

Università di Roma Tor Vergata

Appunti di Operations Management 1

Corso 2017/18

Premessa:

Questa raccolta di appunti non nasce con l'intento di sostituire le dispense o le lezioni del docente, né tanto meno di dare insegnamenti.

Vuole piuttosto offrire delle domande su cui riflettere.

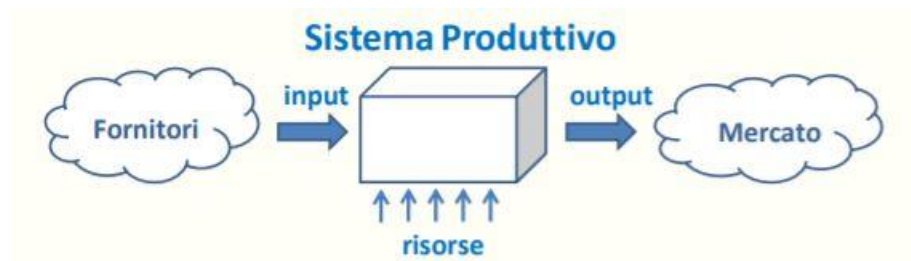
Il suo scopo non è pertanto quello di fornire risposte, ma quello di invitare il lettore nel porsi le giuste domande per poter apprendere nella corretta maniera i concetti del corso.

Indice

INTRODUZIONE	3
PERFORMANCE	7
COSTI	8
TEMPI	12
QUALITA'	16
SPAZIO	16
LEAD TIME DELLA DOMANDA E DEL PRODUTTORE	18
PROCESSI PRODUTTIVI	25
OEE	29
TASSO DI QUALITA'	36
DISPONIBILITA' A GUASTO	47
EFFICIENZA DELLE PRESTAZIONI	51
IL PROBLEMA DEL SETUP	52
LOTTO ECONOMICO	52
CASO MONOPRODOTTO, CONSUMO DISGIUNTO DALLA PRODUZIONE	53
CASO PLURIPRODOTTO, CONSUMO DISGIUNTO DALLA PRODUZIONE	58
CASO GENERICO, CONSUMO CONGIUNTO DALLA PRODUZIONE	60
GESTIONE DELLE SCORTE - Economic Order Quantity	63
TECNICHE SMED (Single Minute Exchange of Dies)	70
TEORIA DEI TEMPI E METODI (Method Time Measurements)	74

INTRODUZIONE

Il sistema produttivo è una scatola in cui giocano un ruolo fondamentale fornitori, mercato e terminologie (il gergo) che risulta fondamentale per la comunicazione e le considerazioni in merito.

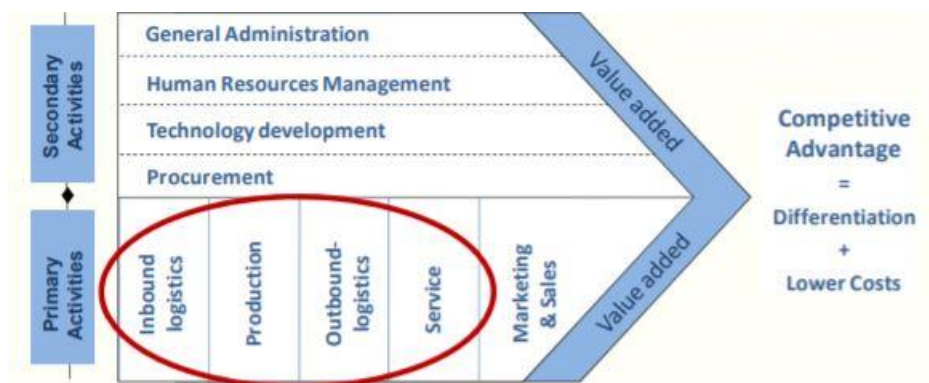


Si immagini di collocare l'azienda produttrice di beni all'interno di una filiera produttiva, dove sia a monte che a valle vi sono altri attori della filiera. A seconda della posizione occupata dall'azienda all'interno di tale filiera si distinguono due tipi di aziende:

- **B2B** (business to business) aziende che vendono ad altre aziende. Generalmente le aziende vendono i loro prodotti a poche aziende clienti, ma con quote di domanda rilevanti;
- **B2C** (business to consumer) aziende che vendono direttamente ai consumatori. La domanda, nel caso di un bene di consumo è formata pertanto da tante unità indipendenti.

Tendenzialmente si è abituati a pensare al B2C, ma in ottica Operations si è più interessati a capire le dinamiche del B2B, in cui l'approccio alla domanda è radicalmente diverso.

La **catena del valore** descritto da Porter permette di suddividere l'organizzazione interna aziendale in un numero finito di attività, precisamente 5 primarie e 4 di supporto. ciascuna di queste attività indispensabili per la creazione del valore per l'azienda.



Attività primarie:

- **Logistica in ingresso:** comprende tutte quelle attività di gestione dei flussi di beni materiali verso l'interno dell'organizzazione: flussi che alimentano le attività operative.

- **Attività operative:** attività di produzione di beni e/o servizi.
- **Logistica in uscita:** comprende quelle attività di gestione dei flussi di beni materiali verso l'esterno dell'organizzazione: flussi che portano sul mercato i risultati delle attività operative.
- **Marketing e vendite:** attività di promozione del prodotto o servizio nei mercati e gestione del processo di vendita.
- **Assistenza al cliente e servizi:** tutte quelle attività post-vendita che sono di supporto al cliente (ad es. l'assistenza tecnica).

Attività di supporto:

I processi di supporto sono quelli che non contribuiscono direttamente alla creazione dell'output ma che sono necessari perché quest'ultimo sia prodotto e sono:

- **Approvvigionamenti (Procurement):** l'insieme di tutte quelle attività preposte all'acquisto delle risorse necessarie alla produzione dell'output ed al funzionamento dell'organizzazione.
- **Gestione delle risorse umane:** ricerca, selezione, assunzione, addestramento, formazione, aggiornamento, sviluppo, mobilità, retribuzione, sistemi premianti, negoziazione sindacale e contrattuale, ecc.
- **Sviluppo delle tecnologie:** tutte quelle attività finalizzate al miglioramento del prodotto e dei processi. Queste attività vengono in genere identificate con il processo R&D (Research and Development).
- **Attività infrastrutturali:** tutte le altre attività quali pianificazione, contabilità finanziaria, organizzazione, informatica, affari legali, direzione generale, ecc.

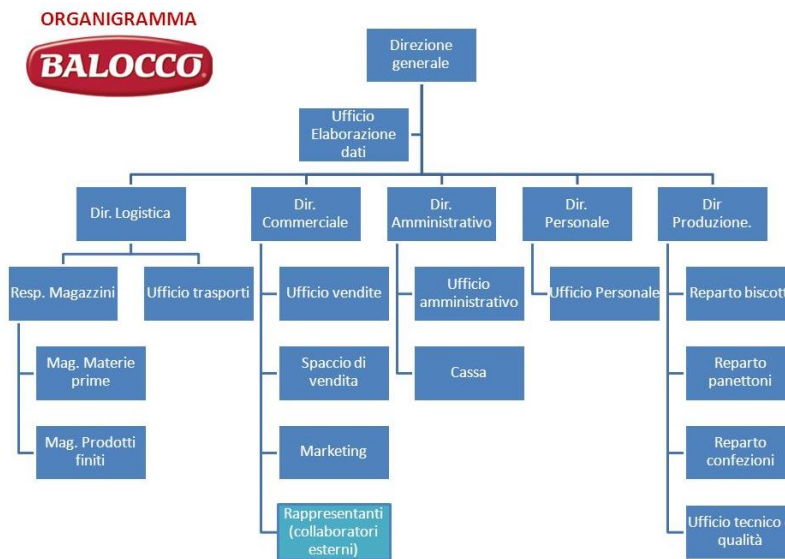
Tra le attività primarie che non sono di competenza dell'operation manager è importante definire la differenza che esiste tra la funzione marketing e quella commerciale delle vendite.

- **Il marketing** è il ponte tra l'azienda ed il cliente per tutto quello che non è commerciale, in poche parole il marketing deve capire il cliente e farsi capire dal cliente (comunicazione, pubblicità, strategie di mercato, analisi dei consumatori, dare agli operation degli input per cambiare i prodotti in funzione del cliente...). Quindi è relazione bidirezionale tra azienda e cliente.
- **La funzione del commerciale** è quella che entra in gioco nel momento in cui il marketing ha suscitato un qualcosa nei confronti di un cliente e decide di firmare un contratto, si occupa quindi di vendita. Il marketing non fa fatturato, ma lo fa fare al commerciale.

Differenza fondamentale → Marketing non fa fatturato, non essendo il suo scopo quello di vendere.

Per capire un sistema produttivo e le figure che lo compongono bisogna conoscerne l'organigramma. L'organigramma è quel grafico della struttura articolata dei vari organi e delle fasi di svolgimento delle operazioni di un'azienda, di un'amministrazione o di un ufficio.

Dall'organigramma si possono notare relazioni verticali fra elementi, anche dette **funzioni di linea** ed esistono laddove si produce del valore, mentre si hanno relazioni orizzontali ovvero le **funzioni di staff** e sono a supporto di quelle di linea.



In un Operations forte tutte le attività fanno riferimento ad un unico capo, appunto un Operations Manager. Un Operations risulta debole se si riducono tutte le attività fino alla sola produzione, perché in questo caso tutte le altre funzioni (Engineering, Health&Safety, ecc.) sarebbero di suo pari grado.

La “funzione Engineering” dentro una azienda è la funzione legata alla progettazione (prodotti e macchinari in primis).

In generale è possibile individuare l’engineering di prodotto e quella di processo la prima non può avere vincoli con il mondo Operations, quella di processo invece vi si può facilmente legare.

È una funzione da non confondere né con la produzione, né con la manutenzione (per ulteriori approfondimenti si suggeriscono le slide sul DFMA).

Volendo classificare le industrie in funzione del loro sistema produttivo, si distinguono due categorie:

L’industria di processo è un’impresa in cui si ha una continua trasformazione, tendenzialmente gli impianti vanno da soli perché caratterizzati da una produzione continua. Ciò vuol dire che non è possibile fermare la produzione in un particolare stadio una volta avviato, il processo non può essere arrestato (si pensi ad un processo chimico come la fermentazione), pertanto non può essere interrotto durante una sua fase intermedia, e quindi non può essere neanche stoccato. Un prodotto di un’industria di processo può essere riconosciuto perché non può essere scomposto a ritroso nei suoi componenti originari, dal momento che non sono più distinguibili o perché hanno cambiato natura. La birra è un tipico prodotto di un’industria di processo. In questi tipi di impianti i processi sono forzati, per via dei vincoli temporali o fisici e tecnologici per la trasformazione della materia, motivo per cui il principale compito dell’operation manager è quello di occuparsi di sicurezza e di manutenzione, avendo poca libertà di modifica delle dinamiche di processo. Generalmente un’industria di processo realizza infatti alti volumi con pochissima varietà di prodotto finito, molto spesso si produce un solo prodotto, motivo per cui il ruolo dell’operation manager è limitato alla sicurezza e alla manutenzione.

Al contrario ***nell'industria manifatturiera*** si realizzano volumi inferiori di ciascun prodotto finito, ma con un maggior mix produttivo. Caratteristica principale dell'industria manifatturiera è la possibilità di individuare uno stadio intermedio di lavorazione dei prodotti, da qui la possibilità per l'operation manager di decidere di **stoccare** semilavorati anche in una fase intermedia del processo produttivo, scelta impraticabile nel caso invece dell'industria di processo.

Tale caratteristica consente inoltre di soddisfare le domande di ciascun prodotto finito, tramite la produzione a lotti, la cui dimensione viene determinata dall'Operation Manager in funzione della tipologia di domanda da soddisfare per il singolo prodotto (quantità stabile o oscillante) e del tempo necessario per realizzarlo (Production Time, da confrontare con il Delivery Time), nonché dei costi di realizzazione del lotto stesso.

Avrebbe senso pertanto produrre a lotti in un'industria di processo?

Non molto, anche se riuscissi a stoccare prodotti intermedi costerebbe moltissimo tenerli fermi. Supponiamo che voglia stoccare del metallo fuso appena uscito da un forno e volendo stoccarlo devo tenerlo a quella temperatura, altrimenti si raffredderebbe, quindi per poterlo stoccare in quello stato dovrei impiegare enormi costi per poterlo conservare.

Avrebbe senso produrre per lotti anche nel caso di realizzazione di un unico prodotto?

In realtà potrebbe averla, per i seguenti motivi:

- Nel caso in cui ho una domanda più bassa di quanto avessi previsto, pertanto mi conviene sospendere la produzione, in quanto ho prodotto il necessario per soddisfare la domanda.
- Altro motivo potrebbe risiedere nella natura del processo: supponiamo che la macchina richieda una pulizia straordinaria ogni tot pezzi lavorati, a questo punto è necessario dover lavorare per lotti per effettuare il setup, anche se la produzione è monoprodotto.

Nell'industria manifatturiera invece il lotto è impiegato per permettere durante un setup su una macchina, di continuare a far lavorare quelle a valle, avendo materia stoccata.

Come si vedrà di seguito la dimensione del lotto può variare in funzione della strategia dell'operation manager da realizzare: lotti grandi permettono un vantaggio in termini di efficienza, dal momento che permette di spalmare i costi di setup su un volume maggiori di prodotti finiti. Maggiori sono il numero di setup maggiore il tempo in cui la macchina è rimasta senza produrre e quindi maggiore il tempo per poter rientrare dell'investimento dell'asset. Al contrario lotti piccoli consentono una maggior flessibilità del mix produttivo, con conseguente riduzione dell'inventario.

In generale, le produzioni per processo sono caratterizzate da un ciclo tecnologico **ben definito e vincolante** (si parla di produzioni a ciclo tecnologico obbligato). Le produzioni manifatturiere sono invece caratterizzate da una **grande varietà** di cicli tecnologici, che si traducono in numerose varianti di prodotti finiti (si parla di produzioni a ciclo tecnologico non obbligato).

La gran parte delle aziende hanno tuttavia un processo produttivo misto, ovvero hanno al loro interno una prima parte di processo (produzione della materia) e poi una seconda parte manifatturiera (ad esempio il packaging). La produzione della birra è anch'esso un processo misto.

12/03/2018

PERFORMANCE

Cos'è la performance e quali sono gli indicatori di performance?

Se nelle startup non esistono problemi di Operations, questi iniziano a sorgere quando la domanda diventa “importante” e non si ha a che fare con un sistema capace di misurarsi.

L'esame all'università lo giudico dal voto sul libretto. La misura indirizza e deve indirizzare il comportamento, e deve farlo in maniera credibile. La media risultante dai voti sul libretto mi dice se sono allineato o meno rispetto ad un obiettivo fissato relativo al voto di laurea.

Kaplan ha inventato il concetto di performance, e da lì sono nati i **KPI** (Key Performance Indicator). Attenzione, i KPI devono essere pochi (“chiave”) ed essere indicativi del processo (spesso vengono abusati dalle aziende), devono perciò: dire qualcosa; essere collegati all'obiettivo finale; e indicare e indirizzare il verso di “marcia” dell'azienda.

Essendo misure devono misurare la performance, ovvero tutto il quadro complessivo, non devono ridursi al solo risultato finale espresso in cifra, spesso riduciamo tutto ad un semplice numero, perché è più semplice.

Se un ufficio di formazione dice che ha fatto 5000 ore di formazione è indicativo o no?

Da solo dice poco, diverso sarebbe se andassi ad indagare l'efficacia della performance, ovvero indagare su quanto sono preparati i dipendenti.

La misura è uno strumento potente, pertanto non dobbiamo misurare a caso, ma là dove vogliamo capire e concentrare i nostri sforzi per un miglioramento (perseguire uno scopo finale).

Il CEO ad esempio è valutato sull'**EBIT** (risultato netto prima dell'aggiunta di tasse ed interessi) oppure **EBITDA** (prima ancora di ammortamenti e deprezzamenti, il cosiddetto risultato operativo lordo).

Se un CEO vuole migliorare l'EBIT allora “divide” l'obiettivo tra le varie funzioni (acquisti, produzione, personale, ecc.). Da qui il responsabile acquisti se dovrà scegliere tra due prodotti, uno migliore e più costoso e l'altro di qualità inferiore, ma più economico, sceglierà il secondo se più vicino ai suoi obiettivi (ridurre il costo degli acquisti) anche a costo di dare problemi in un altro contesto (ininfluente però sull'EBIT).

L'**MBO** (Management By Objective) non è così logica come sembra, c'è il rischio che talvolta la misura radicalizzi gli obiettivi.

Le misure principali dell'operato del lavoro sono: **COSTI, TEMPO, QUALITÀ, SPAZIO** (qualche volta al posto di qualità troviamo delivery o service).

COSTI

Contabilità

La contabilità generale è quella legata alla redazione del bilancio, che all'OM non interessa.

Nel bilancio ci sono dati a consuntivo che sono poco indicativi (il fatto che io spenda "x" per il personale, o per l'energia, ecc non mi dice nulla). Ricordiamo che giocando col bilancio si introduce il concetto di "nero" [*Come si crea il nero?* Società off-shore, gestite in modo da far circolare il guadagno dove la politica di tassazione è minore (quindi quelle off-shore), giacenze vendute a prezzi più bassi nell'anno successivo, specialmente se si è in utile si pagheranno meno tasse. Mentre se si è in rosso e non si è quotati in borsa si possono vendere le giacenze a prezzi più alti gli anni successivi per avere un pareggio di bilancio]. Quindi dal bilancio non si vedono performance.

Ai fini del OM interessa la contabilità analitica: ovvero parliamo di costi e ricavi delle Operations, un documento non pubblico, ma utilissimo all'uso interno. Questa è orientata a ciò che si vuole misurare e si costruisce a seconda dell'esigenza visto che si utilizza come indicatore di performance. Non è così banale da realizzare: Sono facili trovare le informazioni per i costi diretti, difficili invece e costose trovare informazioni sui costi indiretti.

Per esempio, un conto è ripartire il costo di una macchina che lavora su un prodotto, diverso è ripartire il costo della mensa, oppure di personale che si occupa di supervisione di più prodotti.

Da qui il concetto di **Centro Di Costo**, sono centri di calcolo tradizionali che distinguono i costi diretti (ovvero i costi direttamente sostenuti dalla produzione di un prodotto) e costi indiretti (ovvero i costi che contribuiscono al costo del prodotto ma per natura imprecisi ed è molto costoso reperire tali informazioni) spiegato col seguente esempio:

Se una linea accoglie il 2% dei lavoratori, allora assegnerò il 2% della mensa a quel prodotto.

Non è detto che sia totalmente giusto, però un modo per ribaltare tali costi si dovrà pur trovare.

Il concetto di costi indiretti risulta fondamentale nella contabilità analitica e nel calcolo delle performance di un impianto. Aziende tendono a basare le loro strategie su prodotti che hanno il margine più elevato, essendo quelli che incrementano l'utile aziendale, se non si riesce a fare una stima più o meno sicura dei costi indiretti di prodotto si rischia di basare strategie di mercato su prodotti che realmente non hanno quel margine, causando conseguenze a volte anche catastrofiche per la sopravvivenza nei mercati.

Costi fissi e costi variabili

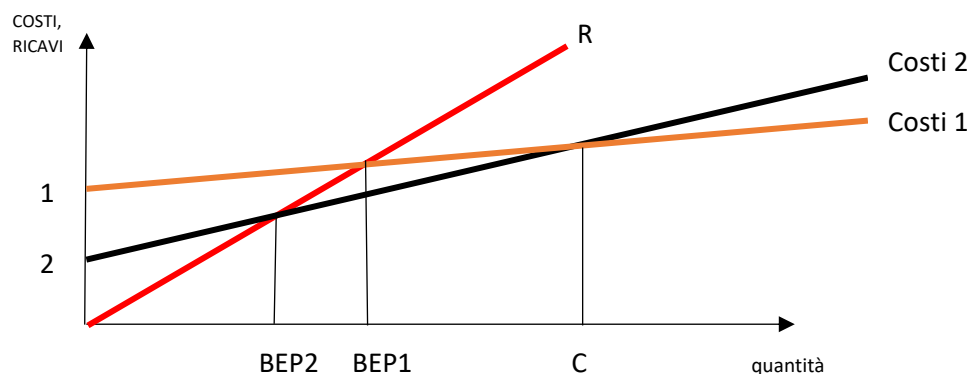
Costo fisso è un costo indipendente dalla quantità di prodotto che vado a realizzare.

Costo variabile è un costo valutabile e in funzione della quantità di prodotto da produrre.

Per esempio, la materia prima è un costo variabile, il costo d'impianto è invece fisso.

Fisso e variabile si distinguono sul costo totale, non sul costo unitario (il costo variabile è un costo unitario costante).

In funzione della strategia, adoperata da un'azienda, si potrebbero preferire gli uni o gli altri, a parità di costi totali. Ultimamente si preferiscono i variabili, per ridurre la possibilità di rischio.



Nell'esempio in figura l'azienda 2 ha investito meno (meno CF), quindi recupererà prima gli investimenti fatti, non esponendosi e riducendo il rischio (trend dal 2008 ad oggi: Investire meno per recuperare prima, inducendo un calo drastico negli investimenti che solo ora nel 2017 con Industry 4.0 sembra riprendersi).

Fino al BEP2 (**Break-even Point**) l'azienda 2 è in perdita, tra il BEP2 e il BEP1 l'azienda 2 è in utile la 1 in perdita, dopo il BEP1 entrambe in utile.

Dal punto C in poi (dove sono già entrambe in utile) l'azienda 1 consegue utili maggiori, perché ha meno costi variabili.

Quale dei due è migliore? Come sempre non si può rispondere in maniera oggettiva, "dipende", poiché gli investimenti sono una diretta conseguenza delle strategie di impresa.

Potendo avere una previsione del mercato in cui si ha una visione ottimistica (grandi volumi) risulta favorevole utilizzare strategie di investimenti come l'azienda 1, mentre se si ha una visione pessimistica o comunque conservatrice si sceglierà l'opzione adottata dall'azienda 2: si investe meno, si recupera, per poi "galleggiare".

È fondamentale dunque avere una giusta ripartizione tra costi fissi e variabili, in quanto la loro interazione determina la flessibilità dell'azienda alla variazione dei volumi di domanda.

Per esempio, i costi del personale possono essere trasformati in costi variabili se ho una manodopera multifunzionale. Infatti, se posso far lavorare sempre gli operai su uno dei prodotti i costi legati ad esse possono essere ripartiti con la produzione. Questo è un vantaggio in quanto mi permette di diminuire i costi fissi e di raggiungere prima il BEP, incentivando le imprese ad assumere.

Tale costo variabile viene quantificato con il **FTE** (Full Time Equivalent) cioè il costo di un operatore a tempo pieno. Se ho un sistema flessibile posso usare l'FTE in modo tale da sapere quanti operai attribuire alla macchina o alla linea, perché essi sono sempre a lavoro su qualcosa.

Produttività

La produttività è quel rapporto che mi indica quanto viene sfruttata bene una risorsa, tendenzialmente espresso come; fatturato per addetto. Le leve per aumentare questo indicatore sono l'aumento del fatturato, la diminuzione degli addetti, oppure la riduzione del costo del personale (spesso l'indicatore è fatturato/costo personale).

Per utilizzare una strategia Win-Win (entrambe le parti vincono) non si deve diminuire il costo del personale, ma consentir loro di fare di più, questo sarà possibile solo tramite investimenti. *Confindustria ha utilizzato una tipologia di contratto firmato dalle parti sociali in cui si prevede uno stipendio legato alla produttività ma legato anche a investimenti da parte della proprietà.*

Si può a questo punto differenziare l'**aspetto economico** da quello **finanziario**.

Il primo è legato ai flussi di denaro in funzione delle risorse utilizzate (Voci Costi e Ricavi nei bilanci, non sono soldi che per forza sono “usciti dalla cassa”), l'altro è invece l'indice di dove sono i soldi, (dove c'è per l'appunto un'operazione finanziaria). Per gli asset [qualsiasi bene di proprietà di un'azienda (macchinari, merci, ecc.), che possa essere monetizzato e quindi usato per il pagamento di debiti] utilizzo l'ammortamento: può essere visto come un recupero della svalutazione del bene, come un costo deducibile dagli utili, come un modo per tenere costante il valore dei cespiti.

I costi degli asset rimangono sempre fra i costi fissi.

In Industry 4.0, il decreto Calenda ha stimolato gli investimenti tramite un piano di ammortamento agevolato per le aziende con misure come il super-iper ammortamento, ossia imputando nel periodo d'imposta, quote di ammortamento e di canoni di locazione più elevati. Riuscendo ad incrementare la crescita del paese.

A titolo d'esempio, il costo del personale è un costo fisso o variabile?

Sembrerebbe fisso: a fine mese pago sempre la stessa cifra. Se però riesco ad avere una manodopera giovane, flessibile e dinamica, posso spostare una risorsa lì dove mi occorre, in maniera produttiva.

In questo modo, nell'ottica di contabilità analitica, il personale diventerà un costo variabile.

Renzi con il “Jobs Act”(2014-2015) ha cercato di rendere il costo della manodopera, costo puramente fisso dato la rigidità del lavoro come costo variabile cercando di renderlo flessibile abrogando ogni contratto a progetto, associazione in partecipazione e il contratto di lavoro ripartito e istituendo un nuovo tipo di contratto a tutele crescenti, contratto che da al neoassunto, fin dal primo giorno di lavoro, tutte le tutele previste nei vecchi contratti e in caso di licenziamento il dipendente avrà un indennizzo pari all'anzianità del servizio. Le aziende che assumono con questo tipo di contratto a forma indeterminata potranno avvalersi di sgravi contributivi.

Come quantifico il costo variabile del lavoro?

FTE (Full Time Equivalent): costo di un operatore a tempo pieno, e non numero di operai.

Quindi 3 FTE significa tre operai a tempo pieno oppure sei operai a tempo dimezzato mentre 2,5 FTE significa che qualcuno è utilizzato per esempio un po' su una linea e un po' su un'altra.

Le scorte

Per scorte possiamo riferirci a materie prime, semilavorati o prodotti finiti, ma anche alla capacità produttiva (ad esempio una macchina di riserva).

Le funzioni principali sono quelle di riserva di capacità e disaccoppiamento delle dinamiche (macchine lasche, ponendo dei buffer fra di esse).

Esse rappresentano materiali o risorse che vengono pagate con il capitale circolante, ottenuto dalle banche o dalla proprietà. Le scorte non sono gratuite anzi, l'azionista percepisce il ritorno sul valore e la banca ha un interesse sul dividendo.

Con il WACC (Weighted Average Cost of Capital) misuriamo il costo medio pesato, esso rappresenta il costo medio del capitale che l'impresa paga a tutti i suoi investitori, azionisti e creditori.

