

Zapasy w przedsiębiorstwie

1. Zasady zarządzania zapasami
 - a. Minimum całkowitych kosztów działania systemu zarządzania zapasami
 - b. Koncepcja *Just-In-time*
 - c. Metoda *Pull*
 - d. Metoda *Push*

Materiały do nauki:

1. Notatka
2. Repetytorium A.30 str. 69-70

Zadanie na ocenę (prześlij w pliku doc. na adres: kamila.hecold@gmail.com, tytuł wiadomości: nazwa przedmiotu, klasa, imię i nazwisko)

Notatka

MODELE KSZTAŁTOWANIA ZAPASÓW W PRZEDSIĘBIORSTWIE

Klasyczne (tradycyjne) modele zarządzania zapasami służą do precyzyjnego określenia wielkości partii dostawy i wyznaczenia momentu złożenia zamówienia dla danego dobra materialnego (towaru). Nie uwzględniają one powiązań pomiędzy poszczególnymi towarami lub rodzajami zapasów występujących w obrębie przedsiębiorstwa, traktując zapas konkretnego dobra w określonym miejscu w sposób autonomiczny. Algorytmy wyznaczania dwóch podstawowych parametrów tych modeli, tj. wielkości partii oraz okresu między zamówieniami, w klasycznych modelach zarządzania zapasami bazują na znalezieniu minimum funkcji całkowitych kosztów działania systemu. Funkcję tę można przedstawić następująco [Carlson 1987]:

$$C = K_1 * C_1 + K_2 * C_2 + K_3 * C_3$$

gdzie:

C - całkowity koszt związany z zapasami,
C₁ - koszt nadmiaru zapasu,
C₂ - koszt braku zapasu,
C₃ - koszt inicjacji procesu zamawiania (np. koszty złożenia zamówienia, koszty planowania),
K_i - współczynniki zależne od innych parametrów uwzględnianych w danym modelu np. od wielkości partii.

Klasyczne modele sterowania zapasami przez długi okres były stosowane zarówno na wejściu do przedsiębiorstwa, a więc w fazie zaopatrzenia, jak również na wyjściu, a więc w fazie dystrybucji. Rozwój metod zarządzania opartych na zasadach współczesnej logistyki doprowadził do uwypuklenia współzależności pomiędzy procesami planowania produkcji a zarządzania zapasami. W planowaniu produkcji, uwzględniającej jednocześnie planowanie poziomów zapasów, można wyróżnić dwa podejścia [Bonney 1994]:

1. Ustalony zostaje poziom zapasów, a następnie opracowuje się tak plany produkcji, aby ten poziom był zawsze utrzymany.

2. Za punkt wyjścia przyjmuje się założony plan produkcji, z którego wynika poziom utrzymywanych zapasów. Wielkość tych zapasów powiększana jest o planowany poziom zapasu zabezpieczającego.

Alternatywą statystycznego sterowania zapasami, opartą na pierwszym podejściu, jest metoda planowania potrzeb materiałowych, która umożliwia dokładne określenie zapotrzebowania na dany surowiec lub część. Przykładem drugiego podejścia jest strategia JiT (ang. *Just-in-Time*), która polega na maksymalnym zsynchronizowaniu w procesie produkcji momentu dostaw materiałów (elementów, podzespołów, itp.) do danego stanowiska roboczego z momentem zaistnienia na nie potrzeby (popytu) [Sarjusz-Wojski 2000]. JiT zakłada zatem, że zapasy powinny być traktowane jako ewidentna strata i stawia sobie za cel skrócenie cykli dostaw i produkcji.

