity at lead time j from order $q_t^{J,i}$ can be denote as $o_{t,j}^{J,i} := \min(U_t^{J,i}, \tilde{q}_t^{J,i}) \rho_{t,j}^{J,i}$. The LLT arrival quantities are defined similarly and denoted as $o_{t,j}^{L,i} := \min(U_t^{L,i}, \tilde{q}_t^{L,i}) \rho_{t,j}^{L,i}$.

In brief, the overall sourcing processes from the initial order quantity to final arrivals can be modeled as

$$o_{t,j}^{J,i} := \min(U_t^{J,i}, f_p(q_t^{J,i}, \mathbf{M}_t^{J,i})) \rho_{t,j}^{J,i}, \quad (4)$$

$$o_{t,j}^{L,i} := \min(U_t^{L,i}, f_p(q_t^{L,i}, \mathbf{M}_t^{L,i})) \rho_{t,j}^{L,i}. \quad (5)$$

Internal Inventory Dynamics I_{t-}^i and I_t^i denote the on-hand inventory for product i at the beginning and end of a period t , respectively. The inventory update rule is given as follows,

$$I_{t-}^i = I_{t-1}^i + \sum_{j=0}^{L_1} o_{t,j}^{J,i} + \sum_{j=0}^{L_2} o_{t,j}^{L,i}, \quad (6)$$

and

$$I_t^i = \max\{I_{t-}^i - d_t^i, 0\}. \quad (7)$$

Reward Function We formulate the reward realized at t taking into account the current period inventory costs and sales. The construction of reward in our problem is such that it measures the periodic cash outflows caused due to replenishment and holding costs of inventory, and inflows are attributed to customer sales. It is defined as

$$R_t^i \triangleq p_t^i \min(d_t^i, I_{t-}^i) - c_t^{J,i} \sum_{j=0}^{L_1} o_{t,j}^{J,i} - c_t^{L,i} \sum_{j=0}^{L_2} o_{t,j}^{L,i} - h_t^i I_t^i. \quad (8)$$

Hence, computation of reward is essentially a function of history vector H_t and policy parameters θ , $R(\mathcal{H}_t, s_t^i, \theta)$.¹

2.2 The Dual Sourcing Optimization Problem

In our setting, we aim to maximize the total discounted² reward across the T length time horizon in expectation, while accounting for other constraints.

In the following, we state the dual sourcing optimization problem $\mathcal{P}_1$,

$$\mathcal{P}_1 : \max_{\theta \in \Theta} \mathbb{E} \left[\sum_{i \in \mathbb{A}} \sum_{t=0}^{T-1} \gamma^t R_t^i(\theta) \right], \quad (9)$$

s.t.

$$I_0^i = \bar{I}^i, \quad (10)$$

$$\text{Equations}(3 - 7), \quad (11)$$

where Eq. (9) is the expression for initial inventory.

¹Reward is a function of current period action a_t^i which is computed via policy π parameterized by θ , and input $\mathcal{H}_t$.

²This discounted reward implies time value of actual cash flows. The discounting factor here is γ .

2.3 Scaling the Learning by Forecasting Sourcing Processes

In practice, we do not actually observe the full supply and arrival share processes for either the JIT or LLT sources. Under the IDP model described in the previous section, we only observe the arrivals share processes when an order was placed historically and we only observe the supply process when we do not receive the full order quantity. To handle this censoring we directly forecast the *arrivals* (Andaz et al., 2024) instead of the supply and arrival shares processes.

To see why this makes sense, note that the dynamics (6) and reward function (8) depend only on the arrivals $o_{t,j}^{J,i} := \min(U_t^{J,i}, f_p(q_t^{J,i}, \mathbf{M}_t^{J,i}))\rho_{t,j}^{J,i}$ and $o_{t,j}^{L,i} := \min(U_t^{L,i}, f_p(q_t^{L,i}, \mathbf{M}_t^{L,i}))\rho_{t,j}^{L,i}$. Thus, for the purposes of constructing our simulator from historic data, we forecast arrivals conditional on the action a_t^i rather than estimating the post-processing behavior and the supply and arrival share processes.