Se ci sono tante scorte di alto valore, ho per quanto detto finora, generato un costo tanto più alto quanto è di valore la scorta. Dovendo restituire più di quanto ho ricevuto (senza considerare poi il rischio di obsolescenza delle scorte) si ha il rovescio della medaglia che le caratterizza.

Quindi se da un lato le scorte sono un valore aggiunto (riserva e disaccoppiamento) dall'altro generano inevitabilmente dei costi.

Generalmente il costo delle scorte è espresso come:

$$C = V * i$$

Dove:

V = Valore del bene stoccato che sia materia prima, semilavorato o prodotto finito.

i = Percentuale del valore del bene imputabile al costo di stoccaggio, tiene in considerazione i fattori esposti in precedenza (rischio di invenduto, obsolescenza, costi di stoccaggio, costo del capitale circolante, assicurazione,).

Costi differenziali e affondati

È importantissimo sapere su cosa possiamo agire; ci sono elementi (costi) di un problema che sono indipendenti dalle scelte e dalle soluzioni trovate (costi affondati).

Bisogna individuare i **costi differenziali**, che possono cambiare in funzione delle nostre scelte.

Questa differenza può dipendere dal tempo, dalla posizione gerarchica del decisore, ecc.

13/03/2018

TEMPI

Cosa vuol dire misurare le performance rispetto ai tempi e quali tempi sono importanti?

D-Time	Quanto tempo il cliente è disposto ad aspettare per usufruire di un bene o di un servizio.
P-Time	Tempo in cui si riesce a reagire alla domanda, formato da tutti i tempi che caratterizzano la produzione dall'istante in cui è stata acquisita la domanda (Approvvigionamento, produzione e consegna/distribuzione).
Lead Time	Quanto tempo prima devo cominciare una determinata azione affinché si rispetti una data prestabilita in tempo (data per una consegna, approvvigionamento, produzione etc...).
Time To Market	è il tempo che intercorre fra l'ideazione di un prodotto e la sua effettiva commercializzazione (momento in cui arriva al cliente)
Tempo Ciclo	(riferito alla linea intera o ad una singola macchina) è il tempo che intercorre fra l'ingresso e l'uscita del prodotto nella linea/macchina
Throughput Rate	È la velocità di emissione dei prodotti lavorati, valutabile come il reciproco del tempo ciclo, ovvero quante unità vengono prodotte nell'unità di tempo. Può essere regolata in funzione della domanda, tuttavia non può eccedere il valore della capacità produttiva della macchina: $\text{Throughput Rate} \leq \text{Capacità Produttiva}$
Takt Time	è il ritmo della produzione settato sulla domanda. Si tratta del tempo necessario a produrre un singolo componente o l'intero prodotto, noto anche come Ritmo delle Vendite. È un dato esogeno dell'impianto, viene dal mercato. Mi serve per adattare il ritmo di produzione alla domanda richiesta.
Throughput Time	è il tempo di attraversamento dell'intero impianto
Tempo Idle	è il tempo in cui una macchina funzionante non è utilizzata. È funzionante dunque non è guasta potrebbe lavorare, tuttavia non viene utilizzata. Questo vuol dire che né il tempo di guasto né il tempo di setup sono considerabili tempi di idle.

Lead time (o tempo di risposta)

È una famiglia di tempi: tempi di anticipo, quanto prima devo cominciare a progettare, ad approvvigionarmi, a produrre o a consegnare affinché io riesca a soddisfare la richiesta del cliente, il lead time può riferirsi quindi alla progettazione, approvvigionamento, produzione o consegna, ma si può riferire anche alla domanda, in quel caso lead time di domanda o D-Time che rappresenta quanto i clienti sono disposti ad aspettare per usufruire del bene o servizio.

Time to market

è il tempo che va dall'idea, all'affaccio sul mercato di un prodotto (poco interessante per un OM)

Tempo di evasione ordine, oppure il Tempo che il mercato mi concede per evadere l'ordine.

In quest'ottica è molto importante capire il contesto competitivo in cui operiamo.

Se si parla di concorrenza perfetta (B2C) nessuna azienda fa le regole, ma le subisce.

Se perciò il cliente ordina una macchina personalizzata e di norma ci vogliono tre mesi per realizzarla tutti i competitors più o meno ci impiegano lo stesso tempo, se arriva però un competitor che ci mette due settimane allora questo sbaraglierà il mercato.

Se invece il mercato è un oligopolio, le regole vengono stabilite dai giocatori, ognuno è abbastanza forte da dettare le regole.

Nel mercato concorrenziale vince chi consegue utile: è un gioco combattuto sulle curve di domanda.

Nell'oligopolio vince chi ha maggiore quota di mercato, l'utile è un aspetto successivo.

Vendere sottocosto per conquistare quote di mercato (vietato dai garanti della legittima concorrenza) è una forma di investimento, al pari della pubblicità.

Nell'oligopolio il P-Time è una grande leva competitiva, aspetto in cui è chiaramente molto influente l'Operations Manager.

Anche nel B2B ciò ha una notevole importanza: un'azienda una volta inviato un ordine al fornitore è vincolata da aspetti contrattuali, se il fornitore ha tempi di consegna di un mese si deve emanare l'ordine un mese prima, di conseguenza aumentano i rischi per il richiedente, ovvero l'azienda, che potrebbe accusare le condizioni mutevoli del mercato in cui opera (variazione della domanda, ripensamento da parte del cliente finale, problema tecnico, ecc). Avere un fornitore che riduce al minimo il Lead Time di approvvigionamento (ovvero il tempo che dovrò aspettare quando emetto un ordine al fornitore) è un punto di forza, significano meno rischi, meno vincoli, probabilità inferiore di imprevisti generici. Poter ordinare all'ultimo minuto ad un fornitore (a parità di altre condizioni) è un vantaggio, in termini finanziari e di rischi, potendo così venire più incontro alle esigenze dei miei clienti.

Si può addirittura rispondere in PULL: il meccanismo di produzione si attiva solo quando arriva la domanda, significa che il P-Time è minore o uguale al Lead Time della domanda (D-Time).

Oggi FIAT manda gli ordini ai fornitori 40 minuti prima della consegna, ovviamente è essenziale avere i fornitori vicino (JIT).

Ordinare all'ultimo minuto flessibilizza e consente di ridurre le scorte (in 40 minuti ci saranno meno imprevisti che in una settimana).

Tempo Ciclo

Con tempo ciclo si intende il tempo che passa tra l'entrata in una macchina di un oggetto e la sua uscita, può essere riferito anche ad una linea, ovvero il tempo ciclo della linea che sarà data dalla sua stazione collo di bottiglia.

Throughput Rate (Capacità Produttiva)

Tempo con cui escono i prodotti dalla linea (u/sec, u/min, ecc) è quindi una capacità produttiva.

Attenzione il Throughput Rate viene deciso a discrezione del Operations Manager se l'impianto riesce a rispettare ampiamente la domanda, se invece non si riesce a rispettarla il Throughput Rate sarà posto alla capacità produttiva dell'impianto, ovvero la velocità massima che può raggiungere l'impianto, (Upper Bound), cioè a quanto teoricamente potrebbe andare la macchina/impianto.

I "dati di targa" per questo motivo hanno poco senso, le macchine e le loro caratteristiche possono essere modificate nel contesto in cui vengono poste (impianto).

Se la domanda non è altissima è inutile sfruttare la macchina al massimo delle sue potenzialità, anche perché maggiore è lo sforzo a cui è sottoposta la macchina e maggiori sono i numeri di possibili guasti.

Ma il tempo di attraversamento è il tempo dato dalla somma dei tempi ciclo delle singole stazioni?

Non sempre, i tempi ciclo non tengono conto di eventuali buffer posti tra le macchine per disaccoppiare le dinamiche, quindi bisogna prima definire il tipo di connessione in cui si vuole operare.

Definiamo:

- Connessione **rigida**: no buffer;
- Connessione **lasca**: con buffer.

Per esempio:



A regime un pezzo potrebbe avere un tempo di attraversamento di 59 secondi così composto:

4 secondi per la prima lavorazione, 50 secondi nel buffer, 5 secondi per la seconda lavorazione.

Attenzione, avere un buonissimo OEE non è sempre indicativo, potrebbe essere "drogato" dalle scorte, più la connessione è rigida e più i macchinari sono collegati, quindi se risulta esserci un guasto in un macchinario avrà ripercussioni sull'intera linea.

Takt time

Ritmo a cui deve lavorare una linea produttiva. Sarà dettato dalla domanda e lo calcoleremo (più o meno) come il reciproco della domanda.

Quindi dovrà essere che:

$$\frac{1}{\text{Throughput Rate}} \leq \text{Takt Time}$$

Per monitorare l'andamento della linea rispetto alla domanda il miglior indicatore è quello del ritmo (Takt Time) piuttosto che la verifica dell'esatto numero di pezzi.

Usare il Takt time consente di gestire accelerazioni e rallentamenti, ma attenzione non deve diventare un'ossessione → *“basta che non diventa il rullo di tamburo per scandire il tempo agli schiavi”*.

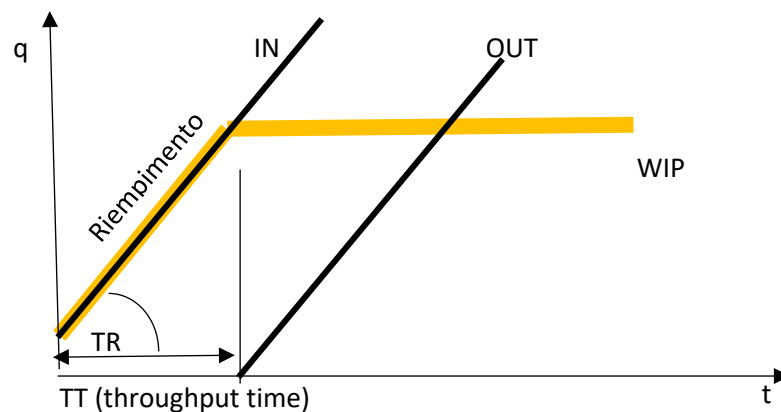
Throughput time (Tempo di Attraversamento)

Tempo di attraversamento dell'intero impianto, coincide con il tempo ciclo della linea più le attese ai buffer. Tale tempo può avere due interpretazioni diverse in base al tipo di linea analizzato:

1. **Linee asincrone**, il materiale viene fatto avanzare da una stazione alla successiva non appena è disponibile all'output di una fase di lavorazione. Il tempo di attraversamento delle stazioni coincide quindi, con il tempo ciclo delle macchine che lo compongono. Nel caso in esame, risultando tale tempo coincidente con il tempo ciclo della stazione di lavoro generica TC si ha che il tempo di attraversamento della linea coincide, assumendo che non si abbiano delle attese dei semilavorati lungo la linea stessa, con la somma dei tempi ciclo delle stazioni di lavoro che la compongono. Questa condizione di assenza di code di semilavorati è rispettata solo se il pezzo non deve attendere di essere lavorato a monte delle stazioni più lente, è necessario alimentare la linea ogni intervallo di tempo t maggiore o uguale al tempo ciclo della fase del processo più lenta. In caso contrario i pezzi tenderebbero ad accumularsi a monte della fase di lavoro più lenta e il tempo di attraversamento aumenterebbe progressivamente nel tempo fino ad un'eventuale interruzione della produzione delle fasi a monte.
2. **Linee sincrone**, invece, l'avanzamento di tutti i semilavorati presenti lungo la linea avviene contemporaneamente ad intervalli di tempo pari al tempo ciclo TC della linea (ad esempio mediante l'avanzamento di un nastro trasportatore). Di conseguenza, il tempo di attraversamento delle singole stazioni di lavoro diventa uguale per tutte e pari al tempo ciclo TC, provocando un aumento del tempo di attraversamento della linea.

Si ritiene inoltre opportuno precisare che tra le due differenti tipologie di linea non varia il tempo ciclo e quindi la capacità produttiva, la differenza principale è costituita dal fatto che la linea asincrona, grazie ad un tempo di attraversamento minore, minimizza la quantità e quindi il valore dei semilavorati presenti lungo la linea (work in progress WIP), mentre la linea asincrona consente di semplificare i sistemi di movimentazione.

Legge di Little:



Se dopo la fase di riempimento tanto entra quanto esce (“conservazione materia”) con un coefficiente angolare pari al TR, avrò che il Work In Progress sarà pari a:

$$WIP = TT \times TR$$

Il WIP critico è il livello di WIP necessario a far sì che il processo raggiunga il massimo valore di Throughput pur garantendo un tempo di attraversamento minimo.

Per diminuire il WIP non posso abbassare il TR (la capacità produttiva), anche se con un WIP alto ho minori tempi Idle per le stazioni non collo di bottiglia, al tempo stesso però con un WIP alto aumenterò la lunghezza delle code alle stazioni di lavoro quindi maggiore tempo di attraversamento medio a regime dei prodotti in lavorazione.

QUALITA'

Quanti pezzi difettosi fa ogni singola macchina, oppure pezzi buoni sul totale dei pezzi.

Yeld: % di pezzi buoni sul totale.

SPAZIO

Prima quando si progettavano gli impianti, si moltiplicava per 5 lo spazio utile (parametro che si impone sempre con più forza), in previsione di una domanda che sarebbe cresciuta.

Oggi non è più così, anzi le aziende vanno in crisi di spazio.

Riuscire a fare le stesse cose con meno spazio è il miglior compromesso.

Lo spazio si dice che “nasconde” i costi, ovvero: sprecarne toglie mobilità e questo rappresenta una spesa (non sempre esprimibile in maniera diretta).

La mentalità occidentale tende a dare importanza allo studio sul recupero dell'investimento per decretare sprechi di spazio (ma come già detto è difficile quantificarlo), ed in alternativa si preferisce non intervenire; diversamente la Lean Production (mentalità giapponese) esprime come gli sprechi di qualsiasi natura vadano eliminati, senza necessariamente dover fare un passaggio intermedio che li quantifichi.

Impianto

È un insieme di risorse produttive che devono rispondere ad una domanda, realizzando prodotti.

Bisogna perciò capire come ci si posiziona la capacità produttiva rispetto alla domanda: se maggiore, minore o uguale.

LEAD TIME DELLA DOMANDA E DEL PRODUTTORE

Come risponde l'impianto alla domanda?

Ci sono varie risposte a questa domanda, in base ai fattori tenuti, fin ora, in considerazione.

L'impianto è un insieme di risorse produttive che deve **rispondere** ad una domanda, realizzando un prodotto. La **domanda** è la richiesta di un certo prodotto dal mercato. La domanda dipende dal **mercato** in cui mi trovo. Per un'azienda con un unico impianto è come da definizione, ma per un impianto di una Corporate prima di tutto la domanda è riferita a tutto il gruppo. Dopo averla raccolta la Corporate suddividerà la domanda tra i vari impianti, ciascuno dei quali produrrà diversi prodotti (per ridurre il rischio). Tale domanda verrà quindi filtrata in base alla capacità produttiva disponibile dell'impianto. Quindi se un impianto riesce ad incrementare la sua CP può chiedere una domanda maggiore. In questo senso quindi la CP è la capacità di rispondere alla domanda.

Rispondere alla domanda in funzione della capacità produttiva:

Un impianto può lavorare in due ambiti rispetto alla domanda:

- **CP>D:** Solitamente non succede mai perché vuol dire avere più costi fissi e quindi più rischi. Può avvenire quando un mercato si evolve (ridotto). In questo caso la priorità non è l'OEE ma sarà sul ridurre i costi.
- **CP<D:** è la normalità in un impianto, ed in questo caso la priorità è l'OEE perché così aumento la CP e posso prendermi la domanda che avanza nel mercato. Nel caso di domanda stagionale posso fare scorte, posso fare **outsourcing** (chiedere CP a concorrenti), posso aumentare i turni di lavoro. In questo caso se sono in una linea vado a lavorare sull'OEE della macchina collo di bottiglia mentre sulle altre, siccome ho CP in eccesso, posso svolgere azioni per minimizzare i costi.

Le aziende possono essere classificate per il metodo di risposta alla domanda, producendo con continuità o intermittenza.

Rispondere alla domanda in funzione del sistema produttivo:

Come avviene quindi la produzione con continuità o ad intermittenza?

- **Produzione continua:** si produce sempre lo stesso prodotto, saranno quindi quegli impianti in cui ho una domanda maggiore della capacità produttiva. Tutto ciò comporta un rischio; queste tipologie di produzioni sono comuni nell'industria di processo, in cui si hanno investimenti elevati (dedico tutta una linea [macchinari, spazio ecc] ad un solo prodotto) legati ad una domanda più o meno certa, ma se così non fosse e la domanda risulta mutevole nel tempo il rischio potrebbe essere troppo alto.
- **Produzione per lotti (intermittenza):** quasi sempre si produce ad intermittenza: per lotti.

Nessun prodotto satura totalmente la capacità produttiva della mia linea/impianto (oppure sono io a volere che questo non accada) così da bilanciare il rischio di mercato producendo altri prodotti.

In inglese il lotto è definito **Batch**, ossia il numero di prodotti uguali tra loro che vengono lavorati insieme e con continuità.

Seguire un lotto da un punto di vista logistico-informatico è comodo mentre seguire ogni singolo prodotto nel suo iter realizzativo è davvero oneroso. Un lotto è facilmente rintracciabile (*si pensi alle esigenze in ordine ai cibi avariati da richiamare, rintracciando il lotto si risale ai punti in cui è stato distribuito*).

D'altro canto, passare da un lotto ad un altro implica inevitabilmente degli sprechi (Set-up o tempi di stop) che possono ridurre notevolmente la produzione della linea/impianto.

I problemi di Set-up possono essere scomposti in:

- cambio formato: produco lo stesso bene in formato diverso.
- cambio prodotto: Solitamente passare da un prodotto ad un altro richiede maggiori costi e tempi.

In questa seconda ipotesi il riattrezzaggio è molto più critico (lavaggi, igienizzazione, sprechi di prodotto, ecc.).

Lavorare per lotti ci consente di poter differenziare la nostra produzione al variare della domanda dei miei prodotti sfruttando la capacità produttiva totale dell'impianto (non è escludendo che mi metta a produrre per terzi → Finalità far lavorare il mio impianto)

Nel farmaceutico, per esempio, capita che un'azienda produca packaging per un marchio concorrente, ovviamente la concorrenza tra marchi è sulla molecola, non sulla scatola, quindi nulla di anomalo o eclatante.

Sempre in tema di farmaceutico, per essere certificato dall'AIFA (in Italia) devi rispondere a determinati requisiti, ad esempio in tema di modalità di setup (norme GMP).

Rispondere alla domanda in funzione del tempo a disposizione (D-Time):

Un'altra considerazione è legata al tempo. Quando si risponde alla domanda? E come?

Introduciamo i concetti di **PULL** e **PUSH** largamente utilizzati per discriminare la tipologia di produzione.

PULL: inizio la produzione quando ricevo la domanda e quando ho finito consegna (non comincio la produzione finché non ho una domanda).

Una produzione di questo tipo è la produzione per commessa. Una produzione in pull può concretizzarsi anche in un "accordo aperto": Fiat chiede i pezzi ai fornitori quando ne ha bisogno (dipende cioè dalle vendite) e i fornitori rispondono quando ricevono l'ordine.

Nell'ambito automotive i fornitori vengono scelti "per gara per ciascun modello". Nelle gare vengono invitati alcuni fornitori. Dal momento che un fornitore si aggiudica la vittoria quel vincitore è fornitore per il determinato componente del determinato modello, per l'intera vita del modello prodotto (al massimo c'è un fornitore di riserva).

Il pull abbate i costi di stoccaggio, minimizzando il rischio di invenduto e l'immobilizzo di capitale circolante netto. Altri vantaggi si avranno nella richiesta di prestiti che saranno caratterizzati da interessi più bassi (quando si chiede un prestito si mostra la "certezza" della commessa), avrò quindi in generale dei benefici finanziari.

PUSH: non si aspetta la domanda, si comincia a produrre per poi depositare i prodotti finiti in magazzino, ovviamente si cercherà di stimare la domanda per tarare il volume di produzione da realizzare, pur rispondendo al fabbisogno tramite le scorte di prodotto finito del magazzino (per evitare spese di magazzino), non seguendo quindi fedelmente l'andamento della domanda.

Si calibrerà la produzione infatti sulla domanda media dell'orizzonte temporale considerato, in questo modo si può produrre in maniera più efficiente lavorando sul lotto economico (quando le condizioni lo permettono), dal momento che le dinamiche tra la produzione e la domanda sono disaccoppiate dal magazzino.

Naturalmente ciascuno dei due metodi ha dei vantaggi e degli svantaggi.

Per esempio, con una produzione di tipo PUSH si può:

- produrre come e quando si vuole, senza essere schiavi della domanda;
- razionalizzare la produzione programmando nel migliore dei modi setup e riattrezzaggi, generando ottimi tempi di risposta (non c'è attesa) quindi un servizio migliore (nel Pull spesso accade che se arriva la richiesta di un solo pezzo devo riattrezzare tutto per fornirlo).
- è indispensabile quando il D-Time è pari a 0.

Gli autosaloni cercano di giocare sulla pronta consegna (o quasi pronta consegna, grazie ad analisi previsionali). Il Km 0 è l'inventuto di questo meccanismo: hanno acquistato dalla casa madre che dopo un complessivo di acquisti e immatricolazioni genera dei bonus. Se non riescono a venderlo realmente, propongono l'inventuto ai clienti con uno sconto per l'avvenuta immatricolazione.

Il mercato del Km 0 si regge su tali strategie di trade-off.

Il Lead Time della domanda ($\underline{LT}_D = D\text{-time}$), come già definito, è costituito dal mercato: se è pari a zero (significa che il cliente non è disposto ad aspettare i tempi di lavorazione) bisognerà necessariamente lavorare in PUSH, altrimenti si può ragionare, almeno in parte, in PULL.

Di fronte al D-Time devo valutare il Lead time del produttore ($\underline{LT}_P = P\text{-time}$), ossia il tempo che si impiega a reagire alla domanda. Ovviamente per lavorare esclusivamente in PULL, la condizione deve essere che $P\text{-Time} \leq D\text{-Time}$.

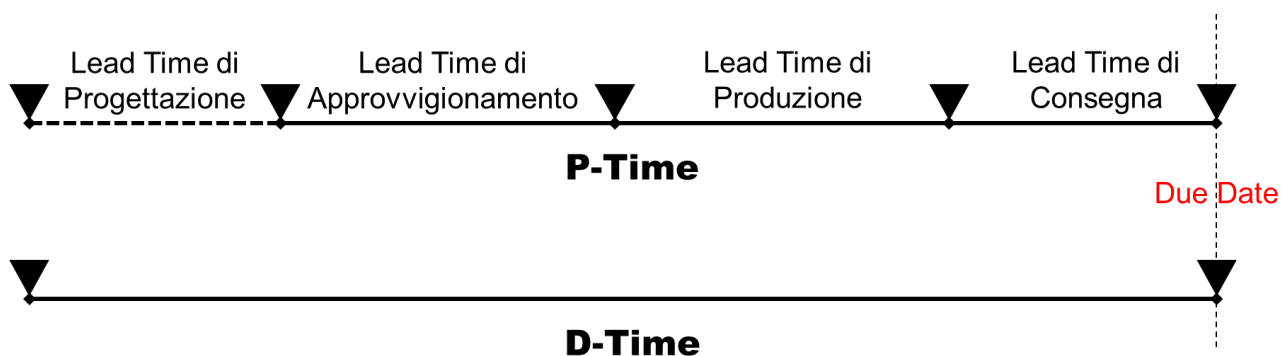


Figura 1. Composizione del Lead Time del produttore e della domanda

Analizziamo il P-Time:

Attenzione: non sono tempi di attraversamento.

Se il cliente mi concede un D-Time pari o maggiore allo schema riportato (Figura 1) allora posso lavorare in pull.

Ma se il D-Time è inferiore al P-Time (Figura 2: Casi 2,3 e 4) allora non si può aspettare la domanda per iniziare la produzione (produzione intesa come fase realizzativa del prodotto, ovvero progettazione- approvvigionamento- produzione- consegna/distribuzione).

Se si ha una configurazione del genere $P\text{-Time} > D\text{-Time}$ parte della catena di produzione dovrà essere svolto in PUSH.

Ma cosa significa progettare in PUSH? Significa fornire un catalogo di prodotti da cui si può far scegliere il cliente, che ciò implica (ancora una volta, come tipico dei sistemi PUSH) un potenziale rischio.

Se il D-Time fosse ancora minore non si dovrebbe solo progettare in PUSH, ma si dovrà anche pensare ad approvvigionamenti di tipo PUSH, facendo previsioni di domanda delle materie prime o semilavorati attraverso i piani di produzione.

Da qui si deduce che un fornitore che ha tempi di consegna minori è da preferire perché mi permette di lavorare in PULL.

Accorciando ancora di più il D-Time si arriva quindi a produrre in PUSH.

Si cerca di prevedere la domanda, compro le materie prime, produco e metto in magazzino, quando arriva la domanda consegnerò in pull dal magazzino.

Ci si può chiedere se vale la pena a questo punto consegnare anche in PUSH. Consegnare (distribuzione) in PUSH comporta notevoli rischi. Ipotizzando di lavorare in mercati globali, la consegna comporta costi elevati, se si produce come unico produttore un certo bene (vendendolo in tutto il mondo) utilizzando una strategia di consegna PUSH si dovrà far arrivare il prodotto in tutti i punti di vendita del mondo prima che arrivi la domanda, implicando un rischio di invenduto elevato. Quindi più elevato è il costo di trasporto migliore sarà adattare strategie di consegna in PULL (**Ricorda che le scorte si misurano in tempi di durata e non in quantità**).

Se il D-Time scende ancora è inevitabile che anche la consegna avvenga in PUSH.

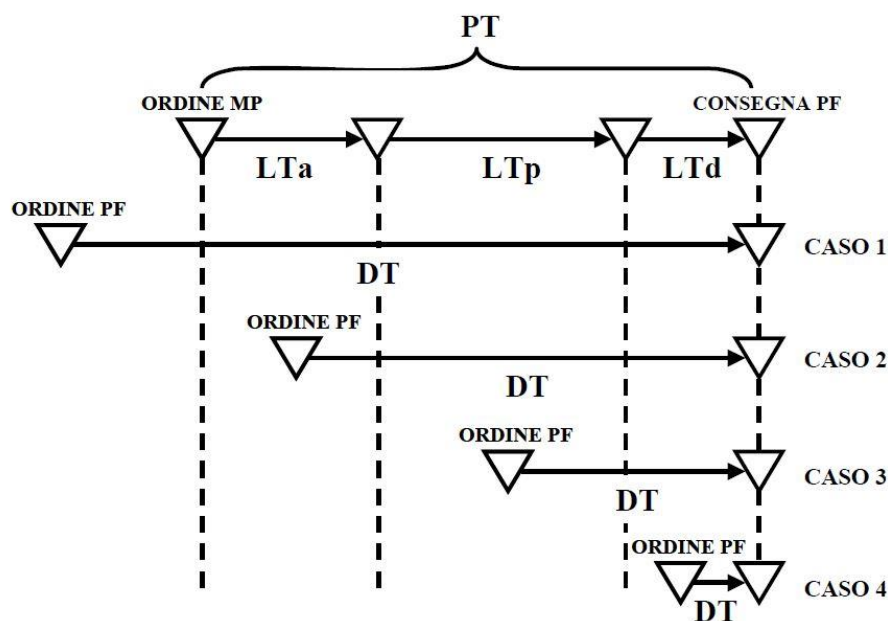


Figura 2. Casistiche generali

Pensiamo ad un esempio quotidiano: se devo cambiare il cofano della macchina dovrò aspettare che il meccanico lo ordini, se invece devo cambiare una lampadina molto probabilmente l'avranno subito disponibile (e se qualcuno non l'avesse io cambierei posto perché mi aspetto di trovarla).