Sposobem klasyfikacji modeli zarządzania zapasami najczęściej spotykanym w literaturze [Bonney 1994, Coyle, Bardi i Langley 2002, Forgarty, Balcstone i Hoffmann 1991] jest wyróżnienie systemów typu *pull* oraz *push*. Systemy typu *pull*, zwane również "systemami ssącymi" [Terminologia logistyczna 1995], określane są mianem systemów reaktywnych, ponieważ opierają się na obserwowanym popycie rynkowym, w przeciwieństwie do systemów typu *push*, pro aktywnych, bazujących na popycie antycypowanym. W systemach typu *pull* zapotrzebowanie zgłaszane jest do dostawcy przez ogniwo znajdujące się najbliżej rynku w łańcuchu dostaw, wymuszając w ten sposób przepływ towarów przez cały łańcuch (na zasadzie reakcji łańcuchowej). Podstawową cechą takiego rozwiązania jest możliwość szybkiej reakcji na nagłe i niespodziewane zmiany popytu. W grupie modeli należących do tej grupy wyróżniamy: klasyczne modele zarządzania zapasami oraz modele oparte na strategii JiT. Systemy typu *pull* mogą funkcjonować autonomicznie, a ich największą wadą jest brak możliwości skoordynowania popytu zgłaszanego przez poszczególne ogniwa sieci dystrybucji, gdyż zamówienie generowane jest na podstawie stanu zapasów w danym miejscu lokalizacji. Przepływ informacji w tego typu systemach odbywa się najczęściej tylko w jednym kierunku, między odbiorcą a dostawcą. Przeciwnieństwem są systemy typu *push*, zwane inaczej "tłoczącymi" [Terminologia logistyczna, 1995], które opierają się na szczegółowych planach powstających na podstawie prognoz zapotrzebowania, sporządzanych dla wszystkich ogniw znajdujących się najbliżej rynku. Prognozy te są sumowane i służą do tworzenia harmonogramów produkcyjnych. Dostępny zasób jest tłoczony do poszczególnych ogniw w kierunku rynku. Przepływ informacji w takich systemach odbywa się dwukierunkowo. Przykładem tego typu modeli są modele oparte na metodzie MRP.

W obszarze dystrybucji można spotkać również tradycyjne modele zarządzania zapasami, które zaadoptowały elementy charakterystyczne dla systemów typu *push*. Przykładem takiego podejścia może być model zwany modelem zapasu podstawowego [Forgarty, Balcstone i Hoffmann 1991] lub model ciągłego uzupełniania zapasów (ang. *continuous replenishment*) [Continuous Replenishment 1996, Hałas i Swarczewicz 1996]. Idea tych systemów polega na tym, że każdy punkt sprzedaży detalicznej należący do danego przedsiębiorstwa i każdy magazyn ustalają okresowo poziom zapasów dla poszczególnych pozycji asortymentowych. Informacje o sprzedaży przekazywane są codziennie lub najrzadziej co tydzień poprzez poszczególne ogniwa sieci dystrybucji do magazynu wyrobów gotowych. Dzięki temu zarówno fabryka, jak i poszczególne ogniwa w sieci dystrybucji, mogą planować i reagować na podstawie rzeczywistego popytu na rynku, a nie zamówień wynikających ze stosowanego systemu uzupełniania zapasów. Podstawowy poziom zapasów, ustalany dla każdego ogniwa łańcucha dystrybucji, jest równy przeciętnemu popytowi w cyklu dostawy powiększony o zapas zabezpieczający, który jest liczony według klasycznych metod.

Przykład metody *pull*- „system ssący”

W magazynie znajduje się zapas 1000 sztuk towaru. Wysłano 300 sztuk, a wyprodukowano kolejne 300 sztuk, aby uzupełnić zapas. Przy produkcji tych 300 sztuk zużyto materiały, które również trzeba uzupełnić. W ten sposób tworzy się łańcuch dostaw oparty na systemie *pull*, poczynając od klienta, przez montaż, podmontaż, kończąc na dostawcy komponentów.

Przykład metody *push*- „system tłoczący”

W każdym procesie są realizowane określone zadania, a po ich wykonaniu system „wypycha” wyroby do kolejnego procesu.

Zadania

1. Funkcja całkowitych kosztów działania systemu

wzór

$$C = K_1 * C_1 + K_2 * C_2 + K_3 * C_3$$

gdzie:

C - całkowity koszt związany z zapasami,
C₁ - koszt nadmiaru zapasu,
C₂ - koszt braku zapasu,
C₃ - koszt inicjacji procesu zamawiania (np. koszty złożenia zamówienia, koszty planowania),
K_i - współczynniki zależne od innych parametrów uwzględnianych w danym modelu np. od wielkości partii.

- a. Oblicz całkowity koszt związany z zapasami, jeżeli masz następujące dane:

K₁ – 0,6,
C₁ – 3000 zł,
K₂ – 0,8,
C₂ – 6000 zł,
C₃ – 0,4,
K₃ – 200 zł.

- b. Oblicz całkowity koszt związany z zapasami, jeżeli masz następujące dane:

K₁ – 0,4,
C₁ – 2000 zł,
K₂ – 0,6,
C₂ – 8000 zł,
C₃ – 0,2,
K₃ – 300 zł.

2. Wyjaśnij istotę metod *pull* i *push*.
3. Zastanów się, czy znasz firmę stosującą koncepcję *Just-in-time*. Na przykładzie postaraj się przeprowadzić analizę założeń koncepcji.