Arrivals For arrivals, denoting by $H_{t,O}^i$ the observed components of H_t^i , one can forecast

$$p(o_{t,0}^{J,i}, \dots, o_{t,L_1}^{J,i} | H_{t,O}^i; \psi_1). \quad (12)$$

Note that conditioning on $H_{t,O}^i$, θ is equivalent to conditioning on $H_{t,O}^i$, ψ_1 and the past JIT order actions a_s^i for $s \leq t$.

2.4 The Deep RL Algorithm for Training the Buy Policy

Firstly, recall that parameters $\theta \in \Theta$ at any time t essentially dictates the RL agent’s policy as expressed in eq. (3), therefore our problem reduces to learning optimal parameters $\theta^* \in \Theta$. Specifically, we leverage a Deep Neural Network (DNN) architecture for $\pi(\cdot, \cdot; \theta)$. So, θ in our DRL framework essentially implies the weights and parameters of the constituent neurons in the policy network. We applied the same DNN architecture as used in Madeka et al. (2022). The detailed training algorithm can be found in Appendix B.

3 Experimental Evaluations

3.1 Training/Evaluation Configurations

Real-world Dataset We use the same real-world dataset as used in Madeka et al. (2022) with approximately 80,000 products for 124 weeks from April 2017 to August 2019. Out of the 124 weeks, we treat the first 72 weeks as training dataset and the remaining 52 as the backtesting dataset. However, in future iterations, we plan to train the DRL models on 104 weeks so that the policy agents can better track seasonal patterns.

Baselines We compare our Dual Sourcing RL (DualSrc-RL) buy policy against several baselines, i.e. BaseStockHorizonTip (BSHT), *Tailored Base Surge (TBS)*, *Just-in-time RL (JIT-RL)*. The first two are classic operation research baselines and the *JIT-RL* is the RL baseline where the RL model is trained as a single-sourcing model (Madeka et al., 2022). More detailed description of those baselines are reported in Appendix A.1.

3.2 Main Results and Analytics

We use the training dataset to train the RL algorithms and perform evaluation experiments for the compared algorithms over the test dataset. The evaluation results are reported in Table 1.

Setting	Method	% of BSHT
JIT Policy	BSHT	100
	JIT-RL	104.78
Dual Sourcing	TBS	117.69
	DualSrc-RL	121.54

Table 1: Cumulative discounted rewards (as % of BSHT) for 52 backtest periods for different policy methods.

144 From the results above, we can observe that overall dual sourcing strategies DualSrc-RL, TBS are
 145 favorable in terms of rewards. Furthermore, DualSrc-RL is most profitable in all the run scenarios,
 146 and has 4% reward gains over TBS.

147 4 Capacity Management with Neural Coordination

148 Having formulated the unconstrained inventory control problem in $\mathcal{P}_1$, now we consider a constrained
 149 situation, where network capacity constraints are introduced, and considered as part of the exogenous
 150 process. We are interested in studying the constrained problem because it is a challenging variant in
 151 real-world supply chain applications. A large retailer typically manages a supply chain for multiple
 152 products and has limited resources (such as storage) that are shared amongst all the products that
 153 retailer stocks. Specifically, we have the following set of formulas representing the storage capacity
 154 constraints,

$$G := \{K_0, K_2, \dots, K_{T-1}\}. \quad (13)$$

155 The constrained problem $\mathcal{P}_1$ can be obtained by adding the capacity constraints (15) to $\mathcal{P}_1$,

$$\mathcal{P}_2 : \max_{\theta \in \Theta} \mathbb{E} \left[\sum_{i \in \mathbb{A}} \sum_{t=0}^{T-1} \gamma^t R_t^i(\theta) \right], \quad (14)$$

s.t.

$$\sum_{i \in A} v^i I_t^i \leq K_t, . \quad (15)$$

156 For the reward function, we modify the unconstrained reward to incorporate a penalty according to
 157 the capacity prices,

$$R_t^{\lambda, i} \triangleq R_t^i - \lambda_t v^i I_t^i, \quad (16)$$

158 where λ_t is the storage capacity price at time t .