Questo esempio lega costi, spazio di stoccaggio e mercato in cui si opera → se tutti operano in una certa maniera ecco che si plasma il D-Time del mercato delle lampadine e del cofano.

Ovviamente la presenza di un certo pezzo di ricambio nel magazzino regionale, piuttosto che in quello nazionale o addirittura europeo e mondiale dipenderà dalla rotazione del prodotto, principalmente quindi dalla frequenza con cui quel pezzo si rompe (tipo di strategia).

Le farmacie, invece, hanno delle scorte, ma se necessitano di un prodotto (domanda di un prodotto non in magazzino) sono in grado di fornirlo nella mezza giornata successiva. Anticamente le farmacie erano più fornite, con enormi costi di stoccaggio, ma qualora fosse mancato qualche prodotto si avrebbe avuto un tempo di consegna dai 2-3 giorni.

Tipicamente le società di distribuzione farmaceutiche gestiscono più punti vendita nello stesso territorio. Queste utilizzano magazzini di prodotti finiti in cui si aggregano le domande delle singole farmacie per quei prodotti che vendono poco (aggregando le domande aumenta la rotazione di un prodotto per il magazzino comune del territorio).

Posso rendere la consegna un po' più PUSH?

*Sì, grazie alla **Milk Run**: ogni giorno il fornitore fa lo stesso giro (calcolato di costo minimo), a prescindere dalle singole richieste, in questo modo distribuisce tutti gli ordini arrivati dalle farmacie entro una certa ora (ad esempio le 12), senza costi aggiuntivi alla farmacia che richiede il servizio pagando cioè un abbonamento anziché ogni singola corsa.*

Anche Zara e Amazon Prime funzionano con questo meccanismo. Magazzino in PUSH, consegna in PULL.

Nel punto di collegamento che segna la transizione tra le due logiche (PULL-PUSH) sarà necessario predisporre l'elemento di disaccoppiamento, quindi buffer di scorte o, più propriamente, un magazzino opportunamente dimensionato.

Analizziamo più nel dettaglio il caso 3 della Figura2. Caso in cui il D-Time mi concede parte della fase di produzione in PULL. Partendo dal presupposto che in generale la fase di produzione è costituita da più step sequenziali, si può pensare quindi di spezzare il processo produttivo costituendone una parte in PUSH e una in PULL (se lo trovo conveniente).

Mettere a scorta un semilavorato è sicuramente meglio che metterci un prodotto finito, la considerazione però che deve essere fatta va ad analizzare il tipo di lavorazioni svolte nella fase di PULL.

Tendenzialmente risulta conveniente tale tipo di logica se le lavorazioni ultime (quindi quelle svolte in PULL) risultano le operazioni di maggior valore (valore = materie prime + tutte le operazioni), quindi più le operazioni costose stanno in fondo più mi conviene spezzare e fare questo spaccettamento PUSH/PULL, spostando la cerniera verso la parte finale del processo di realizzazione del prodotto finito (Rischio è maggiore dove i costi sono maggiori).

Se un impianto produce tanti prodotti diversi l'ideale è che a monte del punto di disaccoppiamento il processo produttivo sia meno differenziato possibile, ovvero che esso sia posto nel punto in cui il numero di differenti articoli semilavorati sia minimo (o almeno minore dei differenti articoli di prodotti finiti) al fine di minimizzare la quantità di semilavorati da tenere in giacenza (Figura 3).

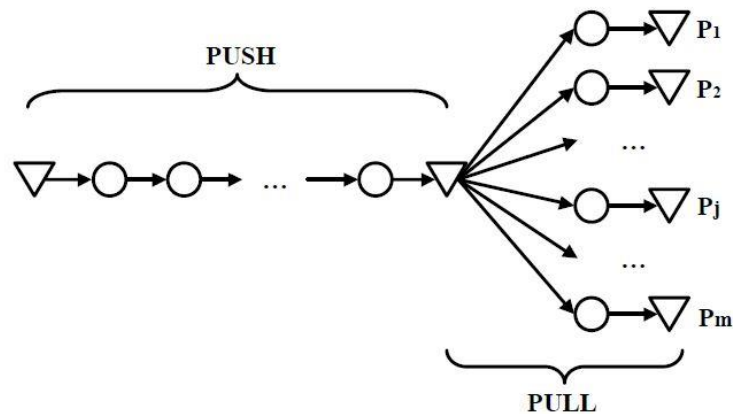


Figura 3. Meccanismo PUSH-PULL in un impianto di produzione multi-prodotto

Benetton ha usato questo meccanismo. La famiglia Benetton era costituita da 3 fratelli: Luciano, il più famoso e grandissimo uomo di marketing; suo fratello grande uomo industriale e sua sorella particolarmente dotata per le tecnologie chimiche.

Il contesto in cui l'azienda Benetton opera è il mercato del prêt-à-porter (pronto a portare) mercato fortemente influenzato da sfilate che venivano organizzate con tre mesi d'anticipo (e non con un anno come l'alta moda). Le aziende del settore generalmente aspettavano per iniziare a produrre, ma con così poco tempo nessuno riusciva a raggiungere grossi volumi (nessuno voleva lavorare in PUSH). Benetton è stato il primo a scoprire che tale tipo di mercato è fortemente influenzato dal tipo di colore che va di moda, e non tendenzialmente dal tipo di vestito (modello). Cerca quindi di utilizzare strategie di tintura dei capi dopo averli prodotti. Riescono a disaccoppiare le dinamiche lavorando parte in PUSH con capi non colorati (tanto di anno in anno cambiava la tendenza del colore e molto poco il modello) e dopo le sfilate li tingevano (Figura 4). In questo modo sono riusciti ad aprire negozi monomarca, generalmente difficili da aprire per l'incapacità di produrre grossi volumi. Utilizzando un forte piano d'attacco grazie a imponenti campagne di marketing (sostenibili proprio perché spalmate su grandi volumi) sono riusciti ad incrementare notevolmente il prestigio dell'azienda. **Abbinamento vincente: Tecnologia; Operations; Marketing.**

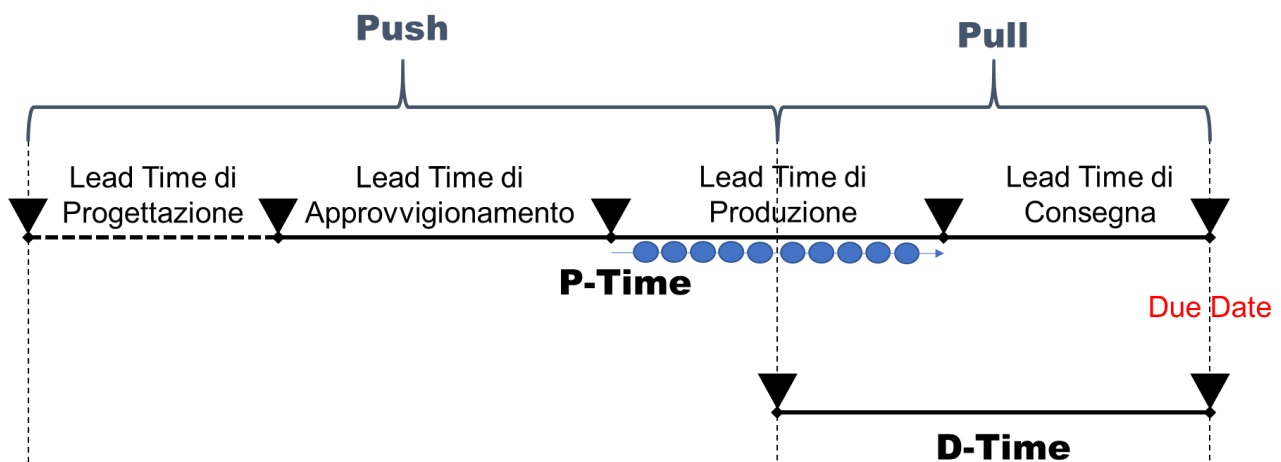


Figura 4. Tipo di produzione PUSH-PULL utilizzato da Benetton

Questo meccanismo misto è chiamato **MTS-ATO** (make to stock – assemble to order), ovvero come prima detto, consiste nel frazionamento ulteriore del processo in una prima fase di tipo PUSH seguita da una seconda fase di tipo PULL, interponendo tra le due parti un magazzino di semilavorati opportunamente dimensionato. Questo meccanismo prevede due modalità gestionali distinte; la produzione su previsione di sottogruppi standard e la successiva personalizzazione del prodotto finito in fase di assemblaggio finale (effettuata non più su previsione ma solamente a fronte di ordini acquisiti).

Tipicamente il meccanismo di PUSH è chiamato “**make to stock**” mentre quello di tipo PULL è chiamato “**make to order**”.

Per questi meccanismi è molto importante anche la fase di progettazione. *Ci sono addirittura raffinerie (produzione per processo) che hanno adottato questi meccanismi, il loro problema è la volatilità del prezzo del greggio, dunque hanno cominciato ad applicarsi su semilavorati da lavorare definitivamente all'arrivo della domanda.*

Un esempio analogo è quello del modello **Hub&spoke** per cercare di consegnare/distribuire in logica PUSH-PULL: parte della consegna consiste nel distribuire i prodotti fino ad un Hub principale (PUSH) per poi spedirli all'arrivo della domanda (PULL), equivalente del MTS-ATO della produzione.

Se ho operazioni con tempi di setup elevati allora devo cercare di inserirli nella parte PUSH, pianificando la produzione e pensando alla loro lottizzazione.

Nel PULL è il mercato che detta i tempi di produzione. **Ogni caso è specifico**, c'è chi preferisce riconvertire da una forma all'altra.

Un'azienda di prodotti biomedicali lavorava in PULL, scelta condizionata dalla data di scadenza che partiva dall'inizio della produzione (dalla sterilizzazione, che avviene all'inizio) e di conseguenza non vi erano le possibilità di fare scorte, inoltre i costi di setup erano elevati ed avevano un margine ridottissimo. La soluzione fu dedicarsi al cambiamento tecnologico, trovando un modo per sterilizzare i cateteri non all'inizio, ma alla fine della produzione, dentro la confezione. In questo modo la sterilizzazione (che è la fase costosa e impegnativa) è stata posta nella fase PULL mentre produzione e setup sono stati realizzati del tutto in PUSH.

PROCESSI PRODUTTIVI

Il meccanismo di produzione Made To Stock (PUSH) si occupa di produrre per il magazzino pensando alla domanda media da dover soddisfare (da previsione), non focalizzandosi sulle dinamiche d'oscillazione della domanda (l'equivalente del serbatoio di impianti). Il focus sarà sul volume produttivo, pertanto si guarderà maggiormente al tempo ciclo ed all'OEE per raggiungere il volume target.

Vincolo → Capacità produttiva reale almeno pari alla domanda media in questo caso (OEE molto influente)

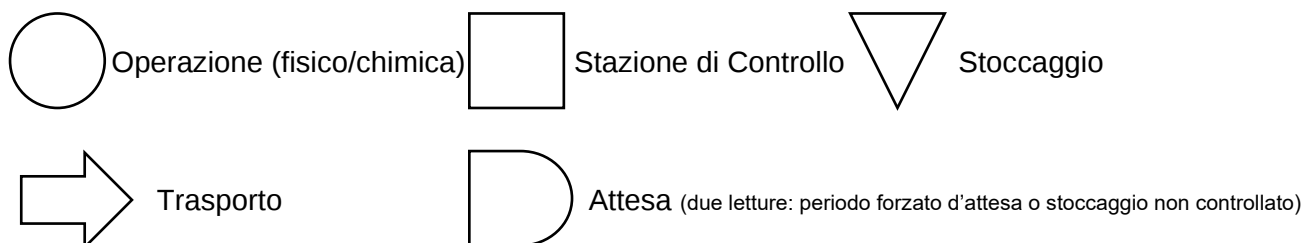
In meccanismi di produzione Assemble To Order (PULL) invece l'attenzione non si focalizza più sulla capacità produttiva e di conseguenza sull'OEE, perché non basta, conta soprattutto la velocità di reazione diventando quasi più una questione di Lead Time di produzione. Assume infatti maggior rilevanza la capacità di rispondere ad un mix produttivo sempre più variegato nei tempi stabiliti dal cliente. Si giunge al limite con la realizzazione di una sequenza di prodotti l'uno diverso dall'altro, è il cosiddetto **One piece flow**: ogni pezzo è diverso dall'altro (nessuna garanzia di lotto) e rischio di dover fare un setup per ogni pezzo, il setup diventa quasi un tempo ciclo.

Spesso per avere una reazione rapida è possibile che si adottino misure che possano danneggiare l'OEE pur di soddisfare nei tempi la domanda: mentre in PUSH gestisco i setup, in Pull devo farli in funzione della domanda che mi arriva.

Cominciamo ad entrare nell'impianto: separiamo processo produttivo da layout: poiché si deve poter lavorare su processo e layout indipendentemente.

- **Il processo** è una sequenza di attività che il mio sistema produttivo realizza.
- **il layout** è la disposizione delle macchine nell'impianto di produzione.

Rappresentazione ASME:



LO STOCCAGGIO

Vi è un'enorme differenza tra buffer e magazzino. **Prima tra tutte è che tendenzialmente un magazzino è monitorato, mentre un buffer no** (ovviamente ci sono le eccezioni: magazzini non controllati e buffer controllati).

Infatti, solitamente, i pezzi in un buffer sono soggetti a successive lavorazioni quindi non sono soggetti a controlli, mentre quando immagazzino i controlli (codici a barre ed altri strumenti) possono avvenire tramite gli R-feed, ossia i tag che sono sui prodotti al supermercato o nelle etichette dei prodotti.

Ci sono **r-feed** passivi (quelli del supermercato non alimentati da ulteriori informazioni che, sollecitati da radiazioni, restituiscono informazioni di base) e gli attivi (che essendo alimentati [aggiornati] forniscono informazioni a grande distanza).

Gli **r-feed** passivi sono arrivati a costare 1-2 cent, comincia quindi ad essere conveniente metterli ovunque.

Ciò comporta non fare più l'inventario perché permettono di conoscere all'istante quante cose ci sono e dove.

Un'altra applicazione importante di questi strumenti si ha nella catena del freddo (ovvero il mantenimento dei prodotti surgelati ad una temperatura costante e comunque inferiore a -18°C lungo tutto il percorso dalla produzione alla vendita), essendo in grado di garantire che il prodotto alimentare non scenda mai sotto un certo valore di temperatura.

I buffer dipendono anche dalla tipologia di layout che si utilizza nell'impianto. Prima tendenzialmente si utilizzavano dei canoni per cui se si produceva per flow shop (elementi/prodotti seguono sempre lo stesso percorso) si utilizzava un layout per linea (molto rigido) mentre se si produceva in job shop (elementi/prodotti possono avere differenti sequenze di lavorazioni) si utilizzava un layout per reparti. Errore! Con l'automazione posso avere macchine flessibili che possono quindi compiere diverse attività anche solo con piccole modifiche (*possono fare da tornio, fresa, trapano, ecc.*).

Ogni postazione ha un suo tempo ciclo (quanto tempo è impiegato per una lavorazione).

Una postazione è una stazione fisica in cui ci sono una o più macchine e uno o più operatori.

Il caso più semplice è: una stazione con una macchina.

Nell'impianto le risorse sono le macchine e gli operatori (anzi, ore macchina e ore uomo).

Le tipologie di macchine presenti possono essere:

- Conduzioni totalmente manuali (operatore sempre sulla macchina);
- Macchine parzialmente automatizzate (operai che conduce il macchinario per parte del tempo);
- Macchine totalmente automatizzate (nessun operaio, esisterà solo la figura dei supervisori).

Se la capacità produttiva di una macchina non è abbastanza alta allora metterò più macchine in una stessa postazione.

Una linea può essere formata da una sequenza di: due macchine, poi una, poi tre, poi di nuovo una macchina, il numero di macchine in una stazione dipende quindi dalla sua capacità produttiva. La macchina singola nella sequenza sarà una di quelle che ha la capacità produttiva più alta.

Come si calcola il tempo ciclo di una stazione di lavorazione formata da più macchine in parallelo?

$$TC_{stazione} = \frac{TC_i}{N}$$

Va da sé che maggiore è il numero di macchine che lavorano contemporaneamente all'interno della stessa stazione di lavorazione, maggiore sarà il throughput rate con cui usciranno i pezzi lavorati.

IL CONTROLLO

Un controllo potrebbe essere una stazione di controllo manuale o una macchina; a campione o al 100%, inoltre una macchina potrebbe essere in autocontrollo (incorporato).

Sono diverse le possibili soluzioni di layout che si devono adottare in funzione dell'obiettivo che si vuole raggiungere. Se si vuole concentrare l'attenzione sull'OEE è importante allora che la stazione di controllo sia prima della macchina collo di bottiglia, così da non essere influenzata negli scarti dalle macchine che ha a monte della linea, aumentando la produttività (non deve lavorare pezzi difettosi).

TRASPORTI

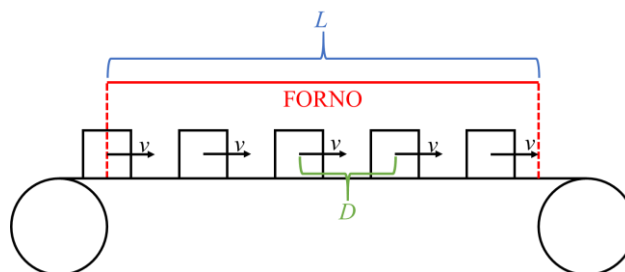
Partiamo con degli esempi particolari in cui un trasporto può essere direttamente connesso ad una lavorazione come un forno a nastro.

In un forno a nastro, lungo L e veloce v , si avrà un tempo di attraversamento pari a L/v .

Poniamo invece attenzione al tempo impiegato a far uscire due pezzi consecutivi, che sarà tutt'altro discorso e dipenderà dalla distanza inter-pezzo (D), cioè la distanza minima tra i centri di due pezzi D/v .

L'inter-pezzo è una grandezza tecnologica, dipende quindi dal tipo di macchinario che si utilizza, tempi di cottura e altro.

Il tempo ciclo in questo tipo di forni non coincide col tempo di attraversamento.



Vincoli tecnologici si trovano in molte operazioni come nella vulcanizzazione della gomma, ossia il processo in cui si riscalda la gomma ad una certa temperatura per un determinato periodo temporale, per poter godere di alcune caratteristiche tecniche (come la resistenza). Se si modifica il processo tenendo la gomma meno tempo a temperatura più alta o più tempo a temperatura più bassa,

nell'immediato avrò le stesse caratteristiche, ma sul lungo periodo la gomma avrà una resistenza diversa.

Prendiamo il caso di macchinari che lavorano in batch, ad esempio un forno. Attenzione, parlare di lotti significa che un numero di prodotti (stabilito in fase di planning) vengono processati contemporaneamente, o più o meno nello stesso periodo.

Le macchine che lavorano in batch, sono macchine in cui il tempo di lavorazione è indipendente dal carico: più carico, più aumenta la capacità produttiva media (ovviamente i pezzi escono tutti insieme).

Per esempio, mettendo 100 prodotti nel forno, per un tempo T : la capacità produttiva media risulta essere di $T/100$ (la variabilità è il numero di pezzi inseriti) il tempo di attraversamento risulta quindi essere T , ma il tempo ciclo sarà $T/100$.

Per potere avere continuità nella linea produttiva a valle dovrà essere predisposto un sistema di movimentazione che preleverà i prodotti in uscita dal forno con un ritmo di almeno $T/100$.

Per le movimentazioni (trasporto) il tempo di attraversamento è quello che intercorre da quando si caricano uno o più pezzi a quando li si smontano:

$$T_{Attraversamento} = \frac{D}{v}$$

il tempo ciclo (tempo medio tra due pezzi successivi) dipende da quanti pezzi vengono caricati (considerando anche il ritorno):

$$T_{Ciclo} = 2 \cdot \frac{D}{v} \cdot \frac{1}{n}$$

dove n è il numero di pezzi trasportati.

OEE

Le possibilità di dover realizzare da zero un impianto sono decisamente inferiori rispetto a quelle di dover ottimizzare un impianto già esistente, su cui sarà possibile realizzare piccole modifiche.

Le possibilità sono diverse a seconda della necessità e delle possibilità; può capitare che vada aumentato il numero di macchine in una stazione di lavoro non solo per aumentare la Capacità produttiva, ma anche per aumentarne l'affidabilità (macchine a ridondanza), quest'ultima configurazione spesso usata negli impianti di servizio (dove la manutenzione non è banale) oppure nel farmaceutico, per esempio, nel caso di farmaci salvavita la cui produzione deve essere sempre garantita.

Cos'è l'OEE? C'è chi lo distingue tra due tipi **OEE 1** e **2**.

L'OEE di tipo 1 risulta molto difficile da calcolare in una linea di produzione, è molto teorico e poco pratico mentre l'OEE di tipo 2, definito anche di Nagashima, è un indicatore relativamente recente, ormai comune e standardizzato capace di misurare le perdite di tempo e l'efficienza tramite l'utilizzo delle ore macchina.

Per esempio, se vengono scartati 5 pezzi ogni 100 si ha interesse nel sapere quanto tempo viene perso in quei 5 pezzi, solo in un secondo momento si vorrà indagare sul numero di pezzi che effettivamente non si possiedono più.

Attenzione, non ha senso parlare di OEE di una linea, ma ci riferiremo alla macchina.

L'OEE di una linea sarebbe infatti un'informazione troppo aggregata, nel caso in cui volessimo migliorare per esempio tale indicatore da dove si dovrebbe cominciare? Per questo si utilizza solitamente un OEE riferito alla singola stazione di lavorazione.

Prendiamo il periodo lavorativo di un giorno, come mostrato in Figura 5, il tempo solare (le 24 ore) conterrà il tempo di apertura che potrà coincidere con quest'ultimo o no (dipende dal numero di turni al giorno).

La macchina potrebbe essere caricata per tutto il tempo di apertura o meno (Tempo Carico), nel tempo in cui è caricata potrebbe fermarsi per le fermate pianificate come setup, manutenzione pianificata, breaks & pranzi, audit sulla qualità e meetings/training.

Il tempo operativo lordo è costituito dal tempo carico sottraendo tutte le fermate pianificate.

Continuando la scomposizione, il tempo operativo lordo potrebbe contenere ulteriori micro-fermate e rallentamenti, ovvero tutte le fermate non pianificate.

La micro-fermata non è misurata come un'attività di manutenzione: in una realtà ancora non totalmente automatizzata non si possono registrare anche le fermate di un minuto (rischiamo che ci fornisca un dato sporco, alterato, non veritiero).

Il tempo operativo netto sarà il tempo operativo lordo a cui vanno sottratte le perdite di tempo come rallentamenti o micro-fermate. All'interno del tempo operativo netto possono esserci ulteriori perdite di tempo legate agli sprechi dovuti alla lavorazione di scarti o all'esecuzione di rilavorazioni.

Togliendo queste ulteriori perdite di tempo si arriva al tempo operativo netto a valore aggiunto.

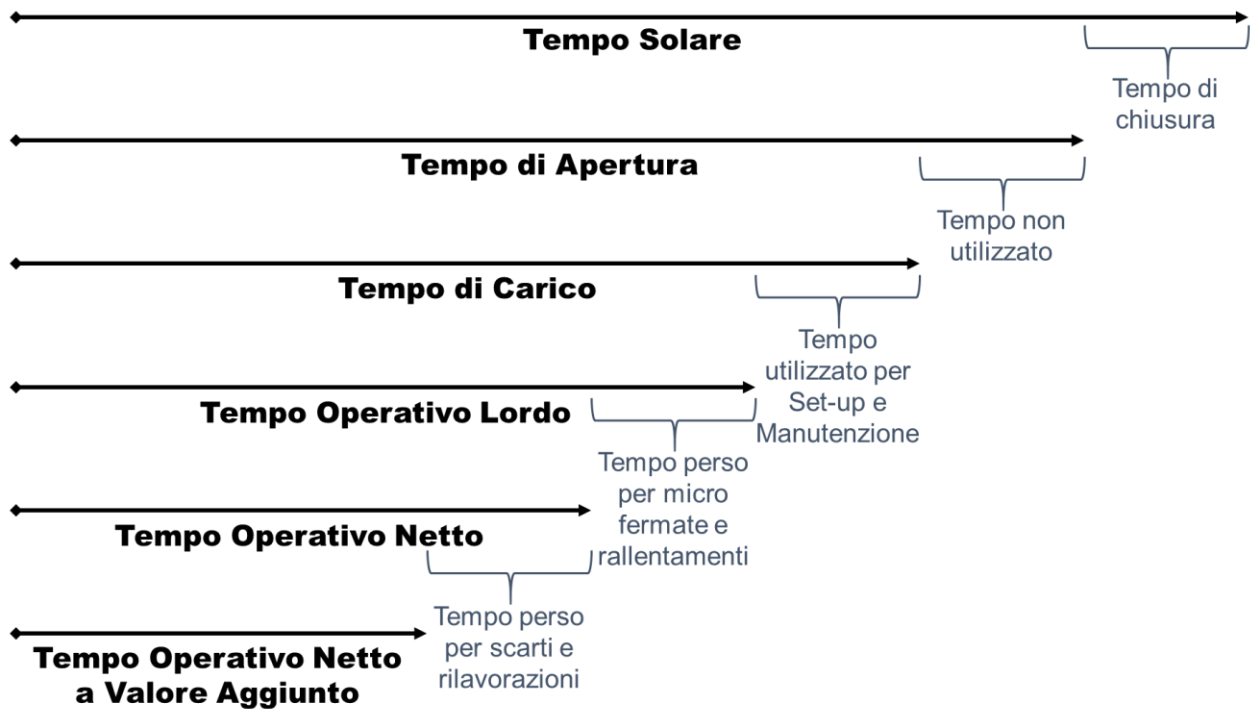


Figura 5. Scomposizione dell'OEE

$$OEE = \frac{\text{Tempo operativo netto a valore aggiunto}}{T_{\text{carico}}} = Q \times Ep \times D$$

Supponendo di fare N pezzi buoni al giorno, con un TCt (tempo ciclo teorico):

$$T_{\text{op.n.v.a.}} = N \times TCt$$

$$T_{\text{op.netto}} = \frac{N \times TCt}{Q} \quad \text{dove} \quad Q = \frac{T_{\text{op.n.v.a.}}}{T_{\text{op.netto}}} \quad (\text{tasso di qualità})$$

$$T_{\text{op.lordo}} = \frac{N \times TCt}{Q \times Ep} \quad \text{dove} \quad Ep = \frac{T_{\text{op.netto}}}{T_{\text{op.lordo}}} \quad (\text{efficienza delle prestazioni})$$

$$T_{\text{carico}} = \frac{N \times TCt}{Q \times Ep \times D} \quad \text{dove}$$

$$D = \frac{T_{\text{op.lordo}}}{T_{\text{carico}}} \quad (\text{disponibilità}) \rightarrow T_{\text{carico}} = \frac{T_{\text{op.n.v.a.}}}{OEE} = \frac{N \times TCt}{OEE}$$

L'OEE indica quanto tempo in più si ha bisogno di caricare la macchina per fare N pezzi buoni (fondamentale per poter rispondere alla domanda).