159 4.1 Forecasting the Capacity Prices by a Neural Coordinator

160 The role of coordinator agent is to predict capacity prices for the targeting inventory network given
 161 any capacity constraints as input. One example of a coordination mechanism is model predictive
 162 control (MPC), which would use forecasted demand to perform a dual cost search for the next L
 163 periods. However, for an RL buying policy that uses many historical features, model predictive
 164 control would require forward simulating many features and it may be difficult to model the full joint
 165 distribution of all these features.

166 Instead, we apply a coordination approach as used in (Eisenach et al., 2024) by introducing a deep
 167 learning model for this problem. Specifically, we train a neural network to forecast the future prices
 168 of capacity that would be required to constrain the dual sourcing RL buying policy. Specifically, we
 169 learn a neural network to predict,

$$(\lambda_t, \lambda_{t+1}, \dots, \lambda_{t+L}) = \phi_\omega(H_t, G), \quad (17)$$

170 where H_t denotes the historical state vector and G denotes the capacity constraints.

171 In the next, we introduce the optimization problem for learning the neural coordinator. First, assume a
 172 fixed dual sourcing buying policy θ . Below we describe the ways in which the coordinator's Exo-IDP
 173 deviates from the buying agent's Exo-IDP.

174 The coordinator solves the following problem:

$$\mathcal{P}_3 : \min_{\omega \in \Omega} \mathbb{E} \left[\sum_{t=0}^{T-1} \left(\sum_{i \in A} v^i I_t^i - K_t \right)_+^2 + \|\lambda_t\| + \mathcal{L}(\lambda_t, H_t^\lambda) \right], \quad (18)$$

$$(19)$$

175 where $\mathcal{L}(\lambda_t, H_t^\lambda)$ denotes the total capacity price forecast error at time t ,

$$\mathcal{L}(\lambda_t, H_t^\lambda) = \sum_{s=1}^L \|\lambda_t - (\hat{\lambda}_{t-s}^L)_{L-s}\|^2, \quad (20)$$

176 which is the mean squared error (MSE) of all past forecasts for the current cost. The training algorithm
177 for the neural coordinator can be found in Appendix.

178 4.2 Evaluation Results

179 In this section, we backtest our proposed dual sourcing buy policy with neural coordinator for storage
180 capacity management.

181 **Compared Policies and Baselines** We compare our unconstrained Dual Sourcing RL (DualSrc-RL)
182 buy policy (as the baseline) with its two constrained variants. For the first variant, the neural
183 coordinator is used to constrain the DualSrc-RL agent, and a MPC predictor is used to constrain
184 DualSrc-RL as the second variant.

185 **Performance Metrics** In addition to measuring reward achieved by the various policies, we consider
186 several additional measures of constraint violation. They are (M1) mean constraint violation, (M2)
187 mean violation on weeks where either unconstrained policy met at least 90% of the limit, (M3)
188 percent of weeks where the violation exceeded 10% and (M4) percent of weeks where constraint
189 violation exceeded 10% and either unconstrained policy exceeded 90% of the capacity limit.

190 **Results** The Table 2 shows the results on the out-of-training backtesting period, where each combina-
191 tion of policy and coordinator were evaluated against 100 sampled storage constraint paths. Note
192 that under all metrics (both violation and reward) the DualSrc-RL policy with neural coordinator
193 outperforms DualSrc-RL with MPC. Although some of the violation metrics seem somewhat high,
194 one should keep in mind that many of the capacity curves sampled will be highly constraining, much
195 more so than in a real-world setting – in practice if the supply chain were so constrained, one would
196 build more capacity.

Table 2: Out-of-distribution evaluation results.

Initialization	Buy Algo	Coordinator	Violations				Reward
			M1	M2	M3	M4	
	DualSrc-RL	-	28.9%	54.7%	35.8%	67.8%	100
Onhand +	DualSrc-RL	Neural	3.5 %	6.7%	9.1%	17.3%	99.7
Inflight	DualSrc-RL	MPC	13.3%	25.1%	18.2%	34.1%	96.9

197 The Figure 1 below shows two examples of trajectories in the evaluation period for our DualSrc-RL
198 (DS-RL) buy policies with their coordinator settings (unconstrained, Neural coordinator or MPC).
199 We can see that the neural coordinator is able to constrain the on hand inventory within the capacity
200 limit on the out-of-training backtesting period.

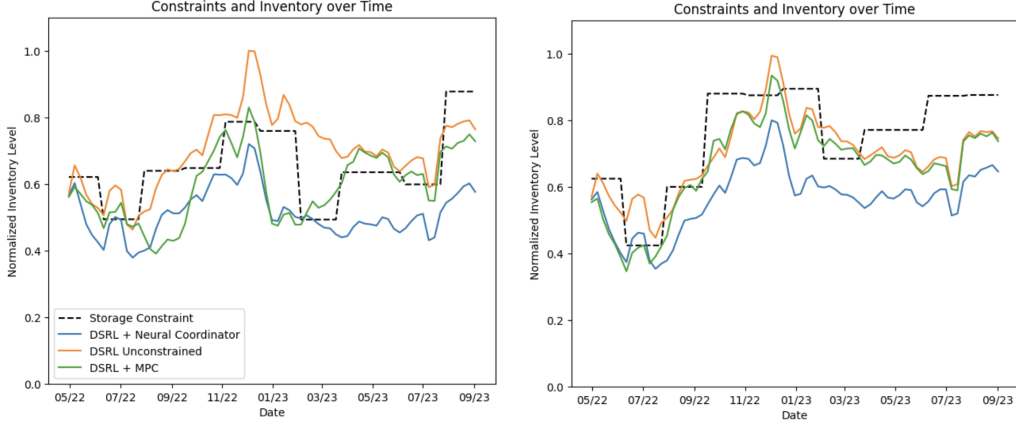


Figure 1: Inventory trajectories under different constraint paths during the out-of-training period.

5 Conclusion and Future Work

In this paper, we investigate ML-based inventory control methodologies with the consideration of stochastic supply chain processes that introduce scalable representation models to mitigate the supply risk and optimize the complex cross-product constraints/resources. Specifically, we investigate how to efficiently build a Deep RL framework for the problem by forecasting the real-world supply chain processes under a range of assumptions. We also introduce a capacity control mechanism, designed to forecast shadow prices given the capacity constraints of the inventory network. We highlight that instead of directly modeling the complex physical constraints into the learning pipeline and solving the problem as a whole, our approach breaks down those supply chain processes into different DL modules, leading to improved performance on larger real-world retail datasets. For future research, it is interesting to develop efficient modules for approximating other cross-product dependencies in real-world supply chain networks, e.g. containerization processes in global transportation, truck-load and placement.