L'OEE di tipo 1 misurato davanti ad una macchina presenta numerose difficoltà di calcolo.

Come calcolo Ep ? Come calcolo il tempo perso per i pezzi difettosi? Non è detto che la macchina lavori in autocontrollo (vedo quanti pezzi scarta), se la qualità è rilevata a valle come imputo la difettosità a livello di singola macchina?

L' OEE più pratico, si calcola davanti alla macchina, prendendo in considerazione il tempo solare, di apertura e carico (tutti e tre misurabili facilmente), cercando, se possibile, di misurare anche le fermate classificandole tra programmate e non programmate.

Le fermate programmate sono ad esempio le pause pranzo (collocate all'interno del tempo carico, ed è bene considerarle nell'OEE per attività di pianificazione [trovare le soluzioni migliori: come schedulare in più pause, ecc.]), i setup, la manutenzione preventiva, riunioni, formazione, assemblea sindacale.

Entrambe i tipi di fermate, programmate e non, impattano sull'OEE, ma in modo molto diverso, la prima la programmerò in momenti strategici dove ho disagi minori, mentre la seconda avviene casualmente durante il tempo di funzionamento della macchina.

Le fermate non programmate possono essere classificate in esterne o interne.

Tra le fermate non programmate interne abbiamo:

- Down Time, ovvero la manutenzione a guasto (durata fra la rottura e il ripristino, comprendendo anche il tempo di fermo in attesa del manutentore), piccole fermate e errori;
- Perdite di Velocità.

Tra le fermate non programmate esterne invece abbiamo:

- **Starving Time** (alla macchina non arrivano risorse, ho quindi un blocco a monte);
- **Blocking Time** (la macchina non può far uscire pezzi perché non sa dove metterli, blocco a valle).

Altre informazioni misurabili sulla singola macchina sono le lavorazioni totali (quindi anche le difettose) e quelle buone, sempre solo se la macchina lavora in autocontrollo (altrimenti sono tutti buoni). Come ricavo a questo punto l'OEE? In genere si ha come dato di targa il TCt della macchina, potendo conoscere il rapporto tra i pezzi buoni e i pezzi totali si ottiene qualcosa di simile al tasso di qualità (calcolo il tempo operativo netto e netto a valore aggiunto).

Tuttavia, si hanno problemi nel calcolo dell'efficienza delle prestazioni (come calcolo micro-fermate e rallentamenti se non sto in completa automazione?). Spostando l'attenzione sul tempo operativo lordo ci si accorge che è possibile calcolare questo tempo come differenza tra il tempo carico e le fermate (quelle misurabili, programmate e non), e utilizzando il rapporto tra l'operativo lordo così calcolato e il netto calcolato precedentemente si ottiene l' Ep .

Spesso si fa riferimento all'efficienza delle prestazioni come tempo "unknown", cioè tempo sconosciuto. La ragione di questo nome deriva dalle poche informazioni che si hanno in merito, non sapendo operativamente cosa sia successo (potrebbe finirci dentro anche una fermata teoricamente misurabile che però l'operaio non ha rilevato). L'unknown ha quindi l'obiettivo di catturare tutte quelle perdite di inefficienza che non sono riuscito a misurare.

Per esempio: una fermata di 15 minuti la metterò nell'unknown? La risposta è dipende da come viene raccolto il dato. Se il dato per esempio viene misurato da un operatore, allora ci si metterà d'accordo con lui sulla soglia di durata oltre la quale segnerò la fermata (*di certo se gli andrò a chiedere di segnare tutte le volte che la macchina si è fermata per meno di due minuti non sarà molto contento,*

rischiando di scrivere un dato per un altro, il che comprometterebbe la misura dell'OEE). Meglio avere meno dati ma precisi piuttosto che molti dati ma “sporchi”. A decidere cosa andrà nell'unknown sarà pertanto la volontà nel tener traccia delle fermate. *Se con l'operatore si pattuisce che segnerà le fermate che dureranno più di 20 minuti allora probabilmente quella fermata di 15 minuti andrà a finire nell'unknown, perché non è stata misurata con precisione da nessuna parte.*

[mi preoccupa quando questo valore sconosciuto assume valori troppo grandi oltre 5-10%].

Il rapporto tra il tempo carico e quello di apertura mi fornisce il grado di utilizzazione:

$$GU = \frac{T_{\text{carico}}}{T_{\text{apertura}}}$$

Il GU si può considerare un indicatore di performance?

Se così fosse mi aspetterei che tanto più è alto il GU tanto migliore sarà la performance della macchina, così come funziona per l'OEE. Ma è veramente così? In realtà il GU aumenta se aumenta il Tempo Carico, il quale a sua volta potrebbe aumentare anche perché la macchina è guasta, oppure perché è aumentato il tempo di setup, manutenzione (preferibile che tenda a 1).

Questo indicatore quindi dipende da come si utilizza la macchina, anche se è alto non è detto che stia facendo bene.

L'indicatore migliore è:

$$Ut = GU \times OEE = \frac{T_{\text{op. n. v. a.}}}{T_{\text{apertura}}}$$

si tratta dell'utilizzo e misura quanto tempo del tempo di apertura sto producendo pezzi buoni (indicatore di performance).

Riepilogando dunque l'OEE è un indicatore di efficienza, in quanto cattura quanto tempo stiamo impiegando per poter produrre un determinato numero di pezzi buoni. Va da sé che, a parità di pezzi buoni prodotti, minore è il tempo impiegato maggiore sarà l'OEE.

Vedendo nel dettaglio quali fattori entrano in gioco, si nota che aumentando uno di questi tre fattori è possibile migliorare l'OEE.

$$OEE = D * E_p * Q$$

Ma è vero per tutti e tre?

L'OEE ha un punto debole: il setup. Il tempo dedicato ai setup dipende da quanti se ne fanno e quanto durano. La durata del setup è l'indicatore di performance (meno durano meglio è), ma il numero no: meno setup significa irrigidimento, significa quindi non sfruttare la versatilità della macchina, significa monoprodotto (e non è detto che sia la soluzione migliore).

Se si è in presenza di due impianti gemelli, che presentano però due OEE differenti, uno 0,8 l'altro 0,95. Non è detto che l'impianto migliore sia quello con l'OEE maggiore! Devo andare a vedere i tempi di setup e il numero di setup fatti da entrambi gli impianti.

Al contrario supponiamo che due impianti che producono lo stesso mix di prodotti finiti, dedichino entrambi 100 ore all'anno per fare setup. Tuttavia, un impianto fa solo 4 setup all'anno della durata di 25 ore l'uno, mentre l'altra fa 100 setup all'anno con una durata di 1 ora ciascuno. A parità di disponibilità D, quali delle due opererà con maggiore efficienza?

Una linea di produzione è una sequenza di macchine o stazioni di operazioni manuali che realizzano un prodotto (sequenza di operazioni che realizza un processo, anche fosse in job o flow shop).

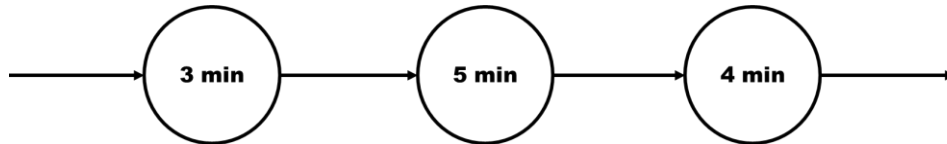


Figura 6. Rappresentazione di una linea in cui vengono rappresentati i tempi ciclo all'interno di ciascuna operazione.

Il Takt Time della linea può essere scelto nell'intervallo:

$$\max(TC_i) \leq TC \leq \frac{1}{D}$$

La scelta deve tener conto che all'aumentare di TC l'impianto viene stressato meno ma diminuiscono i tempi morti a disposizione, utili ad esempio per la manutenzione, inoltre aumenta il rischio di non riuscire a rispondere alla domanda in caso di imprevisti, rischio sensibile soprattutto nel caso in cui la stazione più lenta si trovi nella zona terminale del processo.

Balance delay: misura lo sbilanciamento all'interno della linea, cioè quanta capacità produttiva perdo a causa dello sbilanciamento (per colpa quindi della macchina più lenta **BottleNeck**).

Il Balance delay è un indicatore di performance della linea di produzione.

$$BD = \frac{\sum_{i=1}^m (TC_{BN} - TC_i)}{m * TC_{BN}}$$

ovviamente è 0 se tutte le operazioni all'interno della linea produttiva hanno lo stesso tempo ciclo, più aumenta lo sbilanciamento fra di esse e più tende a 1 (non si arriverà mai ad 1 perché dovrei avere tutti gli altri tempi ciclo tranne la stazione collo di bottiglia pari a 0).

Ho bisogno di portare una piccola correzione alla formula prima vista, quando una stazione è composta da macchine in parallelo, *ad esempio la terza stazione è costituita da due macchine che hanno entrambe un tempo ciclo pari ad 8 minuti (quindi la postazione avrà complessivamente un tempo ciclo di 4 minuti).*

$$BD = \frac{\sum_{i=1}^n m_i * (TC_{BN} - TC_i)}{\sum_{i=1}^n m_i * TC_{BN}}$$

dove m_i è il numero di macchine nella stazione ed n è il numero di stazioni.

Se arrivasse il manutentore per la manutenzione preventiva, in un momento in cui la macchina non era carica, cioè non avevo in programma di farla lavorare, influisce sull'OEE o sul GU? Dipende da noi e da come utilizziamo i parametri. La manutenzione potrebbe essere fatta durante un turno

straordinario, addirittura durante i periodi in cui l'impianto è chiuso (quindi non solo fuori dal tempo carico, ma anche fuori dal tempo di apertura).

Se ho $CP < D$ allora la macchina collo di bottiglia non starà mai ferma, non ci sarà rischio che sto "fuori" dal tempo carico, nell'OEE ci metterò tutto, lo "straordinario" non esiste, è tutto tempo da dedicare alla produzione.

Diversamente se $CP > D$ se ho una manutenzione che potrebbe essere realizzata fuori dal tempo carico forse non la includo nell'OEE, avevamo detto che ci saremmo concentrati principalmente sui costi; metterò nell'OEE ciò che succede quando l'operaio è sulla macchina.

Scelte che dipendono principalmente dall'Operations Manager, l'importante è che si utilizzino in maniera corretta gli indicatori.

In generale tutto ciò che viene incluso nell'OEE diventa di fondamentale importanza ed è soggetto a tentativi di miglioramento {ai fini dell'esame metteremo quasi sempre tutto nell'OEE, quanto meno se gli interventi vengono fatti nell'orario di apertura}.

Per quanto riguarda la macchina con tempo ciclo pari a 3 minuti nella Figura 6, i due minuti per i quali non è utilizzata è OEE o GU? **GU**.

Dal confronto tra BD e GU posso capire la natura della non utilizzazione della macchina, capendo anche il perché del suo livello di GU.

Si può pensare di flessibilizzare così la linea: se si devono fare in un giorno 1000 pezzi, la prima stazione ($TC=3$ min) li farà in 3000 minuti e la seconda lavorerà lo stesso materiale su 5000 minuti. Quindi dopo 3000 minuti si può valutare di far fare altro alla macchina, ovviamente con lo scotto di dover avere un buffer che potrebbe costare tanto o in termini di valore o in termini di volume (è chiaro che pensare di fare per 3 minuti un prodotto e per altri 2 un altro diventa un'impresa difficoltosa e confusionaria).

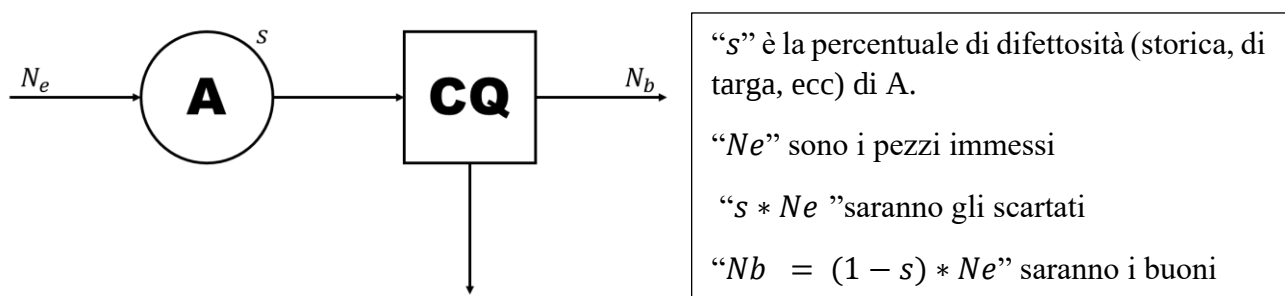
TASSO DI QUALITA'

Si tratta del tempo perso per problemi legati alla qualità.

Dipende molto da dove trovo i pezzi difettosi nella linea [utopicamente se non ho controlli per me è tutto buono, al massimo il cliente mi farà il reso (il cliente stesso diventa il controllo qualità)], si deve tener conto quindi della totalità della linea di produzione, per esempio se una macchina ha un tasso di difettosità pari ad 0 ma è connessa con un'altra macchina che ha un tasso di difettosità strettamente compreso tra 0 e 1, allora anche la prima macchina ne risentirà perché avrà perso del tempo per lavorare dei pezzi che poi verranno scartati per problemi legati non alla sua lavorazione, ma a quelle dovute alle altre macchine presenti nella linea.

La stazione di lavoro Controllo di Qualità (CQ) sarà quella stazione che riesce a scaricare le difettosità delle macchine poste a monte di essa dalle macchine a valle. La stazione CQ, mi garantisce che le macchine poste a valle non risentiranno delle difettosità delle macchine a monte, poiché i pezzi difettosi vengono individuati e instradati in altri tipi di percorsi (vengono scartati o rilavorati in altre stazioni).

Prendiamo come esempio il caso più semplice:



Calcoliamo il tasso di qualità di A:

$$Q_A = \frac{T_{op.n.v.a.}}{T_{op.n.}} = \frac{TC_{tA} * Ne * (1 - s)}{TC_{tA} * Ne} = 1 - s$$

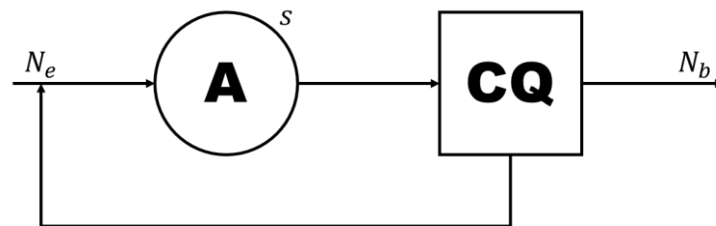
Alla base c'è un'ipotesi, ovvero che il tasso di difettosità di CQ è zero. Questo non significa che il suo tasso di qualità sia pari ad 1, tutt'altro si avrà che:

$$Q_{CQ} = \frac{TC_{tCQ} * Ne * (1 - s)}{TC_{tCQ} * Ne} = 1 - s$$

Comincia ad essere chiaro che due macchine poste in serie senza possibilità di controlli nel mezzo hanno lo stesso tasso di qualità: **lavorano lo stesso numero di pezzi buoni e difettosi.**

Modificando la tipologia di linea potremmo avere diverse configurazioni, qui di seguito ne vengono riportate a titolo di esempio alcune:

1) Il controllo qualità processa infinite volte i pezzi difettosi:



In genere, come vedremo in seguito, esiste un sistema di ripristino (stazione di lavoro/ macchinario) per i pezzi che vengono individuati difettosi da CQ e vogliono essere reimmessi nella linea di produzione. Disinteressandoci di questa piccola parentesi, analizziamo il caso corrente in cui supponiamo di nuovo che il tasso di difettosità di CQ sia nullo, ad ogni passaggio dei pezzi parte di essi verrà scartato e parte sarà buono. Per una certa quantità di prodotti che entrano nel sistema parte di essa tornerà indietro, iterando tale processo n volte si arriverà ad avere che tutti i prodotti immessi saranno buoni. Facciamo i conti sui passaggi dei pezzi attraverso la seguente tabella:

A	CQ	N _B
N_e	N_e	$(1 - s) * N_e$
$s * N_e$	$s * N_e$	$(1 - s) * s * N_e$
$s^2 * N_e$	$s^2 * N_e$	$(1 - s) * s^2 * N_e$
.	.	.
.	.	.
.	.	.
$N_e(1 + s + s^2 + s^3 + \dots) = \frac{N_e}{1 - s}$	$N_e(1 + s + s^2 + s^3 + \dots) = \frac{N_e}{1 - s}$	N_e

Applicando la formula:

$$Q_A = Q_{CQ} = \frac{TC_t * N_e}{TC_t * \frac{N_e}{1 - s}} = 1 - s$$

Ciò significa che nelle rilavorazioni ho lo stesso risultato delle non rilavorazioni (primo caso analizzato semplicistico) in termini di OEE, cambierà solamente il numero di materie prime che dovranno essere comprate per produrre N_B (costo materie prime).

2) Caso in cui la macchina A rilavora soltanto una volta gli scarti del controllo qualità:

Se, diversamente dal caso precedente, il secondo passaggio per la stazione di lavorazione A dei pezzi difettosi produrrà sicuramente pezzi buoni, si avrà che:

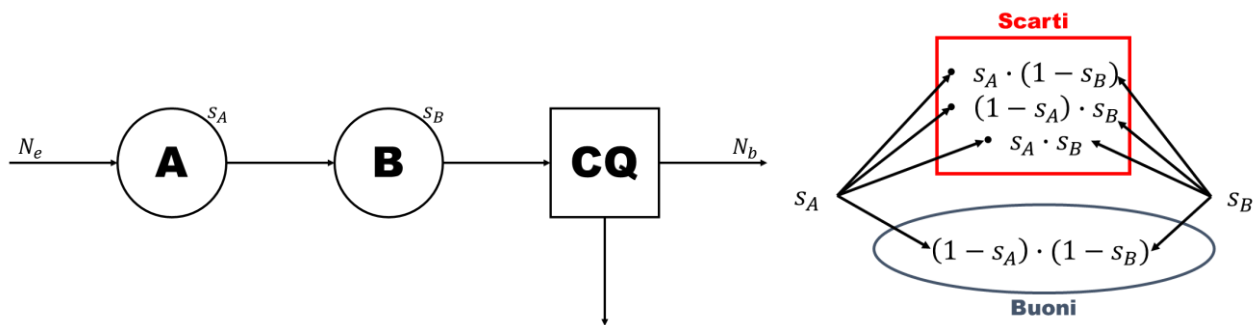
A	CQ	N _B
Ne	Ne	$(1 - s) * Ne$
$s * Ne$	$s * Ne$	$(1 - s) * s * Ne$
$(1 + s) * Ne$	$(1 + s) * Ne$	Ne

Con:

$$Q_A = Q_{CQ} = \frac{TC_t * N_e}{TC_t * N_e * (1 + s)} = \frac{1}{1 + s}$$

3) Caso di due macchine in serie con a valle un controllo qualità:

Caso analogo al caso precedente se si considerano le due macchine A e B un'unica macchina con tasso di qualità pari al prodotto dei due tassi di qualità



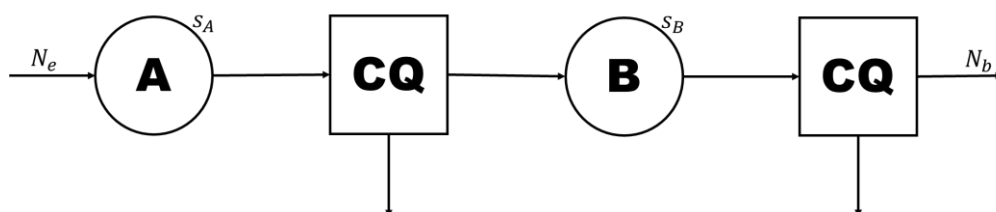
$$Q_A = Q_B = Q_{CQ} = \frac{TC_t * N_e * (1 - s_A)(1 - s_B)}{TC_t * N_e} = (1 - s_A)(1 - s_B)$$

A	B	CQ	N _B
Ne	Ne	Ne	$(1 - s_A)(1 - s_B) * Ne$

Se anche CQ avesse avuto un tasso di difettosità diverso da zero allora sarebbe stato:

$$(1 - s_A)(1 - s_B)(1 - s_{CQ})$$

4) Caso con più controlli qualità in serie:



A	CQ	B	CQ	N _B
Ne	Ne	$(1 - s_A) * Ne$	$(1 - s_A) * Ne$	$(1 - s_A)(1 - s_B) * Ne$

In questa configurazione il primo controllo di qualità scarica la difettosità di A e fa sì che non influenzi la stazione di lavorazione B, che non lavora pezzi difettosi da A.

Se B fosse stata la macchina collo di bottiglia sarebbe la soluzione migliore da adottare.

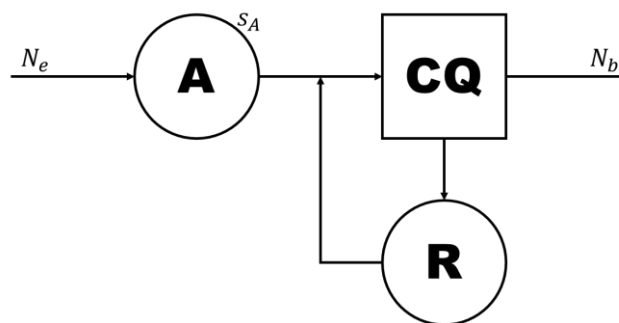
Avremo:

$$Q_A = Q_{CQ1} = (1 - s_A)(1 - s_B); \quad Q_B = Q_{CQ2} = (1 - s_B)$$

Nel caso della stazione di lavoro B si avrà un tasso di qualità più alto, quindi un OEE più alto. Questa stazione se pur posta in serie con una stazione CQ perderà comunque tempo a causa delle lavorazioni difettose della stazione di lavoro B (quindi dal suo tasso di difettosità).

5) Caso di rilavorazioni fuori linea, con re-immissione a monte del CQ:

Vediamo qualche caso di rilavorazione fuori linea, cioè la rilavorazione non la fa la stessa stazione di lavoro A. Analizziamo il caso in cui è possibile pensare che i pezzi, uscenti dalla stazione di rilavorazione, siano sicuramente buoni ma voglio comunque ricontrollarli:



A	CQ	R	N _B
Ne	Ne	$s_A * Ne$	$(1 - s_A) * Ne$
/	$s_A * Ne$	/	$s_A * Ne$
Ne	$(1 + s_A) * Ne$	$s_A * Ne$	Ne

In questo caso $Q_A = 1$ perché pur avendo fatto difetti la macchina A non ha perso tempo per rilavorare i propri difetti, infatti quel tempo che A non ha sprecato per i propri difetti lo ha perso la stazione di rilavorazione R.

La domanda da porsi infatti è: la macchina A ha impiegato del tempo per lavorare un pezzo che uscirà dal sistema senza che questo diventi prodotto finito (perché difettoso)? E ancora, la macchina A sta utilizzando parte del suo tempo per processare un semilavorato che ha già lavorato precedentemente? Se la risposta è no per nessuna delle due domande, allora A impiega tutto il suo tempo per processare oggetti che sono poi diventati prodotti finiti.

Naturalmente nel caso più generico, al tasso di qualità di A andrebbe aggiunto poi il tasso di scarto di eventuali macchine a valle che fanno in modo che, il tempo della macchina A impiegato per la lavorazione di alcuni semilavorati sia perso perché il semilavorato è uscito dal sistema in quanto difettoso.

Anche se sembra anti-intuitivo, la filosofia è che conta il tempo perso e non i pezzi.

In questo caso si avrà che:

$$Q_{CQ} = \frac{1}{1 + s_A}$$

Per quanto riguarda R non si parla di tasso di qualità: è assurdo parlare di perdita di tempo, essendo una macchina che esiste proprio e solo per lavorare i pezzi difettosi prodotti.

Si dovrà invece avere particolare premura nel verificare che R abbia un'adeguata CP al fine di riuscire a rilavorare tutti i prodotti difettosi nel tempo disponibile.

Diversamente da quello che dicevamo ad impianti un sistema di movimentazione può essere collo di bottiglia, ovviamente in fase di progettazione lo sovradimensionerò (visto che di solito la movimentazione non è così costosa).

Può succedere però che in fase di ottimizzazione velocizzo le macchine e il sistema di movimentazione si ritrova "indietro" (se si ha capacità produttiva maggiore della domanda può anche non interessare).

Anche una stazione CQ può essere collo di bottiglia, una stazione di rilavorazione no.

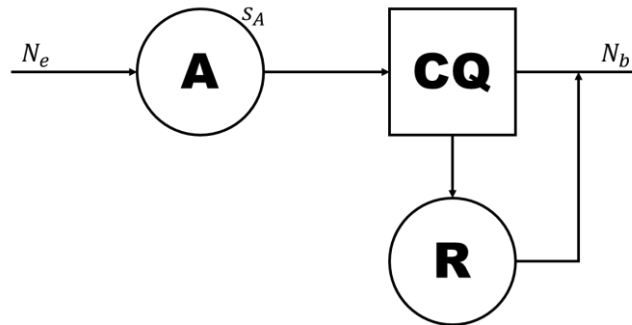
Una stazione di rilavorazione che risulta essere collo di bottiglia significa che la produzione viene rallentata perché non si è in grado di rilavorare i pezzi (arrivando al paradosso che per poter rispondere alla domanda sarebbe meglio scartarli e non rilavorarli tutti, o almeno venderli come sottoprodotto o altre opzioni che sicuramente risultano migliori rispetto al perdere domanda, [eccezion fatta per i pezzi di altissimo valore che devono essere rilavorati, quindi, in questa ipotesi, la stazione rilavorazione potrebbe essere la macchina collo di bottiglia]).

Per la stazione di rilavorazione R non ha senso parlare di OEE, si dimensionerà la stazione sulla base della linea e vedremo (con un riferimento ad un'analisi di costi) se converrà o meno utilizzare tale stazione in caso non convenga significa che converrà scartare direttamente alcuni pezzi.

Discorso simile per la stazione di controllo della qualità CQ: se il fattore qualità non è così fondamentale non sarà necessario controllare tutto (piuttosto che rallentare la produzione).

6) Rilavorazione esterna dei prodotti difettosi, con reimmisione a valle del controllo di qualità:

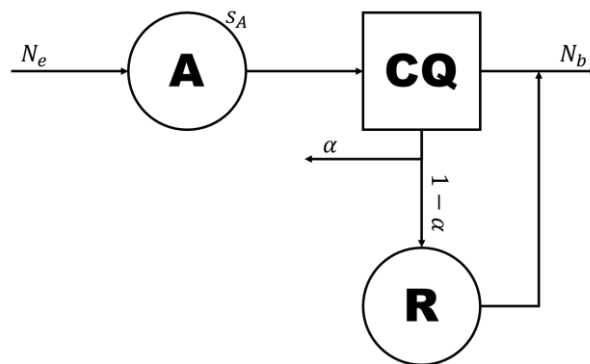
Stazione di rilavorazione R affidabile, non si ricontrollano le rilavorazioni.



A	CQ	R	N _B
Ne	Ne	$s_A * Ne$	$(1 - s_A) * Ne$
/	/	/	$s_A * Ne$
Ne	Ne	$s_A * Ne$	Ne

$$Q_A = Q_{CQ} = 1$$

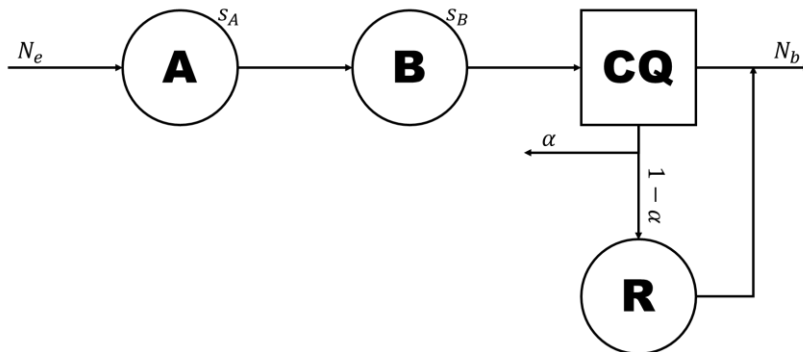
7) Rilavorazione di una sola parte della percentuale $(1 - \alpha)$ dei prodotti scartati:



A	CQ	R	N _B
Ne	Ne	$(1 - \alpha) * s_A * Ne$	$(1 - s_A) * Ne$
/	/	/	$(1 - \alpha) * s_A * Ne$
Ne	Ne	$(1 - \alpha) * s_A * Ne$	$(1 - \alpha * s_A) * Ne$

$$Q_A = Q_{CQ} = \frac{(1 - \alpha s_A) * N_e * TC_t}{N_e * TC_t} = (1 - \alpha s_A)$$

8) Rilavorazione di una sola parte della percentuale $(1 - \alpha)$ dei prodotti scartati con due lavorazioni a monte:



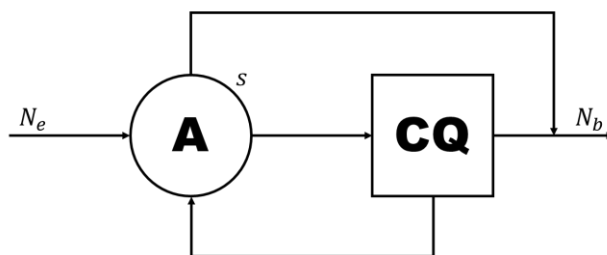
$$Q_A = Q_B = Q_{CQ} = \frac{N_B}{N_E}$$

A	B	CQ	R	N _B
N_e	N_e	N_e	$(1 - \alpha)(1 - (1 - s_B)(1 - s_A)) * N_e$	$(1 - s_A)(1 - s_B) * N_e$
/	/	/	/	$(1 - \alpha)(1 - (1 - s_B)(1 - s_A)) * N_e$
N_e	N_e	N_e	$(1 - \alpha)(1 - (1 - s_B)(1 - s_A)) * N_e$	$[1 - \alpha(1 - (1 - s_B)(1 - s_A))] * N_e$

Sarebbe utile “unire” tutto quello che c’è a monte del controllo qualità come se fosse una sola macchina/stazione di lavoro con tasso di difettosità $\bar{s}$ pari al totale delle macchine che in questo caso si ha, unendo A e B, ottenendo un tasso di difettosità totale pari a: $\bar{s} = 1 - (1 - s_A)(1 - s_B)$.

Con il nuovo tasso di difettosità $\bar{s}$ mi posso ricondurre all’esempio precedente (6) e constatare che il tasso di qualità di A, B, CQ (senza fare tutti quei terribili conti) è $1 - \alpha\bar{s}$.

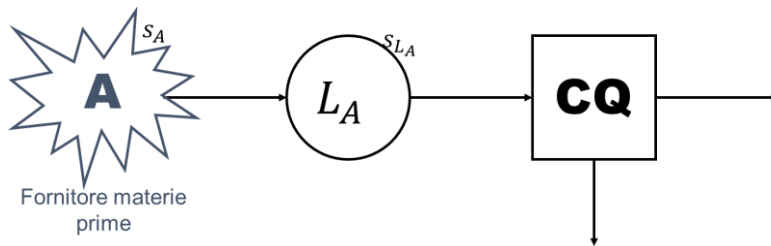
9) Caso in cui la macchina A realizza rilavorazioni che andranno sicuramente a buon fine e non ripasseranno per CQ:



A	CQ	N _B
N_e	N_e	$(1 - s_A) * N_e$
$s_A * N_e$	/	$s_A * N_e$
$(1 + s_A) * N_e$	N_e	N_e

$$Q_A = \frac{1}{1 + s_A} \text{ \& } Q_{CQ} = 1$$

10) Caso in cui la materia prima sia difettosa:

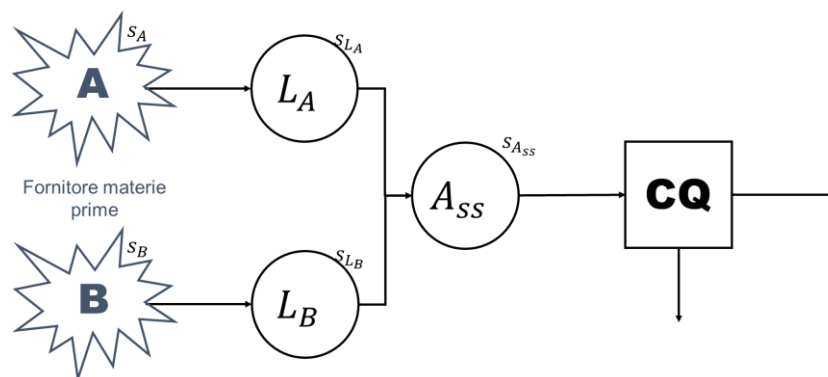


Lo schema è del tutto equivalente a quelli visti in precedenza con due operazioni con due difettosità peculiari. L'unica particolarità che emerge è che al primo step c'è il fornitore, le cui materie prime possono appunto presentare difetti.

$$Q_{LA} = Q_{CQ} = (1 - s_{LA})(1 - s_A)$$

dove s_A è la difettosità della materia prima. Ovviamente non ha senso parlare del tasso di difettosità delle materie prime.

11) Caso di processo convergente con più materie prime difettose:



Il tasso di qualità di L_A (quindi le sue perdite di tempo) è influenzato dal proprio tasso di difettosità, dalla difettosità della materia prima A, dalla difettosità dell'assemblaggio, ma non solo anche del tasso di difettosità della lavorazione L_B e la difettosità della materia prima B.

Tutte queste ragioni determinano la difettosità del prodotto che non passando nel controllo qualità rischia di incrementare criticamente le perdite di tempo della lavorazione L_A .

Ovviamente stesso discorso vale per il tasso di qualità di L_B e di AS, il che comporterà:

$$Q_{LA} = (1 - s_{LA})(1 - s_A)(1 - s_{AS})(1 - s_B)(1 - s_{LB}) = Q_{LB} = Q_{CQ} = Q_{AS}$$

Nel caso in cui A risulta essere un componente di alto valore e B un componente di basso valore metterò un controllo di qualità dopo la lavorazione di B, così che all'assemblaggio arrivino solo pezzi B buoni. Secondo questa nuova configurazione, la lavorazione di A non avrà perdite di tempo per quanto riguarda la parte di linea dedicata a B (non scarteremo i componenti A per colpa dell'assemblaggio con componenti B difettose, $Q_{LA} \nearrow$).

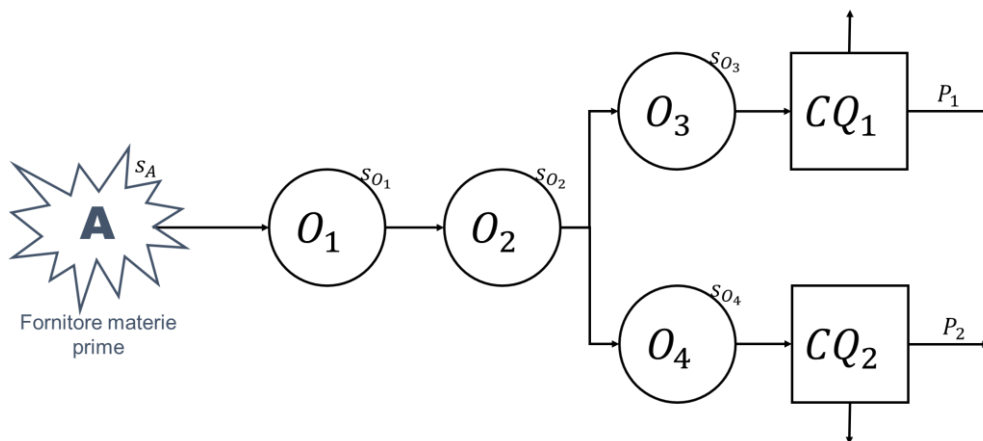
Il tasso di qualità di L_B , invece, non cambierà. Il tasso di qualità di L_A sarà influenzato soltanto dalla difettosità della materia prima di A, delle sue lavorazioni di A e dall'assemblaggio:

$$Q_{LA} = (1 - s_{LA})(1 - s_A)(1 - s_{AS})$$

Generalizzando al caso in cui le componenti di materia prima che compongono il prodotto/semilavorato siano “n” per A e “m” per B, dobbiamo tener conto del fatto che la difettosità dipende dalla singola materia prima sia nel contesto materia prima sia nel contesto lavorazione (consideriamo sempre il CQ come nello schema precedente):

$$Q_{LA} = (1 - s_{LA})^n (1 - s_A)^n (1 - s_{AS})(1 - s_B)^m (1 - s_{LB})^m = Q_{LB}$$

12) Caso di processo divergente con materia prima difettosa:



Il processo divergente è un processo molto comune nelle produzioni in cui si tende ad esaltare la personalizzazione del prodotto. Se si avesse avuto una linea semplice per ogni personalizzazione (quindi non ci sarebbe stata l'operazione 4) allora il tasso di qualità di tutte le macchine era semplice:

$$Q = (1 - s_1)(1 - s_2)(1 - s_3)(1 - s_A)$$

Ma non è così, il caso corrente è più complesso dovendo distinguere le operazioni 1 e 2 a seconda se l'operazione che stanno facendo finirà su P1 oppure su P2 (sono due flussi diversi).

Quando viene lavorato un pezzo che finirà su P1 allora il tasso sarà $Q = (1 - s_1)(1 - s_2)(1 - s_3)$, mentre nella lavorazione di un pezzo che finirà su P2 sarà $Q = (1 - s_1)(1 - s_2)(1 - s_4)$.

Quindi non c'è in un tasso di qualità assoluto per la macchina 1 o 2.

Per poter indicare il tasso di qualità bisogna tenere conto della macchina che compie le lavorazioni (dati di targa), contesto (tipo di linea [dove stanno i controlli qualità]) e flusso (tipo di prodotto che si vuole realizzare).

Secondo questa configurazione (11) avremmo che i due flussi incidenti nelle due macchine 1 e 2 saranno caratterizzati ciascuno da un tasso di qualità diverso, avendo:

$$\begin{aligned} Q_{1P1} &= (1 - s_1)(1 - s_2)(1 - s_3)(1 - s_A) = Q_{2P1} = Q_3 = Q_{C1} \\ Q_{1P2} &= (1 - s_1)(1 - s_2)(1 - s_4)(1 - s_A) = Q_{2P2} = Q_4 = Q_{C2} \end{aligned}$$

Quindi ai fini di ottimizzazione sarebbe utile un controllo qualità della materia prima, considerando che un pezzo difettoso di A mi rovina "n" pezzi di P1 ed "m" di P2. Se la difettosità della materia prima risulta pressochè inesistente sarebbe utile porre una stazione controllo di qualità a fine del processo standardizzato (si controlla l'ultimo semilavorato uguale per tutti). In tutti i casi si tengono in considerazione i tassi di difettosità che nel processo possono influenzare maggiormente le perdite di tempo per rilavorare i pezzi difettosi.

Quante ore macchina e ore uomo ho bisogno per poter rispondere alla domanda?

Analizzare le ore macchina attraverso una media pesata è un ragionamento non banale, proviamo però attraverso l'OEE, cerchiamo quindi di capire quante sono le ore necessarie supponendo di avere una domanda D1 e D2 (del prodotto 1 e 2) attraverso il tempo carico necessario:

$$T_{CAR OP.1} = \frac{DOM1 * TC_{t OP.1}}{Q_{1P1} * D_1 * E_{P1P1}} + \frac{DOM2 * TC_{t OP.2}}{Q_{1P2} * D_1 * E_{P1P2}}$$

Si tratta del tempo che ci occorre per rispondere alla domanda realizzando la lavorazione del prodotto 1 e del prodotto 2 nella macchina 1.

Per capire se è possibile rispettare la domanda si dovrà confrontare questo tempo carico con il tempo di apertura. Analizzando la differenza tra i due termini si deciderà di adottare la misura più opportuna come: interventi di efficientamento; lavoro straordinari; altre strategie di ottimizzazione volte a migliorare l'OEE e quindi il tempo carico necessario, per soddisfare la domanda [casi in cui il tempo carico è pochissimo o poco più grande del tempo di apertura].

(Tasso di qualità sintetico non serve, quindi della linea, si sarà misurati sulla capacità di rispondere alla domanda, quindi l'obiettivo è ridurre il tempo carico)

Supponiamo che O1 e O2 siano operazioni svolte dalla stessa macchina che chiamiamo M1-2. Calcolando il tempo carico, come somma del tempo carico calcolato prima dall'operazione 1 più quello calcolato per l'operazione 2 (in totale 4 termini) si verifica la capacità della macchina di soddisfare la domanda (interessa la macchina, non l'operazione).

Lo scopo finale è analizzare le performance della macchina, quindi non di per sé il tasso di qualità ma l'effetto che questo ha sulla macchina.

Ragionare invece con la media pesata nasconderebbe dei dettagli che risultano fondamentali per la soddisfazione della domanda, ad esempio non saprei se concentrare le attenzioni su un'operazione piuttosto che su un'altra, risulterei miope non capendo per quale operazione o per quale prodotto si sta perdendo più tempo (**quanto tempo mi serve, e su questa base prendere decisioni sul processo**).

Trovare il tempo carico più alto significa anche trovare il bottleneck della linea, ovvero la stazione di lavorazione che richiede più tempo per potere soddisfare la domanda. Conseguentemente essa avrà il GU più elevato.

ATTENZIONE: quanto visto fin ora (figure) rappresentano schemi del processo e non del layout.

Se una stessa macchina compie due operazioni successive è diverso dal punto di vista di layout, poiché viene rappresentata come una singola macchina nel Shop Floor, mentre da un punto di vista di processo viene rappresentata da due operazioni/stazioni diverse.

Sono le attività ad avere un tempo ciclo, una macchina che fa la stessa operazione però su due materiali diverse avrà tempi cicli diversi (*viviamo ormai in una realtà in cui la singola macchina può fare più operazioni diverse su materiali diversi*).

Negli esempi fatti finora si ragiona su tempi carico quindi in una strategia di tipo PUSH, per poter affrontare una strategia di tipo PULL bisogna far riferimento al tempo di attraversamento (P-Time e D-Time).

DISPONIBILITA' A GUASTO

La disponibilità a guasto di un componente/sistema riparabile è misurata dal rapporto tra il tempo cui il componente/sistema può funzionare ed il tempo totale per il quale è richiesto il servizio, questo fattore dipende dalla manutenzione e dai setup che deve effettuare una determinata macchina nel contesto produttivo (politica di manutenzione e politica di setup) **[situazione in cui la macchina è caricata ma non sta funzionando]**.

La manutenzione è un complesso di operazioni necessarie a conservare la conveniente funzionalità ed efficienza dell'impianto. Ogni singola macchina all'acquisto possiede dei dati di targa, quali:

- MTBF (mean time between failure), tempo che passa da un ripristino ad una nuova rottura;
- MTTR (mean time to repair), tempo impiegato per una riparazione.

Questi dati di targa sono teorici, tendenzialmente validi in condizioni ideali.

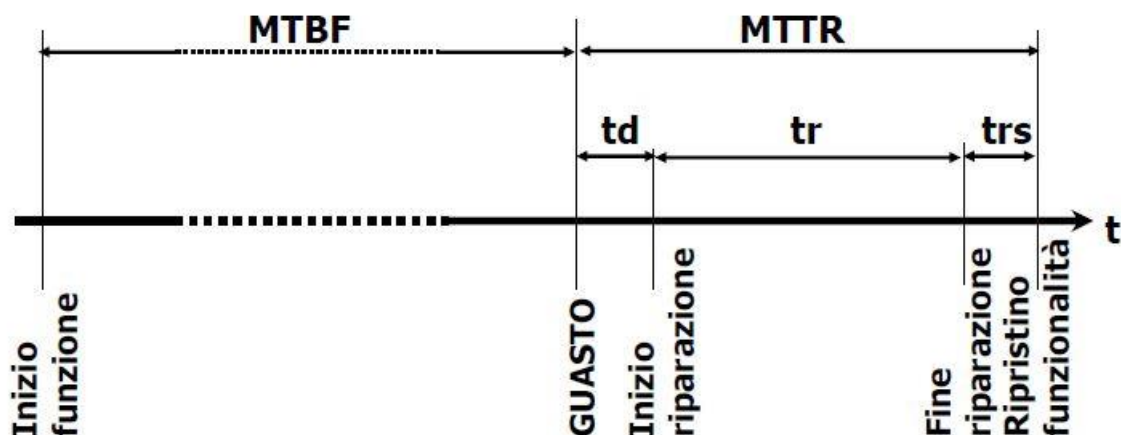
L'MTBF è una misura di affidabilità, che quindi aumenta con l'affidabilità della macchina, ovvero la probabilità che un sistema/componente opererà nel modo richiesto per un periodo di tempo determinato, se sarà impiegato nelle condizioni prestabilite.

L'MTTR è una misura della manutenibilità, ovvero la facilità, economia e sicurezza dell'assolvimento delle operazioni di manutenzione (**quanto è facile ripararlo quando si rompe**). Capita spesso che l'MTTR non sia fornito perché non tutti i guasti sono uguali, e di conseguenza non sono uguali neanche tutte le riparazioni.

Chiamiamo

$$A_i = \frac{MTBF}{MTBF + MTTR}$$

dove A_i sta per Availability ("disponibilità intrinseca di una macchina", ma non è ancora la disponibilità vera e propria, cioè quella che troviamo nell'OEE). Questo fattore A_i sarà il dato riferito alla macchina, le cose cambiano in base al contesto in cui la macchina è collocata non trovandosi più in condizioni ideali (tipo di linea, buffer ecc).



Un dato sicuramente più preciso, ma più difficile da ottenere è la Disponibilità Operativa A_o che utilizza dati non di targa come:

- l'MTBM (mean time before maintenance), che è l'MTBF esteso a qualsiasi tipo di fermo, non solo lo stop per guasto, ma anche per la manutenzione preventiva (perché ovviamente la macchina sta ferma non solo per i guasti). Ovviamente l'MTBM non è più una caratteristica intrinseca, ma dipende dalle politiche di manutenzione (se non c'è manutenzione preventiva coincidono). Si può facilmente intuire che (vedendo la formula a fondo pagina) ponendo al posto dell' MTBF l'MTBM la disponibilità diminuisce (perché mediamente risulta che l'MTBM è più basso dell'MTBF). In questo modo risulterebbe non molto vantaggioso fare manutenzione preventiva, visto che incide solamente sul numeratore, diminuendolo. In realtà influisce anche sul denominatore, in particolare sul MDT, il Mean Down Time, ovvero il tempo che mediamente trascorre dal momento in cui la macchina si è arrestata all'istante in cui è stata riavviata. Il MDT contiene al suo interno il tempo medio di riparazione (MTTR), ma ingloba anche altre tempistiche, come il tempo di diagnosi, il tempo di arrivo dell'operaio per riparare la macchina, il tempo per far venire eventuali pezzi di ricambio. Nel caso di una fermata programmata, tutti questi tempi al di fuori del MTTR rendono l'MDT un dato più sicuro in cui questi tempi possono essere drammaticamente ridotti, si potrebbe infatti già preparare gli strumenti necessari per la fermata, nonché farli già spedire i pezzi di ricambio necessari, o ancora arrestare la macchina nel momento in cui si ha l'arrivo del manutentore. Quindi se è vero che da una parte il numeratore della disponibilità diminuisce perché riduco l'intertempo tra una fermata e l'altra della macchina, dall'altra parte riduco anche il denominatore, in particolare l'MDT.
- l'MDT il tempo totale in cui un sistema resta inoperoso a causa di un guasto, che non è composto esclusivamente dal tempo attivo necessario per la riparazione (MTTR) ma anche del tempo di organizzazione dell'intervento e la manutenzione preventiva. Non è detto infatti che il manutentore sia specializzato per il tipo di problema sopraggiunto (non è detto che capisca subito il problema, non è detto che i pezzi di ricambio siano immediatamente disponibili). Tuttavia, una delle parti più critiche della manutenzione è il settaggio della macchina dopo la riparazione (*molte volte si evita la manutenzione preventiva pur di evitare di interrompere una macchina che sta lavorando in condizioni ottimali*). Solo nella manutenzione di tipo preventivo forse rispecchierò i tempi dell'MTTR perché quando interrompo la macchina sono già pronto per l'intervento. L'MDT (mean down time) viene introdotto come misura di quanto tempo mediamente la macchina sta giù per qualsiasi tipo di condizione.

$$A_o = \frac{MTBM}{MTBM + MDT}$$

Ma qual è l'orizzonte temporale su cui basare questi valori?

I guasti per natura sono di tipo diverso, possono avere frequenze alte o basse, con tempo di relativa riparazione più o meno lunghi. Qualsiasi stima quindi si farà, essa sarà soltanto una rozza approssimazione della realtà.

I principali tipi di guasti avvengono per:

- Usura;

- Numerosità di accensioni/spengimenti;
- Guasti accidentali;
- Legati al tempo solare (da quanto tempo sta ferma senza funzionare);
- Erosione;
- Deformazione per sbalzi termici.

Si è interessati ad una media di questi tempi, per quanto sporca possa risultare, per ottenere il dato orientativo della disponibilità. Per qualche macchina, particolarmente critica (collo di bottiglia) si andranno ad analizzare in dettaglio i singoli tempi che compongono l'MTTR o MDT.

Per esempio, una macchina con GU prossimo ad 1 e disponibilità “sporca”, dovrà essere curata con un'opportuna strategia di manutenzione e opportune scorte.

In generale si usa la disponibilità operativa, se non è possibile ricavarla ci si affida all'intrinseca.

Ma perché ancora non si arriva a definire la disponibilità vera e propria?

Si è parlato nell'OEE 2 di fermate di tipo interno ancora non considerate: lo starving e il blocking.

La disponibilità tiene conto delle fermate interne ed esterne di una macchina, quindi anche della disponibilità che altre macchine a lei connessa (rigida) possono avere. Si valuta quanto è lasca o rigida la connessione.

Se la connessione delle macchine risulta lasca allora la disponibilità dell'OEE coincide con quella operativa.

Per rendere lasca la connessione quanto devo far grande il buffer?

Ci si basa sul tempo di riparazione, questo mi dà l'ordine di grandezza (lower bound).

L'MTTR è una media che tiene conto dei tempi di riparazione della singola macchina, quindi non solo si dovrà considerare i singoli MTTR della macchina ma bisognerà confrontarli con quelli delle macchine tra cui il buffer è posto. I buffer in genere dovranno essere dimensionati come il massimo tra i massimi MTTR delle macchina a monte e a valle di questo (**il massimo dei massimi**).



Una volta che il buffer si svuota come e quando riesco a riempirlo?

La differenza tra tempo ciclo teorico e reale stima la capacità produttiva persa nel lungo periodo a causa delle perdite di tempo analizzate precedentemente, ciò non significa che la macchina viaggerà ad una velocità pari al tempo ciclo reale, ma tendenzialmente si avrà una velocità che sarà pari a quella teorica o decisa dall'Operation Manager (Throughput Rate). La macchina quindi accumula i prodotti in eccesso nel buffer, in modo che quando avviene un guasto o un'ulteriore perdita di tempo la macchina a valle sarà alimentata dallo stesso buffer interposto tra le due macchine.

Quindi alla domanda (unità/anno) si risponde con la capacità produttiva reale, ma quando tutto funziona in maniera corretta l'impianto marcia al valore teorico (*teoricamente quando il buffer è di nuovo pieno c'è la nuova rottura, ma trattandosi di medie [quindi valori irrealistici] potrei dover rallentare le macchine perché riesco a rispondere ampiamente alla domanda oppure avere un guasto mentre il buffer non è ancora pieno*).

Al problema del dimensionamento del buffer non si ha una soluzione (formula precisa, analitica) si utilizza il buon senso, ovvero si tiene conto del massimo tempo di riparazione delle macchine a monte e a valle e delle possibilità logistiche dell'impianto.

Per esempio se il buffer costasse troppo per il valore della merce prodotta si cerca di evitare i guasti attraverso opportune strategie accettando i limiti della connessione rigida.

Ricapitolando, nella connessione lasca si ha $A_o = D$, mentre in quella rigida (pensando a due macchine rigidamente connesse) si ha $D = A_{o1} \times A_{o2}$ (cioè devono funzionare entrambe).

La connessione rigida sicuramente è da preferirsi se la disponibilità operativa risulta essere molto buona. I buffer comportano sempre un costo, legato a spazio utilizzato e "messa in stock" del prodotto.

Pensiamo ad un caso: si ha una macchina con alta disponibilità, ma con un tempo di riparazione molto alto che ci obbligherebbe a tenere a scorta un'ernorme quantità di semilavorato per poter mantenere una produzione continua, anche se verrebbe usato raramente. In generale una macchina con disponibilità alta potrebbe nascondere un MTTR alto.