References

- ANDAZ, S., EISENACH, C., MADEKA, D., TORKKOLA, K., JIA, R., FOSTER, D. and KAKADE, S. (2024). Learning an inventory control policy with general inventory arrival dynamics. [arXiv:2310.17168](#).
- BALAJI, B., BELL-MASTERSON, J., BILGIN, E., DAMIANOU, A., GARCIA, P. M., JAIN, A., LUO, R., MAGGIAR, A., NARAYANASWAMY, B. and YE, C. (2019). Orl: Reinforcement learning benchmarks for online stochastic optimization problems. [arXiv:1911.10641](#).
- BOYD, S., PARIKH, N., CHU, E., PELEATO, B., ECKSTEIN, J. ET AL. (2011). Distributed optimization and statistical learning via the alternating direction method of multipliers. *Foundations and Trends® in Machine learning* **3** 1–122.
- EISENACH, C., GHAI, U., MADEKA, D., TORKKOLA, K., FOSTER, D. and KAKADE, S. (2024). Neural coordination and capacity control for inventory management. [arXiv:2410.02817](#).
- FUKUDA, Y. (1964). Optimal policies for the inventory problem with negotiable leadtime. *Management Science* **10** 690–708.
- MADEKA, D., TORKKOLA, K., EISENACH, C., LUO, A., FOSTER, D. and KAKADE, S. (2022). Deep inventory management. [arXiv:2210.03137](#).
- MAGGIAR, A., DICKER, L. and MAHONEY, M. (2024). Consensus planning with primal, dual, and proximal agents. *arXiv preprint arXiv:2408.16462*.
- WHITTEMORE, A. S. and SAUNDERS, S. (1977). Optimal inventory under stochastic demand with two supply options. *SIAM Journal on Applied Mathematics* **32** 293–305.
- XIN, L. and GOLDBERG, D. A. (2018). Asymptotic optimality of tailored base-surge policies in dual-sourcing inventory systems. *Management Science* **64** 437–452.

A Related work

Dual Sourcing Inventory Control: Dual sourcing is a heavily studied problem in inventory control literature, and has become a common practice supply chain organizations worldwide. In vanilla dual sourcing problem, the sourcing channels i.e., JIT, LLT essentially cause a trade-off in lead time vs ordering costs since usually LLT (direct import sourcing channel) has lower sourcing cost. In this context, Tailored Base-Surge (TBS) policy is a widely used method, wherein at each period, a constant order is placed via the LLT channel, and a dynamic order is placed to match an inventory order-up-to level via the JIT channel respectively. The JIT channel’s order-up-to level essentially implies a *Safety Stock* maintained to address sudden *demand surges* (Xin and Goldberg, 2018). Although, such policies present fairly intuitive approaches to handle the dual sourcing problem, but the optimality analysis of such policies has been a hard problem in general (Whittemore and Saunders, 1977), except for certain edge cases. Interestingly, the optimal policy is shown to be vanilla base-stock policy when the lead time difference between LLT and JIT channels is exactly 1 period (Fukuda, 1964).

Reinforcement Learning for Inventory Control: Recently, application of Reinforcement Learning (RL) methods to produce ordering decisions in large scale inventory management systems has gained significant attention. In this context, Deep RL approaches have emerged more recently over its other *Model-Based* counterparts, due to improved computational scalability and generalization performance of Deep Neural Network (DNN) architectures. Furthermore, DRL methods have been shown to achieve performance gains over benchmarks for the multi-period newsvendor problem under fairly realistic assumptions on costs, prices, demand and stationarity (Balaji et al., 2019). It is worth highlighting that DRL for single sourcing problem has been comprehensively investigated with extensive empirical evidences at Amazon (Madeka et al., 2022).

Capacity Management and Coordination Mechanism Retailers typically manage a supply chain for multiple products and has limited capacity resources (such as storage) that are shared amongst all the products that retailer stocks. The classic method for handling capacity constraints is to call a simulation and optimization process to compute shadow prices on the shared resources. Model predictive control (MPC) consists of using a model to forward simulate a system to optimize control inputs and satisfy any constraints. At each time step, one re-plans based on updated information that has become available in order to select the next control input. Recent work Maggiar et al. (2024) introduced the Consensus Planning Protocol (CPP), which targets problems where multiple agents (each of which is locally optimizing its own utility function) all consume a shared resource. This is closely related to a distributed ADMM procedure (Boyd et al., 2011). Another work Eisenach et al. (2024) presented a new capacity control mechanism for RL-based buying policies and proposed a Neural Coordinator model to generate forecasts of capacity prices. Their formulation of capacitated inventory management can be viewed as a special case of CPP (a central coordinator adjusts prices, and the other agents adjust their plans).