EFFICIENZA DELLE PRESTAZIONI

L'efficienza delle prestazioni riguarda le fermate (perdite di tempo) dovute a microfermate e rallentamenti (parte sconosciuta nell'OEE1). In questo caso possiamo affrontarli con i microbuffer.

Ma l' E_p si propaga?

Nella connessione lasca ogni macchina ha la sua efficienza delle prestazioni, questo significa che non è un dato che si propaga nella linea.

Nella connessione rigida potrebbe non propagarsi grazie ai microbuffer, che essendo micro non possono quasi essere definiti buffer (non generano costi elevati) e comunque sono tali da rendere lasca la connessione solo in circostanza di microfermata. Ma in altri casi il microbuffer potrebbe non essere abbastanza capiente o non può essere implementato, portando l'efficienza delle prestazioni a propagarsi.

Quindi nella connessione rigida in cui non è possibile utilizzare i microbuffer si avrà: $E_p = E_{p1} \times E_{p2}$ (nei nostri esercizi i microbuffer riusciranno sempre a rendere le macchine appena lasche da non far propagare microfermate).

IL PROBLEMA DEL SETUP

LOTTO ECONOMICO

Il problema dei setup sono tra i problemi più articolati che fin ora si è affrontato.

Nei setup rientrano l'insieme delle attività da svolgere prima dell'avvio di un lotto, di solito si parla di setup di riattrezzaggio, ma esso è utilizzato anche per pulizia, caricamento materiale e/o programmazione. Quindi, tendenzialmente si utilizzano i setup per riattrezzare un macchinario per un cambio di lotto, formato o prodotto.

Il setup sebbene indispensabile, danneggia l'OEE. Dover trovare del tempo "libero" all'interno del tempo carico in modo da non danneggiare la capacità produttiva dell'impianto risulta estremamente complesso.

A questo punto sorge spontanea una domanda ovvero: quando e quanti lotti fare?

La disponibilità per setup è una decisione operativa scelta dall'OM, presa in considerazione al tipo di processo e alla politica con cui si risponde alla domanda.

Parlando di setup obbligatoriamente dovremmo associarvi il concetto di PUSH (quantomeno fino al magazzino che potrà essere posto a valle della macchina) in cui la domanda viene soddisfatta dal magazzino e non dalla produzione.

Se ho una domanda di 100 pezzi al giorno, e per 7 giorni non produco quel prodotto, non significa che non soddisfo la domanda, ma dipenderà dalla capacità del magazzino.

Più aumenta il tempo in cui la merce è in magazzino e più aumentano i rischi a cui essa si espone.

Lotti grandi comportano pochi setup, lotti piccoli invece setup più frequenti. Più si fanno setup, più si perde capacità produttiva, danneggiando l'OEE.

Quindi o si aumentano i prodotti a scorta oppure vengono aumentati il numero dei setup (attenzione l'operazione di setup sostiene costi legati alla manodopera, attrezzatura e potrebbe anch'essa avere costi in termini di materiali sprecato).

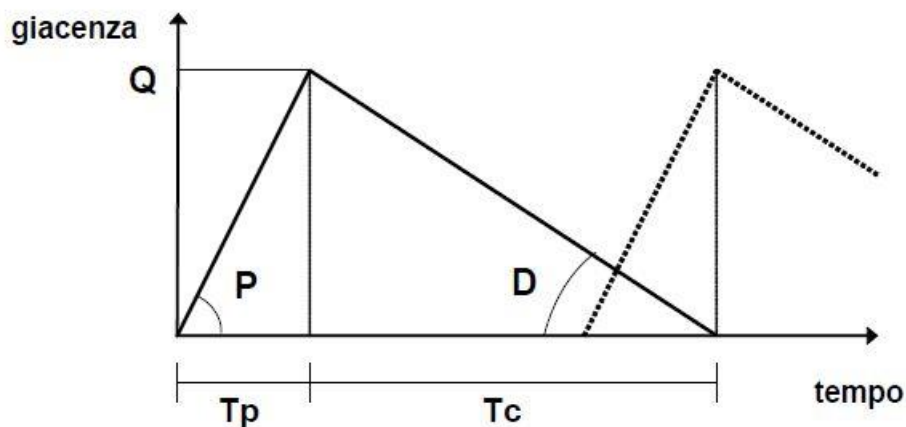
Occorre pertanto valutare il giusto trade-off tra i costi di setup e quelli di stoccaggio del materiale.

Solo in una politica di PULL "teso" a decidere quando effettuare i setup è direttamente la domanda.

CASO MONOPRODOTTO, CONSUMO DISGIUNTO DALLA PRODUZIONE

Linea in cui si produce un solo prodotto per lotti.

Ipotizziamo l'ammortamento di un macchinario a 10 anni, dopo tale periodo il macchinario è totalmente ripagato e non c'è un costo di impianto da caricare sul prezzo del prodotto. È quindi possibile che una linea produca un solo prodotto e ogni tanto stia ferma, non dando peso ai periodi idle (tipico di una visione tradizionale, attenzione però che la filosofia lean [OEE] non lo interpreta così. Indipendentemente dal tipo di macchinario i tempi idle sono uno spreco!).



Siamo in presenza di una produzione a intermittenza, con $P > D$.

Il tempo di anticipo con cui si dovrà riiniziare a produrre sarà proprio pari al tempo di produzione T_p . Per poter considerare i setup si deve considerare che una parte del tempo carico la si dedica a questi [considerando che dal tempo carico tolgo anche le fermate per manutenzione e guasto].

Il tempo carico, considerando che è composto dai tempi di produzione e i tempi per setup & manutenzione, deve riuscire ad essere inferiore al tempo di apertura, ciò significa che si dovrebbe riuscire a fare setup nel tempo idle, in cui la macchina è ferma. In alternativa si avrebbe che il tempo impiegato per eseguire i setup incide negativamente sul tempo di produzione, non potendo rientrare con il tempo carico nel tempo di apertura dell'impianto.

Se c'è un momento in cui la macchina sta ferma significa che il grado di utilizzazione è inferiore ad 1, quindi il tempo carico è minore uguale del tempo di apertura. Se esiste un tempo idle nonostante i setup e la manutenzione allora sicuramente il tempo carico è strettamente minore al tempo di apertura.

Ma quanti setup posso fare?

Considerando la macchina su cui fare setup si potrà avere una capacità produttiva reale (senza considerare i setup) maggiore, minore o uguale alla domanda. Se già la capacità produttiva reale è inferiore alla domanda significa che se inserisco i setup diminuisco la capacità produttiva reale della macchina, diminuendo ulteriormente la capacità dell'impianto di rispondere alla domanda.

Allora possiamo fare lotti continui?

Forse una produzione continua è la soluzione giusta, nei casi in cui la capacità produttiva reale è inferiore della domanda perché non posso perdere ulteriore domanda.

Se invece la capacità produttiva reale è maggiore della domanda?

In questo caso è possibile fare setup entro una certa misura senza perdere domanda. Quindi tutte le volte che si avrà un situazione simile con l'esigenza di dover fare setup si dovrà mantenere il vincolo di non dover perdere domanda. *Se si ha un turno di 8 ore senza poter fare straordinari e riesco a rispondere alla domanda con 7,5 ore allora ho mezz'ora per poter fare setup.*

Ma quanti si posso fare nel tempo?

Il tempo a disposizione per i setup è calcolato come differenza tra il tempo di apertura e il carico di lavoro:

$$T_{AP} - \frac{D * TC_t}{Q * Ep * Dg}$$

Questo primo dato inizia a dare delle informazioni per quanto riguarda una possibile misura della dimensione del lotto.

Ma come?

Supponiamo di ottenere ½ ora dall'equazione analizzata.

È proibitivo allora un setup di 2 ore?

No, si dovrà solo spalmarlo su 4 giorni: si dovrà quindi produrre consecutivamente per quattro giorni, ogni giorno portandomi avanti di mezz'ora.

Il riferimento che si è utilizzato fin ora, cioè giornaliero, è solo un nostro riferimento, non è un riferimento assoluto. Bisogna sempre ricordare che in ottica PUSH il magazzino disaccoppia le dinamiche (produco di più per avere tempi idle sulla macchina), quindi la massima frequenza di setup (o il minimo tempo tra un setup e un altro, o ancora, il lotto minimo) è pari al rapporto tra il tempo di setup e il tempo disponibile ogni giorno.

$$T_{\min(\text{setup})} = \frac{T_{\text{setup}}}{T_{\text{disp/g}}}$$

Se un setup dura 4 ore e ho solo ½ ora al giorno allora il lotto minimo sarà di 8 giorni, potrò farlo anche di 9, 10 o 11, ma non di 7.

Torniamo al caso iniziale: un solo prodotto sulla linea. Si deve scegliere quanto è grande il lotto di produzione ideale. Ovviamente se da una parte si ha un vincolo da rispettare, ovvero la domanda $T_{\min(\text{setup})}$, dall'altra aumentare la frequenza dei setup si paga con una giacenza media maggiore (Trade off dipendente dai costi).

Un indicatore per misurare il costo delle scorte è “H”, la cui unità di misura è €/u·anno. Questo parametro è quasi sempre frutto di un'approssimazione molto brutale. Infatti, misurare il costo delle scorte è davvero molto difficile.

La domanda quindi sorge spontanea; come si fa a calcolare quanto costa un solo prodotto in termini di stoccaggio?

Tendenzialmente non si ha un magazzino dedicato per ogni prodotto, esistono quindi costi condivisi tra più prodotti, come energia, personale, costo movimentazione, quanto occupa di spazio ecc. Per calcolare il valore di H effettivo si dovrà quindi trovare una metodologia che tenga conto di ciascuna voce di costo di immagazzinamento e la si ripartisca per ciascun prodotto. Avere questo complicatissimo sistema di conteggio implicherebbe un enorme costo di consulenza, arrivando ad un assurdo, ovvero pagare troppo informazioni che non hanno quel valore (costo

dell'immagazzinamento. Per le precedenti ragioni si utilizzerà l'approssimazione di H come il rapporto tra il costo totale e la giacenza media (riferendoci all'anno precedente).

Calcolo dell'Economic Batch Quantity:

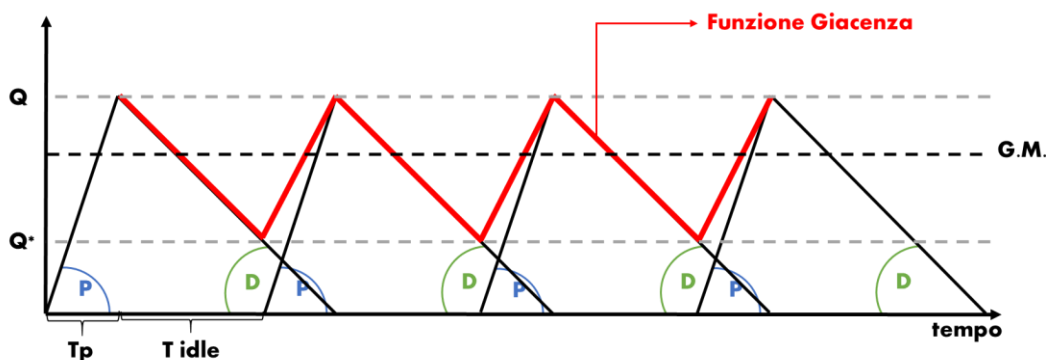
Supponiamo di conoscere:

- H (costo delle scorte); [€/unità*anno]
- Tasso di consumo d, costante nel tempo (di solito dovrebbe essere a "gradoni"), viene associato alla domanda media di consumo nell'unità di tempo considerata (nella figura di seguito sarà la pendenza D della funzione di consumo). [unità consumate/unità di tempo]
- Domanda complessiva annuale D.[unità annuale]
- Tasso di produzione p, anch'esso considerato costante nel tempo.[unità prodotte/ unità di tempo]
- S costo setup (costo fisso); [€/setup]
- CF costi fabbricazione (costo variabile).[€/unità]

Inoltre ci troviamo in una situazione in cui la CPr (senza setup) > D.

La quantità che si vorrà cercare sarà Q ovvero la grandezza del lotto, da cui si potrà risalire ad n_a numero di lotti annui e alla GM giacenza media generata da una determinata politica di setup.

$$\text{Costi Totali} = CF * D + n_a * S + H * GM = CF * D + \frac{D}{Q} * S + H * GM$$



La giacenza media GM è calcolata come media tra il valore massimo dell'inventario (raggiunto al termine del tempo di produzione Tp) e il valore minimo di inventario, denominato con Q^* , corrispondente alla quantità di prodotto stoccato nell'istante di tempo in cui la produzione viene riavviata per poter far risalire l'inventario a quota Q nel tempo di produzione a disposizione Tp :

$$GM = \frac{Q + Q^*}{2}$$

$$Q^* = d * Tp = d * \frac{Q}{p}$$

$$Q = Tp * p$$

$$GM = \frac{1}{2} \left(1 + \frac{d}{p} \right) Q$$

Quindi ottengo:

$$Costi\ Totali = CF * D + \frac{D}{Q} * S + \frac{1}{2} H * \left(1 + \frac{d}{p} \right) * Q$$

Se derivo rispetto a Q ed eguaglio a zero ottengo:

$$Q = \sqrt{\frac{2 * S * D}{H'}}$$

Dove:

$$H' = H * \left(1 + \frac{d}{p} \right)$$

Il Q ottenuto:

$$Q = \sqrt{\frac{2 * S * D}{H'}}$$

Viene chiamato EBQ, Economic Batch Quantity, cioè la quantità che minimizza i costi: sostituendo infatti la formula dell'EBQ alla formula del costo totale si ottiene che i costi di setup eguagliano quelli di stoccaggio, come confermato dal grafico della funzione di costo totale riportata nella *figura 7*. Si può notare dal grafico infatti che il minimo della curva dei costi totali corrisponde al punto di intersezione tra la curva dei costi di stoccaggio e quelli di setup, in altre parole nel punto in cui queste due curve si eguagliano.

Le ipotesi semplificative imposte nel calcolo dell'EBQ, nonché le ulteriori approssimazioni che necessariamente occorrerà fare nella stima/rilevazione dei parametri ipotizzati noti, fanno sì che l'EBQ calcolato sia principalmente un'indicazione, che andrà pesata e utilizzata in virtù delle stesse approssimazioni fatte ed eventuali analisi di sensibilità.

Inoltre, per il momento non si sono tenuti in considerazione le possibili durate dei setup, dovendo portare tutto alla stessa unità (ore/giorni/settimane), il tempo minimo di setup (*o meglio, frequenza; quel valore calcolato prima, pari ad 8 giorni*) è la somma del tempo di produzione, tempo idle e tempo setup

$$T_{setupmin} = T_{prod} + T_{id} + T_{set}$$

somma che a sua volta è pari al tempo di consumo (Tc).

Essendo

$$Q = Tc \times d$$

(dove Q è Q_{MAX}).

Facciamo un esempio: $T_{consumo} = 8$ giorni, $d=1000$ u/g allora il Lotto minimo = $8 \times 1000 = 8000$ unità.

Analizziamo le varie casistiche: se l'EBQ è maggiore o uguale del lotto minimo allora è facile, lavorerò un lotto pari all'EBQ (l'importante è che non scendo sotto il lotto minimo, non avrei tempo). Se invece l'EBQ è minore o uguale al lotto minimo allora lavorerò sul lotto minimo.

L'EBQ quindi è il miglior trade off in termini di costi, ciò non significa che dovrò utilizzare sempre l'EBQ come unica modalità di produzione per lotti non tenendo presente le esigenze dell'impianto in termini di tempo, che non mi permettono di scendere sotto il L_{MIN} (se si scendesse sotto il lotto minimo si avrebbe una produzione che non è capace a rispondere alla domanda).

La curva dei costi totali, presentata nella seguente figura, non è simmetrica intorno al valore di EBQ, ma più inclinata a sinistra e meno a destra rispetto al minimo della funzione. La minor pendenza presente a destra del minimo, riconducibile alla presenza nella formula dell'EBQ della radice quadrata, comporta un vantaggio in termini di robustezza nella scelta delle dimensioni del lotto: risulterà infatti che sarà sempre più conveniente approssimare l'EBQ per eccesso piuttosto che per difetto, in quanto in tale tratto ad una grande variazione della dimensione del lotto corrisponderà una piccola variazione in termini di costo.

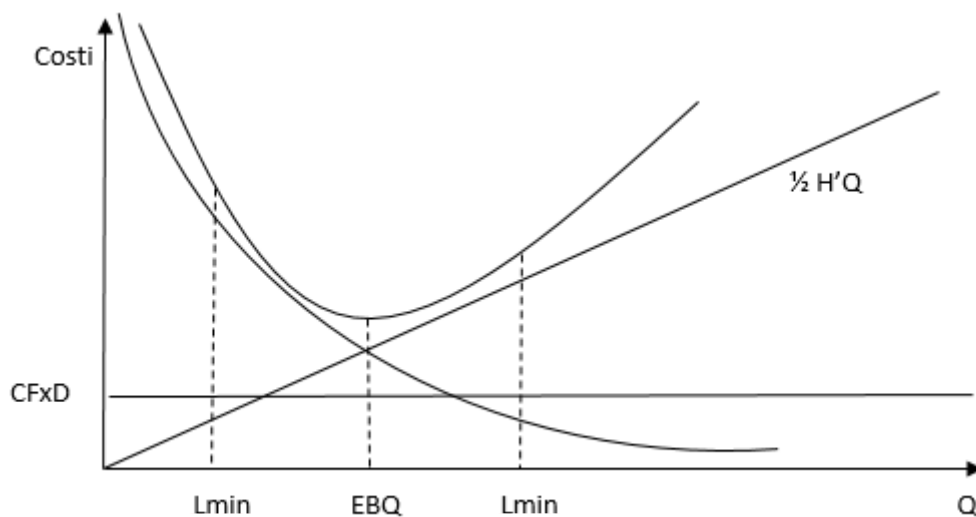
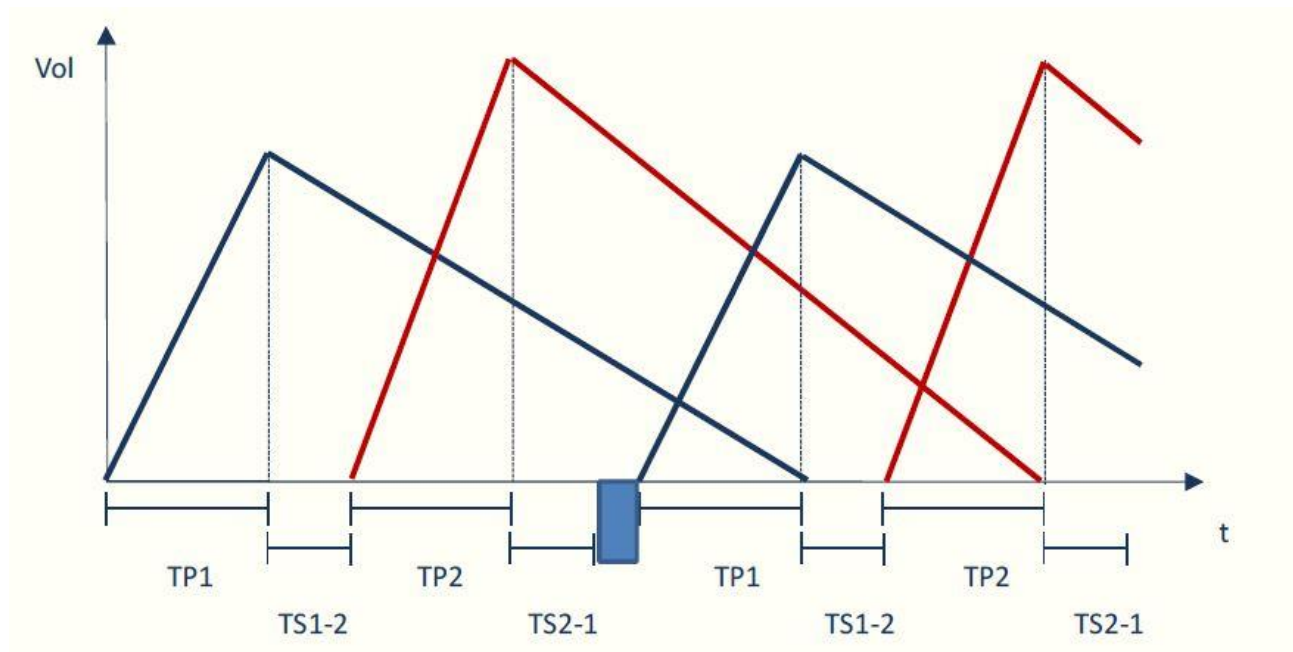


Figura 7. Funzione dei costi totali di Set-up

CASO PLURIPRODOTTO, CONSUMO DISGIUNTO DALLA PRODUZIONE

Vediamo cosa accade se si realizza sulla medesima risorsa produttiva più prodotti, in cui ovviamente le varie domande e produzioni possono non essere uguali.



Il tempo ciclo diventa:

$$\sum T_{p_i} + \sum T_{s_i} + \sum T_{p_i}$$

In particolare il tempo di consumo del prodotto 1 sarà:

$$Ti1 + Ts1-2 + Tp2 + Ti2 + Ts2-1 + Tp1$$

mentre il tempo di consumo del prodotto 2 sarà:

$$Ti2 + Ts2-1 + Tp1 + Ti1 + Ts1-2 + Tp2$$

I costi totali invece saranno:

$$Costi\ Totali = \sum CF_i * D_i + \sum \frac{D_i}{Q_i} * S_i + \sum \frac{1}{2} H'_i * Q_i$$

Questa funzione presenta n variabili pari al numero di prodotti da realizzare sulla medesima risorsa produttiva. Tuttavia tale numero si può ridurre osservando che per la ripetibilità del ciclo di setup/produzione, tutti i prodotti vanno realizzati in quantità tali da avere medesimo tempo di consumo TC_i , in quanto per ciascun prodotto il tempo che trascorre tra l'immissione in magazzino di un lotto e l'immissione del lotto successivo è proprio pari alla durata del ciclo di setup/produzione. Tutti i termini sono uguali, quindi la condizione è

$$TC_j = TC_i = \frac{Q_i}{D_i}$$

La condizione è quindi che il rapporto $\frac{Q_i}{D_i}$ è costante, cioè $\frac{Q_1}{D_1} = \frac{Q_2}{D_2}$.

Utilizzando il vincolo trovato si riescono a ridurre le variabili ad 1:

$$Q_i = Q_1 * \frac{D_i}{D_1} = Q_1 * a_i$$

La forma è la stessa di prima, ottenendo così:

$$EBQ_1 = \sqrt{\frac{2 \sum S_i * D_1}{\sum H'_i * a_i}}; EBQ_2 = EBQ_1 * \frac{D_2}{D_1}; L_{min} = \frac{\sum T_{setup}}{T_{disp}} * D$$

Infine occorre verificare (o imporre) la condizione di sufficienza di capacità produttiva sulla macchina in questione, ovvero che effettivamente:

$$T_{consumo} \geq \sum T p_i + \sum S_{i-j}$$

Nel caso in cui ciò non fosse verificato occorrerà imporre tale condizione aumentando la dimensione dei lotti (e dunque allontanandosi dal valore ottimo) di una m quantità pari a:

$$m = \frac{\sum S_{i-j}}{\frac{EBQ_1}{D_1} - \sum \frac{EBQ_i}{P_i}}$$

Con le nuove dimensioni dei lotti di produzione che diventeranno $Lotto_1 = EBQ_1 * m$, $Lotto_2 = EBQ_2 * m$ e così via.

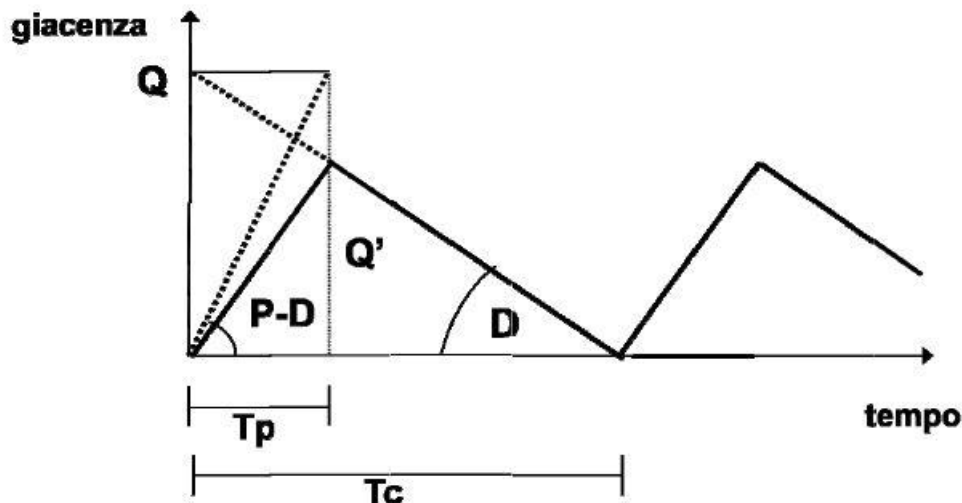
In un contesto legato al mantenimento della capacità produttiva della macchina, fare setup è un'operazione tanto critica quanto sarà critica la macchina su cui si applica. Ciò significa che se la capacità produttiva della linea fosse minore della domanda, ovvero che la macchina BottleNeck ha $Cp_{BN} < Domanda$, non implica che la macchina i , su cui si deve fare setup, non possa avere $Cp_i > Domanda > Cp_{BN}$.

In generale quindi si cerca di fare setup nel tempo libero cumulato, secondo una produzione di tipo PUSH, dove non si intacca la domanda. Si cerca poi il lotto ottimo da un punto di vista economico (trade off di costi) e lo si confronta con la dimensione minima del lotto che mi consente di non incidere sulla domanda.

CASO GENERICO, CONSUMO CONGIUNTO DALLA PRODUZIONE

Negli schemi visti finora avevamo assunto l'ipotesi di dover separare produzione e consumo: prima produco tutto e poi consumo tutto, ma potrei anche consumare e produrre contemporaneamente, cioè avere la possibilità di vendere un prodotto prima che sia terminata la produzione del suo lotto.

(Il caso visto sinora potrebbe essere quello in cui la separazione sia funzionale ad un CQ da effettuare prima di dare l'ok al lotto).



La produzione si ferma non appena viene raggiunto Q (cioè il lotto che volevo non l'effettiva giacenza).

La produzione del lotto successivo può avvenire in concomitanza con l'esaurimento della giacenza del lotto precedente (tanto posso consumare anche il primo pezzo prodotto – One Piece Flow).

Ovviamente la curva di crescita della giacenze in magazzino non sarà inclinata di P, bensì di (P-D), perché produco P, ma contemporaneamente consumo D.