A.1 Baseline Methods

A.1.1 Improved Tailored Base Surge Policy

We adopt a modified version of vanilla Tailored Base Surge (TBS) policy which has been described in Xin and Goldberg (2018). For each product i , TBS policy will place a *dynamic* LLT order $q_{L,t}^{i,TBS}$ every period. In our problem setting, this LLT order will arrive with lead-time of δ_L .

Whereas, for orders through JIT channel, this TBS policy will first calculate a dynamic target order-up-to-level $I_t^{i,Tip}$ via the production “Horizon Tip Calculator” method. Consequently, TBS policy places also places a dynamic order for the JIT channel i.e., $q_{J,t}^{i,TBS}$ that can bring back the inventory level to $I_t^{i,Tip}$. With on-hand inventory at the end of the period, I_t^i , the JIT order at time t is given by:

$$q_{J,t}^{i,TBS} = \max \left\{ 0, I_t^{i,Tip} - I_t^i - \left[\sum_{\tilde{t}=0}^{t-1} \sum_{k=\tilde{t}}^{t+\delta_t^{i,Pred}} (o_{\tilde{t},k-\tilde{t}}^{J,i} + o_{\tilde{t},k-\tilde{t}}^{L,i}) \right] \right\}, \quad (21)$$

where $\delta_t^{i,Pred}$ is the median forecasted VLT at time t for product i . In other words, Improved TBS subtracts on-hand and inflight inventory from the Horizon Tip to compute JIT orders for current

period. The order quantities $q_{J,t}^{i,\text{TBS}}, q_{L,t}^{i,\text{TBS}}$ will arrive according to the underlying quantity over time arrival models.

Choice of LLT Order Input: We set LLT order quantities $q_L^{i,\text{TBS}}$ as scaled 12-week rolling mean of product-level demands from training set, i.e., $q_{L,t}^{i,\text{TBS}} = \alpha \cdot \frac{1}{12} \cdot \sum_{i=t-11}^{t-1} d_t^i$. The order scaling factor α is used in our experiments as a search parameter for getting optimal TBS policy.

A.1.2 Other Baseline Policies

1. We use a single source JIT Base Stock Policy with orders dictated by eq. (21) which we call BaseStockHorizonTip (BSHT).
2. The existing RL buying policy for JIT single source from the literature (Madeka et al., 2022).

B Training Algorithms

B.1 Training Algorithm for the Buy Policy

Observe that constraints of $\mathcal{P}_1$ (9) are in fact definitions for different components of the reward function, and, therefore can be omitted from the optimization problem formulation otherwise. Next, we present the Direct Backpropagation (DirectBP-DualSrc) training algorithm for DRL policy.

Algorithm 1 Direct Backpropagation DRL training algorithm (DualSrc-RL)

Input: set of products $\mathcal{A}$, training batch size M , step size: $\eta, \theta_0 \in \Theta$.
Initialize: Batch Iterator: $b \leftarrow 1$.
while θ is not converged **do**
 Sample mini-batch of products $\mathcal{A}_b$ of size M from $\mathcal{A}$. and set $R^b \leftarrow 0$.
 for $i \in \mathcal{A}_b$ **do**
 $R^i \leftarrow 0, I_0^i \leftarrow \bar{I}^i$.
 for $t = 0, \dots, T^{\text{Train}} - 1$ **do**
 Place orders $(q_{J,t}^i, q_{L,t}^i) = \pi_t^i(H_t; \theta_{b-1})$.
 Sample JIT, LLT arrivals $(o_{t,0}^{J,i}, \dots, o_{t,L_1}^{J,i}, o_{t,0}^{L,i}, \dots, o_{t,L_2}^{L,i})$.
 Update inventory I_t^i according to (6) and (7).
 Collect reward R_t^i according to (8)
 $R^i \leftarrow R_t^i + \gamma^t \cdot R_t^i$.
 end for
 $R^b \leftarrow R^b + R^i$.
 end for
 $\theta_b \leftarrow \theta_{b-1} + \eta \cdot \nabla_{\theta} \mathcal{P}_1^b|_{\theta=\theta_{b-1}}$. // Update Parameters of Policy Network $\pi(\cdot, \cdot; \theta)$.
 $b \leftarrow b + 1$.
end while