La funzione dei costi totali generalizzata sarà:

$$Costi\ Totali = \sum CF_i * D_i + \sum \frac{D_i}{Q_i} * S_i + \sum H_i * GM_i$$

E detto Q il lotto di produzione, il valore di massima giacenza non sarà Q ma bensì:

$$GM = \frac{Q'}{2}, Q' = Tp(P - D), Q = Tp * P$$

Si ha

$$Q' = \frac{Q}{P} * (P - D) = Q \left(1 - \frac{D}{P}\right) \text{ e } GM = \frac{1}{2} Q \left(1 - \frac{D}{P}\right)$$

Di conseguenza si avrà che:

$$Costi\ Totali = \sum CF_i * D_i + \sum \frac{D_i}{Q_i} * S_i + \sum \frac{1}{2} H_i \left(1 - \frac{D_i}{P_i}\right) Q_i$$

dove chiamo

$$H''_i = H_i \left(1 - \frac{D_i}{P_i}\right) \text{ e } Q_i = Q_1 * \frac{D_i}{D_1} = Q_1 * a_i$$

Ottengo:

$$EBQ_1 = \sqrt{\frac{2 \sum S_i * D_1}{\sum H''_i * a_i}}$$

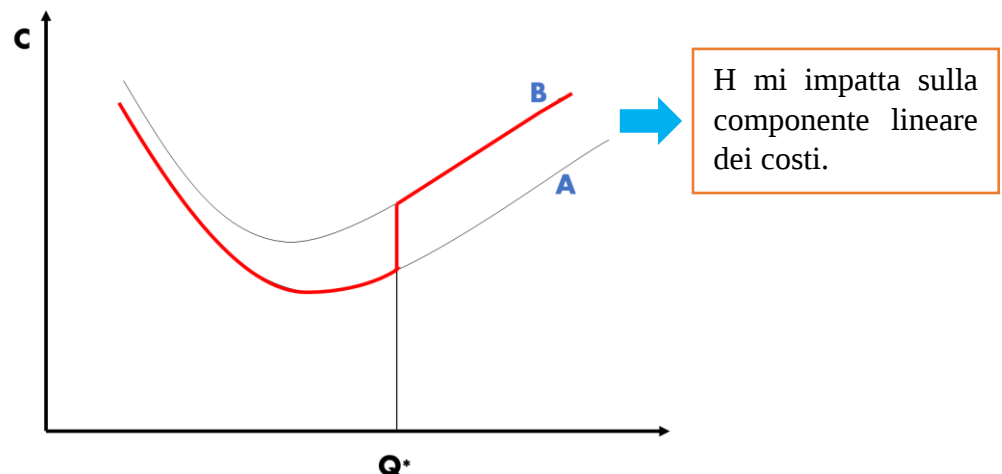
$$EBQ_1 = EBQ_1 * a_i$$

Un buon test di validità si può fare valutando H''_i , ovvero se questa quantità risulta minore di 0 (numero negativo) secondo tale formula risulterebbe impossibile una produzione per lotti, concettualmente si sta dicendo quindi che la Domanda è maggiore della Produzione, di conseguenza si crea un assurdo.

Ma cosa succede se H è costante a tratti, ovvero che dopo un certo livello di scorte il prezzo si dovesse alzare?

Supponiamo di avere quindi, $H = H_a$ se $Q \leq Q^*$ e $H = H_b$ se $Q > Q^*$

Per i costi si avrà:



Si dovrà calcolare il lotto economico con H_a , che esisterà solo per $Q \leq Q^*$ (quindi verifico se si trova nella sua regione di ammissibilità, ovvero rispetto al lotto minimo), proseguo nella ricerca, iterando lo stesso ragionamento con il lotto economico calcolato con H_b .

Calcolo anche il valore assunto dalla funzione nei punti di discontinuità, da sinistra e da destra.

Il valore minimo corrispondente ai lotti trovati sarà il lotto che minimizza i costi totali.

ATTENZIONE: Finora P e D sono stati identificati come la capacità produttiva e la domanda, in realtà raramente coincidono esse, sarebbe più appropriato quindi identificarle come le velocità di riempimento e svuotamento del magazzino.

La velocità di riempimento e di svuotamento torna ad essere una grandezza d'impatto, e in particolare bisogna porre attenzione al tasso di qualità, distinguendo due possibili casi in cui:

1. Se tra setup e buffer c'è un CQ allora la P che entra in magazzino va decurtata dei pezzi difettosi, quindi P va moltiplicato per il tasso di qualità.

2. Se invece il CQ è dopo il magazzino allora viene alimentato anche dai prodotti difettosi (Cp_{reale} senza contare il tasso di qualità = $Cp * D * Ep$).

La domanda invece rappresenta come viene svuotato il magazzino. Come prima bisogna ragionare sul come vengono presi i prodotti dal magazzino, quindi se dopo il magazzino c'è un CQ significherà che si tirano fuori anche i prodotti difettosi (anche se la domanda teoricamente mi chiede solo i pezzi buoni).

Nel lotto economico questi piccoli aspetti causano modifiche consistenti nella funzione dei costi totali.

GESTIONE DELLE SCORTE - Economic Order Quantity

Negli USA, il costo medio delle scorte equivale al 30-35% del loro valore: un'azienda che detenga scorte per 20 milioni di dollari spende per esse più di 6 milioni l'anno. Questi costi sono dovuti a obsolescenze, assicurazioni, costo opportunità e via discorrendo. Riducendo le scorte a 10 milioni di dollari, per esempio, l'azienda risparmierebbe oltre 3 milioni, il che ci porta al cuore della questione, cioè al fatto che un risparmio dovuto a una riduzione nelle scorte equivale a un aumento di profitto. Una scorta è costituita da qualsiasi articolo o risorsa impiegati in un'azienda. Un sistema di gestione delle scorte è l'insieme delle politiche e dei controlli che monitorano la quantità a magazzino e stabiliscono quale livello mantenere, quando reintegrarle, e quali dimensioni debbano avere gli ordini. Le scorte di produzione identificano quegli articoli che concorrono a produrre o a formare l'output di prodotto dell'azienda, e vengono classificate in materie prime, prodotti finiti, componenti, forniture varie e semilavorati in corso di lavorazione.

L'analisi delle scorte, in produzione e nei servizi, è principalmente diretta a specificare quando ordinare gli articoli e quante unità inserire in un ordine.

Tutte le aziende mantengono giacenze di magazzino per i seguenti fini:

1. Per garantire indipendenza tra le fasi;
2. Per far fronte a variazioni nella domanda di prodotto;
3. Per garantire flessibilità al piano di produzione;
4. Per cautelarsi contro le variazioni nei tempi di consegna delle materie prime;
5. Per sfruttare la dimensione ottimale dell'ordine di acquisto (emissione di un ordine comporta costi quindi maggiore è l'entità del singolo ordine, minore sarà il numero di ordini da emettere, inoltre, la spedizione di lotti più grandi è generalmente più economica: maggiore è la dimensione del lotto spedito, minore il costo unitario).

Nel prendere qualunque decisione che impatti sulle scorte, occorre considerare i seguenti costi:

1. Costi di giacenza (o di mantenimento);
2. Costi di setup (o cambio di produzione), produrre un prodotto differente implica azioni quali ottenere i materiali necessari, riattrezzare le macchine, compilare la documentazione richiesta, impiegare opportunamente tempo e materiali, sgombrare la rimanenza;
3. Costi di emissione dell'ordine;
4. Costi di mancanza (o di **Stock-out**), quando una scorta di articoli viene esaurita, un ordine che richieda tale articolo deve attendere fino al reintegro dello stock, oppure essere annullato. Fra il mantenimento di uno stock a copertura della domanda e i costi di stock-out esiste un trade-off, che è difficile risolvere in modo equilibrato, poiché non sempre è possibile stimare attentamente il margine di profitto perso, le ripercussioni dovute alla perdita di clienti, o le penali da pagarsi per il ritardo.

Esistono due tipi di sistemi di gestione delle scorte a periodo multiplo: i modelli a quantità d'ordine fissa (modelli EOQ) e i modelli a tempo fisso. I sistemi di gestione delle scorte multi-periodo sono progettati per assicurare la costante disponibilità di un articolo nel corso dell'anno. Solitamente

l'articolo viene riordinato più volte durante l'anno, in tempi e quantità ottimali, dettate dalla logica del sistema adottato.

L'elemento di distinzione è che i modelli a quantità fissa sono “attivati dall'evento” (devo monitorare le scorte quindi) e i modelli di tempo fisso sono “attivati dal tempo”.

Quindi, un ragionamento simile a quelli visti finora vale per l'approvvigionamento delle materie prime a quantità d'ordine fissa (e non solo delle materie prime).

I modelli di riordino a quantità fissa mirano a stabilire lo specifico punto (o livello) R in coincidenza del quale emettere un ordine e le dimensioni dell'ordine Q. Un ordine di dimensioni Q viene dunque emesso quando le scorte disponibili toccano il punto R. I modelli più semplici di questo genere funzionano quando tutti gli elementi sono noti con certezza: quando, per esempio, la domanda per unità di tempo di un prodotto è costante e lo stesso vale per i costi di setup e di giacenza.

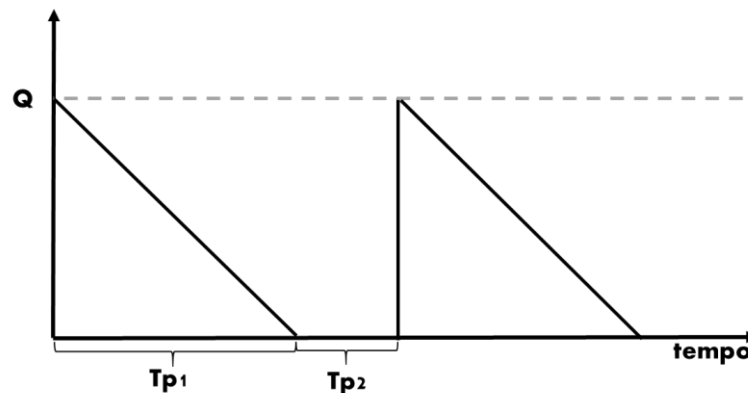
Inoltre, bisognerebbe capire che esistono possibili politiche di sconto (che però quasi mai dipendono dalla grandezza del singolo lotto d'acquisto ma soprattutto dai volumi annui).

Riepilogando il modello dell'Economic Order Quantity si regge sulle seguenti ipotesi:

Ipotesi dell'EOQ:

- Accumulo a lotti:
 $Tp = 0$, $p =$ tasso di approvvigionamento/produzione = infinito,
 ovvero la fornitura del lotto è istantanea e non frammentata, il lotto arriva in un'unica soluzione;
- Capacità del magazzino infinita (in questo modo qualunque sia il valore dell'EOQ si suppone di avere la capacità adatta per stoccarlo);
- Tempo di approvvigionamento costante (indipendente dunque dalla quantità di materie prime che mi occorrono, **ipotesi alquanto forte**, dal momento che il venditore potrebbe impiegare più tempo nel caso in cui i volumi richiesti siano di grandi dimensioni, *basti pensare che per consegnare interamente tutta il lotto nello stesso giorno dovrebbe mobilitare più camion, ciò richiede un maggior tempo di preparazione*);
- Beni non deperibili:
 il tempo di consumo altrimenti potrebbe essere superiore al tempo massimo di consumo per deperibilità;
- Domanda D interamente soddisfatta;
- Prezzo unitario del prodotto costante:
 anche questa ipotesi è alquanto forte, in quanto esclude la possibilità di poter beneficiare di sconti quantità (Il modello che tratta la dimensione del lotto in presenza di sconti quantità è riportato di seguito);
- Decumulo continuo (pseudo-irrealizzabile):
 E quindi tasso di consumo costante da parte della macchina a valle, attenzione quindi al caso in cui la macchina dovesse lavorare per batch e quindi dovesse avere per determinati periodi temporali tassi di consumo pari a 0.
 Quando allora viene meno l'ipotesi di decumulo continuo?
 Un esempio sta nella realizzazione di due prodotti finiti, uno dei quali non richiede la materia prima per la quale si sta calcolando l'EOQ. In tal caso si avrebbe comunque un

andamento a dente di sega, ma tra un dente ed un altro ci sarebbero dei tempi “morti” in cui non vi è richiesta della materia prima considerata. Ciò va a modificare il calcolo della giacenza media, che a sua volta intacca la formula dei costi totali e quindi il calcolo dell'EOQ.



- Tasso di consumo costante:

L'ipotesi di un tasso di consumo costante permette di definire la curva del consumo con un profilo lineare. Ciò è verificato nell'ipotesi in cui per tutto il periodo di consumo T_c la macchina consumi piccole quantità costanti del lotto di approvvigionamento. Se così non fosse supponiamo che il tempo di consumo sia di 10 giorni, e che il lotto di approvvigionamento sia di 100 unità. Se il tasso di consumo non fosse costante potresti supporre che tutto il lotto venga consumato l'ultimo giorno. La funzione di consumo non sarebbe così più lineare, bensì una funzione a gradino, in cui la giacenza avrebbe valore costante per i primi 9 giorni e nell'ultimo avrebbe valore 0. Ciò sconvolge il modello di Wilson¹, in quanto l'accuratezza del tasso medio di consumo nel periodo T_c è garantito dal teorema del limite centrale, che tuttavia è applicabile solo nel caso in cui il numero di prodotti consumati nell'unità di tempo siano in misura ridotta rispetto alla dimensione del lotto. (per completezza il teorema richiede anche l'indipendenza tra le quantità consumate dalla macchina in unità di tempo differenti, ovvero tra un tasso di consumo giornaliero e l'altro non vi sia alcuna correlazione, se assumiamo come unità temporale il giorno);

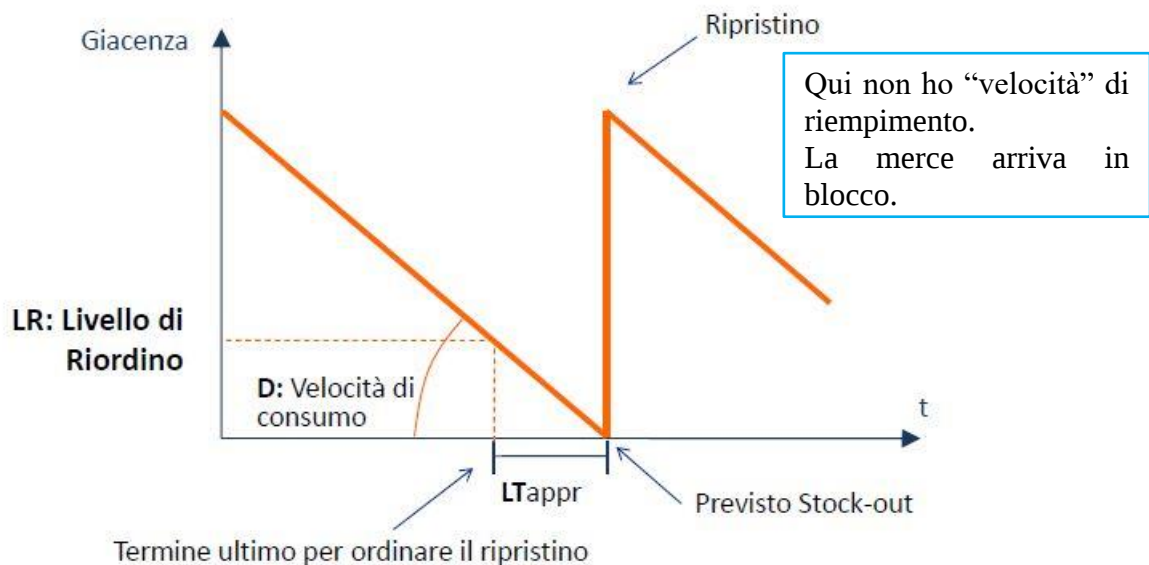
Quando valgono allora le ipotesi di decumulo costante (deterministico) e continuo (non si ferma mai)?

Per esempio, nel caso di approvvigionamento di materie prime, in cui ho un impianto fortemente stabile con una produzione continua: la prima macchina sarà quella che necessita sempre di quella materia prima, per la quale andrò a calcolare l'EOQ.

¹ Il modello EOQ (dall'inglese Economic Order Quantity) è stato proposto da F.W. Harris nel 1913, ma è attribuito principalmente a R. H. Wilson, che per primo studiò il caso.

Nella letteratura economica recente, tuttavia, è conosciuto come modello di HarrisWilson per la gestione delle scorte.

Ipotesi: Consumo costante



Quindi qualunque modello di gestione delle scorte intenda costruire, il primo passo è lo sviluppo di una relazione funzionale tra le variabili rilevanti in gioco e una misura di efficacia del risultato, partendo per esempio dai costi:

$$Costo Tot_{Anno} = Costo Acquisti_{Anno} + Costo Emissione Ordini_{Anno} + Costo Giacenza_{Anno}$$

$$Costi Tot = C_A * D + L * \frac{D}{Q} + H * GM$$

Dove Indichiamo con:

- D = Domanda (annua);
- C_A = Costo unitario della materia prima A;
- Q = quantità da ordinare. Il cui valore ideale è denominato lotto economico d'acquisto (EOQ);
- L = i costi per ordine (costi logistici, amministrativi, di trasporto) [costi fissi];
- H = Costo annuo di mantenimento per unità di scorta media;
- GM = Giacenza media, che per questo caso semplicistico “a dente di sega” sarà $Q/2$.

$$Costi totali = C_A * D + L * \frac{D}{Q} + H * \frac{Q}{2}$$

Quindi, per stabilire la quantità da ordinare per cui il costo totale risulta minimo bisogna calcolare la derivata rispetto alla quantità, ovvero il costo totale risulta minimizzato nel punto in cui la pendenza della curva è 0.

$$EOQ = \sqrt{\frac{2 * L * D}{H}}$$

Poiché questo modello semplice assume domanda e lead time costanti, non occorre alcuna scorta di sicurezza (safety stock). Quindi il punto di riordino è:

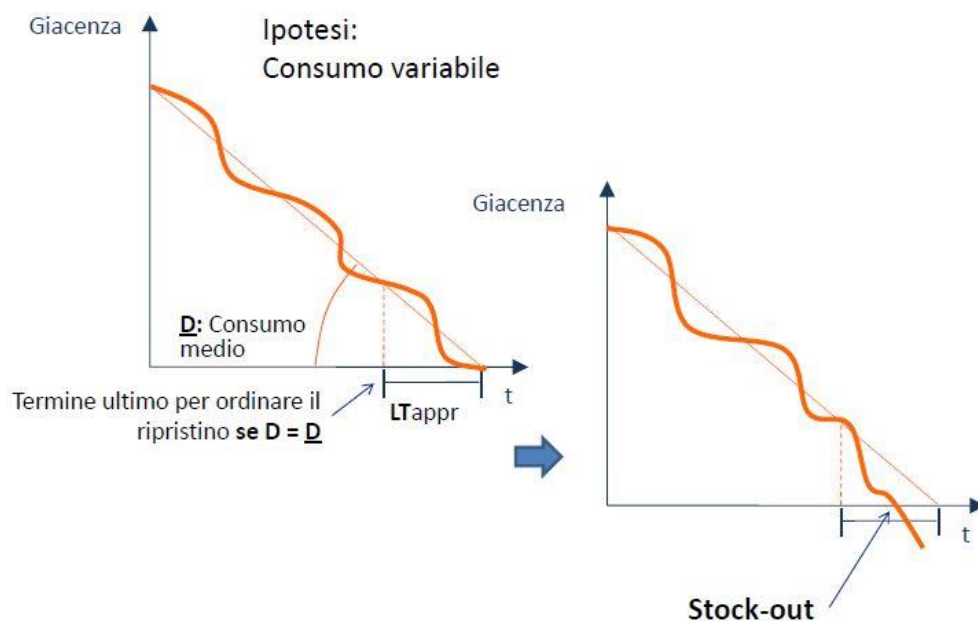
$$R = \bar{D} * Lead\ Time$$

Ovviamente, anche in questo caso semplicistico, possono esserci casi particolari, come il caso in cui si hanno due componenti che vengono ordinati insieme dallo stesso fornitore: nel caso dell'EOQ vengono trattate come se fosse un unico ordine, utilizzando il costo unitario pesato (sto ipotizzando che L non dipenda da Q), domanda pari alla somma delle domande, H potrebbe essere semplicemente la media pesata sulle quantità dei due diversi componenti, ovviamente l'EOQ sarà da ripartire proporzionalmente sui due componenti.

Il modello precedente assumeva una domanda come costante e nota (consumo lineare), ma non è sempre realistico. Potrebbe trattarsi sia di consumo di materia prima che di domanda di prodotto finito. Dunque, bisogna predisporre delle scorte di sicurezza che garantiscono un certo livello di protezione contro lo Stock Out.

Cosa succede quindi nel caso in cui non possiamo più supporre un consumo deterministico?

Come abbiamo visto è necessario, ai fini dell'applicazione del modello supporre un consumo costante, che nel caso del consumo variabile assumerà proprio il valore del consumo medio, con l'accortezza di rendere tale valore il più accurato possibile.



Come già ricordato precedentemente infatti, finché si ha una domanda più o meno approssimabile alla retta di consumo le approssimazioni viste fin ora possono avere un senso, ma se la domanda è un gradone dell'ordine di grandezza del tempo di consumo l'approssimazione comincia a diventare fortemente inadeguata (viene meno una delle ipotesi per l'applicazione del teorema del limite centrale).

Le giacenze medie tra i due rettangoli sarebbero fortemente diverse in questa ipotesi, quindi quando la variazione della domanda è a gradini, paragonabili ai tempi di consumo, la situazione si complica, per esempio, *se ho un prelievo ogni ora con un tempo di consumo di tre giorni, la situazione può*

essere ricondotta con grande approssimazione a quella di domanda costante, ma se ho un prelievo ogni 16 ore e un tempo di consumo di 24 ore allora la situazione si complica drasticamente. In questi casi si farà l'integrale della vera funzione dei costi [negli esercizi sarà solo chiesto di accorgerci di questa particolarità e procedere nel modo richiesto → caso della macchina a valle che lavora per lotti].

Nel caso di consumo aleatorio avremo dunque che la probabilità di soddisfare il consumo nell'unità di tempo con il consumo medio è pari al 50%, per definizione di media stessa.

A questo punto entrano in gioco le scorte di sicurezza: a cosa servono dunque le scorte di sicurezza? In generale a tutelarci dall'incertezza, sia che si tratti di incertezza derivante da un consumo non deterministico e/o per un tempo di approvvigionamento non deterministico.

Ad ogni modo il loro effetto è un aumento del livello di servizio [definito come numero di ordini evasi/numero totali di ordini ricevuti in un dato arco temporale].

superiore al 50%, livello di servizio offerto nel caso in cui mi affidassi solamente al consumo medio, senza alcuna scorta di sicurezza.

Le safety stock possono essere stabilite in base a molti e differenti criteri. Secondo un approccio diffuso, l'azienda può semplicemente decidere di stoccare scorte di sicurezza pari a una copertura di un certo numero di settimane. In realtà, è meglio servizi però di un approccio che consideri la variabilità della domanda.

Servirsi del criterio probabilistico per stabilire le scorte di sicurezza è molto pratico, supponiamo di avere una domanda, osservata in un definito periodo di tempo, normalmente distribuita, con una media e una deviazione standard. Per stabilire la probabilità di stockout nel periodo considerato, ci basta tracciare graficamente la distribuzione normale della domanda attesa e osservare dove si collochi sulla curva di quantità di scorta disponibile. Tale approccio però considera soltanto la probabilità di andare in stock out (esaurire le scorte), non il numero di unità mancanti.

Per esempio, *si attende una domanda di 100 unità per il prossimo mese e di conoscerne la deviazione standard di 20. Se si affronta il mese con sole 100 unità, sappiamo che la probabilità di stockout è del 50%.*

Per imprese che adottano questo criterio, è prassi stabilire la probabilità di non esaurire le scorte al 95%. Ciò significa munirsi di scorte di sicurezza pari a circa 1,64 deviazioni standard. Ciò non significa dover ordinare ad ogni ordine quella quantità, bensì continuare a ordinare scorte per i quantitativi medi mensili di domanda ogni volta, ma organizzando le consegne in modo da attenderci una giacenza a magazzino dell'1,64 deviazioni standard all'arrivo della merce ordinata.

Se invece si ha una domanda costante nel periodo e un lead time di approvvigionamento variabile, dovuto ad un fornitore poco affidabile?

In questa ipotesi, il pericolo quindi di stockout (sotto scorta) si manifesta soltanto nel lead time, che va dall'emissione dell'ordine all'acquisizione della merce. Come si è visto, l'entità della scorta di sicurezza dipende dal livello di servizio desiderato, per calcolare quindi la quantità ottimale da ordinare si può ragionare senza perdita di generalità con l'EOQ visto prima. A questo punto bisogna ricercare il punto di riordino, per coprire la domanda attesa nel lead time, più una sorta di sicurezza imposta nel livello di servizio desiderato.

Perciò, la differenza fondamentale tra un modello a quantità fissa con domanda nota e un modello a quantità fissa con domanda incerta risiede nella modalità di calcolo del punto di riordino e il calcolo delle safety stock che si vogliono assicurare.

$$R = \bar{D} * Lead-Time + z * \vartheta_{Lead-Time}$$

Dove:

- R = Punto di riordino in unità;
- $\bar{D}$ = Domanda giornaliera media;
- z = Numero di deviazione standard associato ad una specifica probabilità di servizio, percentile corrispondente al livello di servizio desiderato; il percentile si calcola supponendo una distribuzione normale in quanto si ritengono valide le ipotesi per applicare il teorema del limite centrale.
- $\vartheta_{Lead-Time}$ = Deviazione standard del consumo durante il lead time.

Il valore $z * \vartheta_{Lead-Time}$ definisce la quantità di scorte di sicurezza. In questo caso le scorte di sicurezza hanno avuto lo scopo di assorbire la sola variabilità del tempo di approvvigionamento, in quanto si è supposto deterministico il consumo.

Più in generale tuttavia vale la seguente formula per il calcolo delle scorte di sicurezza, chiamata anche Harley-Within. Formula che tiene conto di variazioni della domanda e del lead time di approvvigionamento:

$$SS = z * \sqrt{(\sigma_d^2 * LT + d^2 * \sigma_{LT}^2)}$$

Dove:

- z : percentile corrispondente al livello di servizio desiderato; il percentile si calcola supponendo una distribuzione normale in quanto si ritengono valide le ipotesi per applicare il teorema del limite centrale;
- LT : Lead Time di approvvigionamento delle materie prime;
- σ_d : Deviazione standard della domanda;
- σ_{LT} : Deviazione standard del Lead Time di approvvigionamento;
- d : Consumo medio (domanda media).

TECNICHE SMED (Single Minute Exchange of Dies)

Chi ha detto che il setup debba costare e durare una certa quantità?

Si può quindi lavorare su tempi e costi di un setup?