B.2 Training Algorithm for the Neural Coordinator

We implement of the training algorithm for learning the neural coordinator for our dual sourcing problem. Specifically, we train the coordinator by solving the problem ($\mathcal{P}_3$). Similar to the training algorithm for the dual sourcing buy policy, we apply the Direct Backpropagation algorithm to optimize the loss objective 18 in ($\mathcal{P}_3$). The pseudo code of the training algorithm is shown in Algorithm 2.

C Input Features

C.1 Featurization for Buying Policy

In terms of features, we mainly use the following feature list provided to the buying policy:

1. The current inventory level
2. Previous actions aiu that have been taken

Algorithm 2 Training algorithm for the Neural Coordinator

Input: set of products $\mathbb{A}$, training batch size M , step size: η , given buy policy θ , $\{\bar{I}^i\}_{i \in \mathbb{A}}$, δ_L , ϵ , initial neural coordinator $\omega_0 \in \Omega$.
Initialize: Batch Iterator: $b \leftarrow 1$.
while stop criterion is not satisfied **do**
 Sample mini-batch of products $\mathbb{A}_b$ of size M from $\mathbb{A}$. and set $R^b \leftarrow 0$.
 for $i \in \mathbb{A}_b$ **do**
 $R^i \leftarrow 0$, $I_0^i \leftarrow \bar{I}^i$.
 for $t = 0, \dots, T^{\text{Train}} - 1$ **do**
 Collect instantaneous state and history vector $s_t^i, \mathcal{H}_t$.
 Place orders $(q_{J,t}^i, q_{L,t}^i) = \pi_t^i(\mathcal{H}_t, s_t^i; \theta)$.
 Sample JIT, LLT arrivals $(o_{t,0}^{J,i}, \dots, o_{t,L_1}^{J,i}, o_{t,0}^{L,i}, \dots, o_{t,L_2}^{L,i})$.
 Collect reward R_t^i and update inventory I_t^i .
 end for
 end for
 Update global state and compute coordination loss by Equation 18.
 $\omega_b \leftarrow \omega_{b-1} + \eta \cdot \nabla_{\omega} \mathcal{P}_3^b|_{\omega=\omega_{b-1}}$. // Update Parameters of the Neural Coordinator $\phi(\cdot, \cdot; \omega)$.
 $b \leftarrow b + 1$.
end while

- 308 3. Time series features
 - 309 (a) Historical availability corrected demand
 - 310 (b) Distance to public holidays
 - 311 (c) Historical website glance views data
- 312 4. Static product features
 - 313 (a) Product group
 - 314 (b) Text-based features from the product description
 - 315 (c) Brand
 - 316 (d) Volume
- 317 5. Economics of the product - (price, cost etc.)
- 318 6. Capacity costs – past costs and forecasted future costs

319 C.2 Featurization for Neural Coordinator

320 The neural coordinator takes the following aggregate/population level features:

- 321 1. Aggregated actions, inventory, demands for all current and previous times
 - 322 (a) Order quantities
 - 323 (b) Inventory
 - 324 (c) Availability corrected demand
 - 325 (d) Inbound
 - 326 (e) All the above, but weighted by inbound and storage volumes
- 327 2. Forecasted quantities for next L weeks.
 - 328 (a) Mean demand at lead time
 - 329 (b) Inventory after expected drain at lead time
 - 330 (c) All the above, but weighted by inbound and storage volumes
- 331 3. Other time series features
 - 332 (a) Distance to public holidays
 - 333 (b) Mean economics of products - (price, cost etc.), weighted by demand and volumes
- 334 4. Capacity costs (past costs and forecasted future costs)