Sì,

il modo è lo **SMED**. Si tratta di uno slogan inventato da Shigeo Shingo, da cui talvolta prende il nome di riduzione drastica dei tempi di riattrezzaggio la parola infatti sta per “Single minute Exchange of die”. Ovviamente non va preso alla lettera, significa semplicemente che i tempi di setup devono essere sempre messi in discussione per migliorarne l’efficienza.

Un intervento SMED, che è un intervento radicale, deve ridurre drasticamente il tempo di setup.

Analizziamo un esempio: *se una persona buca 4 ruote dopo quanto riesce a sostituirle? Forse mezza giornata. Nella formula 1 le 4 ruote vengono cambiate in 10 secondi. Ovviamente c’è una differenza, differenza nella disponibilità di risorse (ad esempio ruote di scorta), risorse umane, allenamento, metodo e anche tecnologia, nella formula 1 esiste un unico bullone, ecc.*

A livello industriale i setup sono visti come un “tempo affondato”, cioè si accetta passivamente che il tempo sia quello, è molto spesso considerato un tempo non migliorabile. Le attività di setup infatti sono viste come delle attività molto complesse ed altamente specialistiche, in cui si nutre un profondo rispetto di chi le opera, inoltre vi era la convinzione che le attività fossero ogni volta diverse ciò rendeva impossibile la loro standardizzazione.

Shigeo vuole rompere questa visione suggerendo di andare a cercare sprechi e inefficienze anche laddove sembra che non ce ne siano: il manutentore che lavora sembra non perdere tempo, proprio perché lavora di continuo, senza fermarsi, ma non è detto che tutto quello che fa è all’insegna dell’efficienza, magari le stesse cose potrebbero essere fatte in maniera più rapida senza modificarne la qualità dell’intervento (*il manutentore non si rende conto di perdere del tempo per andare a prendere una chiave inglese, per esempio, che poteva benissimo prepararsi prima dell’intervento, perché dal suo punto di vista sta normalmente lavorando, rispondendo ad una necessità giunta in una fase non iniziale del processo*).

Ci sono tanti piccoli accorgimenti da poter adottare con tecniche più o meno varie, come per esempio *avere una macchina che per essere smontata anziché le viti abbia dei serraggi a blocco*, una semplice manutenibilità è fondamentale per migliorare le prestazioni della macchina in termini di OEE.

Molto spesso comprare una nuova macchina per migliorare le prestazioni in termini di setup, non ne vale la pena, quindi si cerca di modificare quella esistente attraverso interventi per un miglioramento continuo, interventi suggeriti e analizzati da persone che lavorano direttamente sulla macchina e.g. *se si hanno le viti devono essere poche ed uguali! avere una gamma vasta già pronta di guarnizioni di ricambio! È quindi importante investire su un’attrezzatura adeguata, essere sicuri che il manutentore quando interviene ha sempre tutti gli strumenti necessari attraverso una gestione a vista [Metodo Lean] anziché cercarli durante l’intervento prolungando inutilmente i tempi di fermo macchina.*

È una questione di metodo, tramite il quale con piccoli accorgimenti si riesce ad avere un'enorme differenza.

La prima volta che si applicano le tecniche SMED in un'azienda portano dal 50-70% di risparmio, solo attraverso piccoli accorgimenti verso gesti e atteggiamenti a cui nessuno dà peso.

E.g. se ho 5 persone che si occupano di manutenzione, è meglio farle lavorare su una stessa macchina anziché su 5 macchine diverse, non risparmio sulle risorse umane, ma riduco drasticamente il tempo in cui la singola macchina è ferma (vale per gli interventi programmati). Le ore uomo sono le stesse: con 5 persone sulla stessa macchina finisco in 12 minuti anziché in un'ora, dopodiché intervengono sulle altre 5 macchine (sempre un'ora a testa ma ogni macchina è stata ferma 12 min. invece che un'ora). Questo esempio è utile per dimostrare che molte volte nelle realtà aziendali si casca nella trappola del "si è sempre fatto così" o che ad interventi di miglioramento dei tempi di setup vi è bisogno di forti investimenti.

Una macchina moderna probabilmente avrà già incorporati molti accorgimenti, ma non l'aspetto organizzativo, che va comunque implementato dall'uomo. La gestione a vista è una soluzione semplice e veloce per migliorare l'organizzazione soprattutto delle fasi pre-setup, la gestione a vista è un approccio tradizionale nel Lean Thinking e viene utilizzato per monitoraggi, controllo e ispezioni in cui si hanno attività programmate con risorse dedicate.



Figura 8. Disponibilità attrezzature con sagome (Gestione a vista).

Ma operativamente come si applicano le tecniche SMED?

Le tecniche SMED devono avere il coinvolgimento dei manutentori, i quali sono orgogliosi e difficilmente accetteranno di essere affiancati. Si dovrà sfruttare delle tecniche di gruppi di lavoro per poter discutere insieme, coinvolgendoli nel farsi riprendere dalle videocamere così che, loro stessi possano capire dove possono migliorare e in che modo, contando il tempo sprecato e ragionare insieme a loro sulle possibili migliorie da effettuare.

È molto importante seguire questo percorso, perché se al manutentore viene chiesto la descrizione delle operazioni che compie, in automatico racconterà solo la successione "ideale", il suo cervello automaticamente rimuoverà tutte le inefficienze e quindi non gli darà importanza.

Come si applicano tali accorgimenti?

È importante distinguere i tempi in interni ed esterni all'interno di un processo, i tempi interni vanno fatti sempre con macchine spente mentre quelli esterni potrebbero essere fatti con macchine accese (e.g. cercare gli strumenti necessari al setup).

Il manutentore non distingue se è stata fatta un'operazione per un'ora a macchina ferma oppure mezz'ora a macchina ferma o viceversa mezz'ora a macchina accesa, dal suo punto di vista sta lavorando ed il resto è indifferente, mentre per noi risulta di fondamentale differenza.

Si deve quindi cercare di insegnare loro a distinguere i vari tempi di setup.

Un esempio di setup potrebbe essere il seguente:



Supponiamo che tutte le operazioni sopra riportate avvengano inizialmente a macchina ferma.

1° intervento: strutturare le tempistiche affinché il primo Tempo sia un T_{esterno} e quindi risulti quello più a sinistra (*un tempo a macchina accesa senza fermare la produzione come andare a prendere la strumentazione*).

2° intervento: cercare il più possibile di cambiare la sequenza di azioni per spingere i tempi esterni verso i margini (destro o sinistro) così da ripetere il 1° intervento.



3° intervento: velocizzare le operazioni che rimangono dentro il tempo in cui la macchina è ferma (accorgimenti quali parallelizzazioni, semplificazioni, pezzi di scorta, ecc...).



Ma su quale macchina applicare le tecniche SMED?

Sicuramente sulla macchina BottleNeck, soprattutto quando la capacità produttiva è al limite rispetto alla domanda.

Ma quando devo migliorare in questo miglioramento continuo? Quando mi dovrò fermare?

Succede che dopo un po' gli interventi, sempre più strutturati, cominciano a diventare costosi.

E.g della Colgate: *in ogni setup si doveva portare al lavaggio la tramoggia con la quale veniva alimentata la linea di dentifricio. Una soluzione interessante e sicuramente costosa è stata avere una tramoggia di riserva da mettere subito mentre si lava la precedente* (esternalizzo il tempo di lavaggio).

Inoltre bisogna cercare di eliminare regolazioni e centraggi velocizzando il riavvio del macchinario.

Quindi, dove fermarsi?

Quando si raggiunge la domanda oppure quando la macchina BottleNeck non è più quella su cui sto applicando le tecniche SMED, quindi si sposterà la mia attenzione su un'altra macchina.

ATTENZIONE: può succedere che dopo un po' di tempo gli interventi fatti diventino vani, quindi ci potremmo trovare a fare dei cicli di SMED, dove i successivi interventi divengono sempre più costosi. Shigeo diceva che l'ultima regola è non fare setup, anche se spesso diventa quasi la prima regola: **la filosofia giusta per poter ottimizzare i setup è farlo solo se necessario.**

Pensiamo agli enormi stampi dell'industria automobilistica: anziché avere uno stampo che ci dà 30 sportelli anteriori sinistra, non sarebbe meglio averne uno che combina gli sportelli in funzione di come è schedulata la produzione? (per esempio: sportelli punto-lancia y-panda, ecc).

Bisogna ricordare in generale che lo SMED è un intervento drastico, in cui se si porta a casa un saving del 10% risulta un fallimento.

La Capacità produttiva reale ($CPT * OEE$, nel senso dell'effettiva) non è solo migliorabile attraverso l'indicatore di prestazione OEE, ma anche chiaramente da quello sulla CPT .

Ma, allora cosa c'è nel tempo ciclo teorico?

Si può lavorare all'interno di esso per ottenere dei miglioramenti?

TEORIA DEI TEMPI E METODI (Method Time Measurements)

La teoria dei tempi e metodi considerano una generica fase del ciclo di lavorazione che presenti una componente svolta manualmente dall'operatore ed una componente effettuata dalla macchina.

Inoltre, i tempi non dipendono solo dalla persona, ma anche dalle condizioni in cui viene fatta lavorare, nacquero così i manuali in cui si studiano i tempi standard di ogni singola azione umana, così da ottenere tempi standard per ogni tipo di movimento. Oggi figlia di quegli studi è l'ergonomia, poiché studia quale soluzione riduce la fatica di un operaio così da avere una produttività maggiore con un minore sforzo.

E.g. *Si pensi all'operaio che deve lavorare su un'automobile che sta in alto (quindi braccia alzate) si chiama ergonomia il cambio logistico di collocamento della macchina posta di fronte all'operatore per agevolare le manovre operative.*



L'approccio della teoria dei tempi e metodi (cronometro in mano!) perde di umanità, non crea quel feeling con gli operai che permette di arrivare a raggiungere gli obiettivi di interesse.

Il vero aspetto interessante di questa teoria è che ha avuto il merito di entrare nel tempo ciclo teorico, mettendolo in discussione, andando a vedere quant'è il tempo impiegato dalla macchina e quanto dall'operatore.

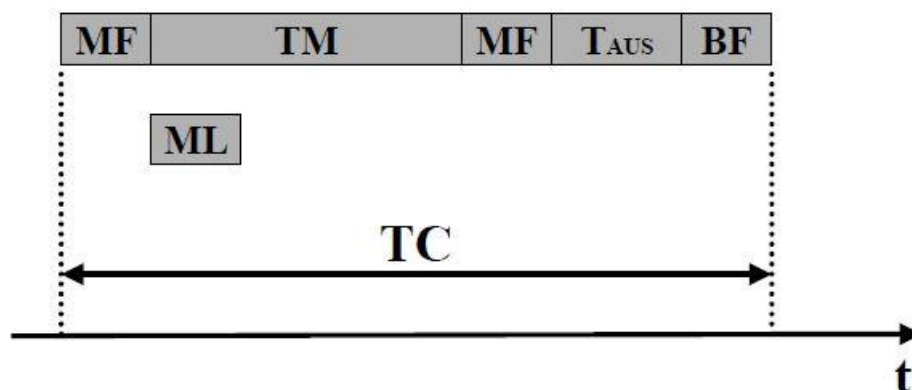
In generale, una lavorazione di questo tipo può suddividersi nelle seguenti operazioni che vengono compiute ad ogni ciclo di lavoro:

1. l'operatore preleva un pezzo da lavorare, lo carica nella posizione di lavoro, lo blocca e quindi avvia la macchina, il tempo impiegato per compiere tale operazione può essere contraddistinto con la sigla MF per indicare che è un tempo durante il quale l'operatore lavora mentre la macchina è ferma;
2. la macchina effettua la lavorazione, il tempo impiegato per compiere tale operazione è detto tempo macchina (TM) per indicare che è un tempo durante il quale la macchina lavora;

3. durante la lavorazione la macchina può subire interruzioni, per regolazioni centrature/spostamento del pezzo, ecc. In tal caso il tempo macchina (TM) non è continuo e viene frammezzato da operazioni a macchina ferma (MF);
4. durante la lavorazione la macchina, è possibile che servano degli interventi dell'operatore che non comportino la fermata della macchina, quali supervisione, input di comandi, ecc. Tali tempi sono "paralleli" al tempo macchina e sono detti macchina lavora (ML);
5. Terminando la lavorazione, l'operatore interrompe la macchina, sblocca il pezzo, lo rimuove, effettua un'operazione di controllo e quindi lo depone nel raccogliatore dei prodotti lavorati anche il tempo impiegato per compiere tale operazione può essere contraddistinto con la sigla MF per indicare che è un tempo durante il quale l'operatore lavora ma la macchina è ferma.

In sintesi, potremmo avere i seguenti tempi:

- Attività a macchina ferma (MF), macchina ferma per un'operazione che rientra nel TCt , come caricare il materiale per lavorare;
- Attività a macchina che lavora (ML), c'è l'operatore che fa operazioni, che sia condurre oppure fare operazioni che ha parallelizzato;
- Tempo macchina (TM), tempo in cui la macchina effettivamente funziona.



Oltre al tempo necessario per lo svolgimento di queste operazioni, che avvengono ad ogni ciclo di lavoro, per la determinazione del tempo ciclo è necessario tenere conto di una serie di tempi ausiliari per i quali la produzione viene interrotta, a causa di setup per cambio prodotto o riattrezzaggi del macchinario, interruzioni di funzionamento per guasti e interventi di manutenzione e microinterruzioni. Questa seconda categoria di tempi (raggruppati sotto un'unica voce chiamata tempi ausiliari), non ricorrendo ad ogni ciclo, deve essere ripartita sul numero di prodotti realizzati nel periodo di tempo nel quale si sono manifestati (T_{AUS}).

A questi tempi, è necessario sommare il tempo in cui il lavoro è interrotto a causa dei bisogni fisiologici dell'operatore, i quali possono assumersi, in condizioni di lavoro normali con lavori non particolarmente affaticanti, pari al 4% del tempo totale di lavoro ma possono aumentare per lavori particolarmente affaticanti o per condizioni di lavoro particolarmente sfavorevoli (BF).

Il tempo ciclo (TC) di una fase di un processo produttivo, inteso come l'intervallo di tempo che mediamente intercorre tra gli istanti di tempo in cui sono disponibili in output due prodotti (componenti processati in successione, può essere determinato quindi attraverso la relazione:

$$TC = MF + TM + T_{AUS} + BF$$

Proviamo ora, dopo aver introdotto in un quadro generale la teoria dei tempi e metodi, a fare delle piccole considerazioni riguardanti i tipi di operazioni su macchina che si potranno affrontare:

- In una macchina totalmente manuale, significherà che $ML = TM$;
- In una macchina totalmente automatizzata, si avrà solo TM ;
- Una macchina parzialmente automatizzata/manuale (caso prima descritto, generico) in cui solamente la prima fase richiede un intervento manuale dell'operatore.

Si Utilizzerà questo schema per capire dove andare ad ottimizzare. **Attenzione qui la terminologia tra visione OEE e visione MTM cambia come abbiamo potuto vedere.**

Secondo la teoria MTM ci si focalizza sul tempo ciclo reale, quindi come abbiamo visto, occorrono anche le “perdite”. Per l'esattezza quello visto finora (visione OEE) è un TCt/Ep perché anche nella visione MTM si considera un TCt mediato, e non un TC ideale (fattore che aggiusta il TC Ep).

Per la disponibilità la teoria MTM considera un “tempo ausiliario”, da spalmare su ogni tempo ciclo (includono, oltre alla disponibilità, anche i bisogni fisiologici).

E.g. *mi fermo 2 min ogni 60pz, allora $T_{AUS} = 120/60 = 2\text{sec/ciclo}$.*

In generale

$$T_{AUS} = \frac{T_{fermate}}{\text{numero cicli tra fermate}}$$

Ovvero il lasso di tempo medio in cui l'operatore (manuale) è indisponibile sulla macchina in ogni ciclo

Se nell'OEE si aveva $D_g = \frac{MTBF}{MTBF + MTTR}$, allora secondo la teoria MTM si avrà $T_{AUS} = \frac{MTTR}{MTBF/TC}$.

Da questa configurazione possiamo tirare fuori qualche indicatore come:

- La saturazione dell'operaio (come il GU), dove al numeratore c'è il Tempo Attivo (operaio utilizzato):

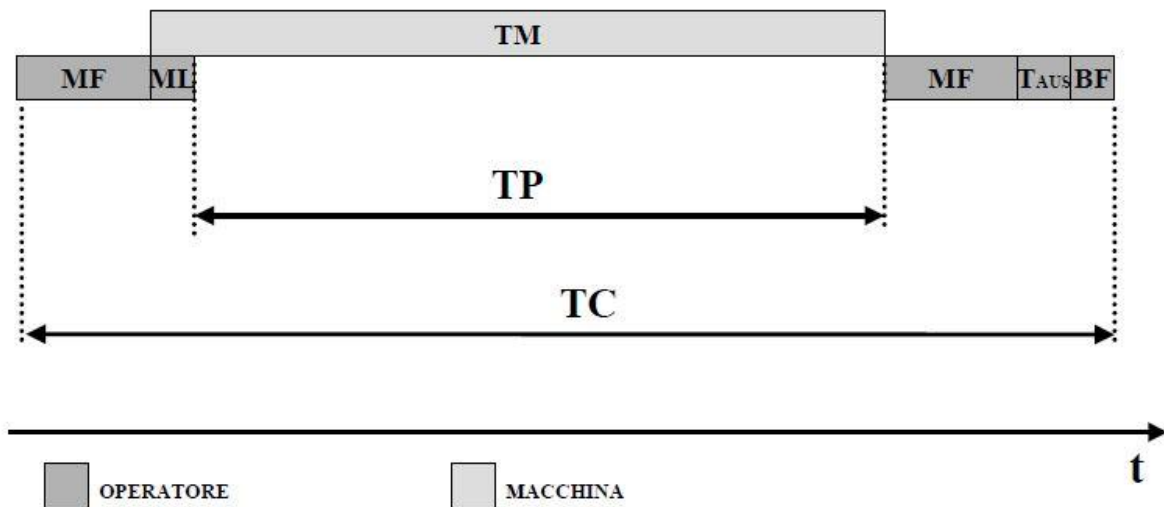
$$S = \frac{\sum MF + \sum ML + (T_{AUS}) + BF}{TC}$$

Dove il tempo ausiliario va considerato solo nel caso (o nella misura in cui) esso sia realizzato dall'operatore stesso (manutenzione autonoma). Il termine al numeratore rappresenta infatti il tempo in cui l'operatore è attivamente impegnato, ed è noto anche come tempo attivo;

- *Tempo attivo* = $\sum MF + \sum ML + (T_{AUS}) + BF$;
- *Tempo passivo* = $TC - TA$ (tempo in cui l'operaio è fermo);
- Indice di utilizzazione, ovvero quanto viene sfruttata la macchina per la produzione:

$$IU = \frac{\sum TM}{TC}.$$

Attenzione, nella teoria MTM non si rileva il tasso di qualità, perché non si guarda alla singola operazione.



Generalmente, come abbiamo sempre ragionato, distinguiamo due situazioni:

- $C_p < \text{Domanda}$ (lavoro per aumentare C_p).
- $C_p > \text{Domanda}$ (lavoro sui costi),

Partiamo dal primo caso, oltre a lavorare sull'OEE come abbiamo già visto, o su un tentativo di accelerare la macchina, c'è qualche altra possibilità?

Come posso fare da un punto di vista delle Operations per accorciare il TC?

Una possibile soluzione risulta essere la **Mascheratura**, ovvero facendo compiere all'operatore alcune delle operazioni mentre la macchina lavora e.g. *l'operazione di controllo del pezzo appena lavorato può essere condotta dopo aver messo in lavorazione il tempo successivo, in tal modo il tempo macchina ferma, e di conseguenza il tempo ciclo, si riducono, del tempo in cui l'operatore effettua l'operazione di controllo, che si è trasformato da macchina ferma a macchina lavora.*

Quindi, un MF iniziale potrebbe essere il caricamento del pezzo, da questo ci si chiede:

Questo caricamento si può fare a macchina che lavora?

Possiamo impegnarci per portare buona parte del MF in TM, parallelizzando al massimo la preparazione del pezzo successivo riducendo i MF al minimo indispensabile.

Questo naturalmente è solo un esempio, la mascheratura cambia di caso in caso. Può succedere che non si lavori sulla mascheratura per abitudine (perché magari $C_p > \text{Domanda}$) o per non saturare ancora oltre il lavoratore.

In generale per la mascheratura può dover servire intervenire sulla meccanica, per poter ottenere una riduzione del tempo ciclo attraverso l'introduzione di piccoli accorgimenti.

Mascheratura Essenziale:

- 1) Riduzione/Compattazione dei tempi a macchina ferma;
- 2) Compatto le operazioni in cui si necessità l'attività dell'operaio (ML) verso sinistra in modo da sfruttare la presenza dell'operaio data da MF;
- 3) Sfrutto l'automazione della macchina durante il periodo di TM. Voglio che la macchina svolga questa operazione senza troppe interruzioni e per farlo devo cercare di rimuovere o compattare a loro volta i tempi dei punti 1) e 2), (operazione utile per parallelizzare le attività mentre la macchina lavora).

Naturalmente per quanto riguarda i tempi ausiliari, si può ottenere una riduzione del tempo ciclo intervenendo:

- Sui tempi di setup cercando di ridurre il numero con un'opportuna programmazione della produzione e di ridurre il tempo attraverso opportune tecniche di riduzione dei tempi di setup (tecniche SMED);
- Sui guasti attraverso una manutenzione preventiva condotta al di fuori dei tempi destinati alla produzione e attraverso miglioramenti dell'organizzazione tesi a ridurre i tempi per il ripristino del macchinario in caso di guasto.

Se $C_p > \text{Domanda}$, si focalizzerà l'attenzione sulla riduzione dei costi.

Uno dei costi principali è la manodopera, riesco a risparmiare qualcosa? la mascheratura in questo caso va bene, produco in meno tempo e se ho flessibilità in termini di lavoratori posso risparmiare qualcosa (magari faccio fare altro ai lavoratori che sono liberi), ma non sempre è possibile.

Inoltre, è importante considerare che la tecnica di mascheratura è fortemente legata alla saturazione dell'operaio. Essa tende chiaramente ad aumentare visto l'aumentare dei compiti ad esso legati.

Un'altra soluzione percorribile è l'**abbinamento**.

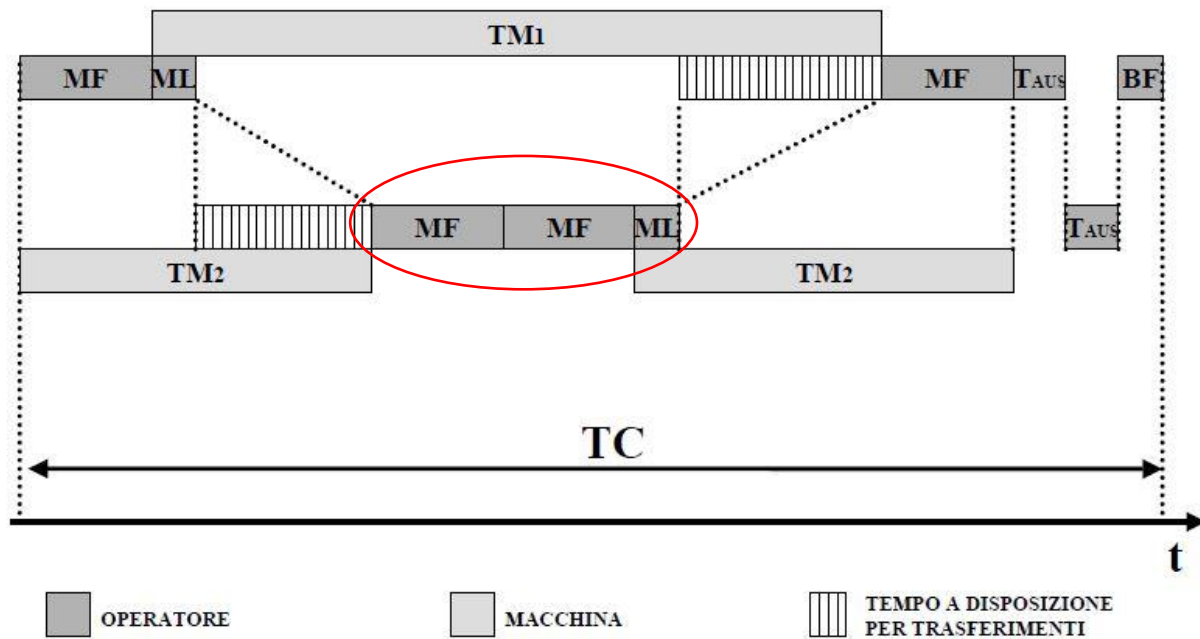
Allo scopo di ridurre l'incidenza del costo della manodopera sul costo del prodotto, si può valutare la possibilità e la convenienza di far condurre più macchine ad uno stesso operatore, soluzione che prende il nome di abbinamento.

Molto chiaro quindi, la principale differenza di questa tecnica rispetto alla mascheratura (in cui l'obiettivo principale è quello di ridurre il tempo ciclo).

Un operatore può difatti lavorare alternativamente o su una macchina o su un'altra, oppure può farlo contemporaneamente; questa seconda ipotesi è proprio l'abbinamento.

Se ho una rappresentazione di due cicli consecutivi come nell'immagine precedente, posso abbinare all'operatore che lavora su una macchina un'altra macchina, purché essa sia compatibile.

Nello schema, mostrato nella figura successiva, l'operatore fa tre operazioni consecutive (MF2, MF1, ML) e poi riposa (Tempo Passivo) cerchiato in rosso. Quindi una volta fatto ML sulla prima macchina l'operatore dovrebbe avere il tempo di spostarsi sulla seconda, fare MF2, poi MF1, poi ML (tutte sulla seconda) per poi tornare nella prima per riprendere a fare MF2.



Devo quindi trovare il giusto sfasamento: i due cicli devono essere simili.

Nel caso in cui le due macchine richiedono lo stesso periodo di attività, l'impegno dell'operatore su ciascuna di esse deve anche essere inferiore al 50% perché andranno considerati anche i tempi di spostamento fra le macchine.

Continuando con l'analisi dell'abbinamento, possiamo ridurre i tipi di tempi prima analizzati alla sola parte del tempo ciclo con frequenza unitaria (MF+TM), si possono analogamente definire:

$$S' = \frac{MF + ML}{MF + TM}; \quad TA' = MF + ML; \quad TP' = MF + TM - TA' = TM - ML$$

Vediamo alcune condizioni affinché si riesca a fare abbinamento:

- 1) Condizione temporale: $TP'_{macch1} \geq TA'_{macch2} + \sum TSpuntamento_i$
- 2) Il tempo passivo sulla prima macchina non deve essere troppo frazionato (ma sufficientemente compatto), vorrei riuscire ad infilare il TA' della seconda macchina in uno dei "pezzi" del TP' della prima. Ricordiamo che stiamo lavorando con tempi medi: non è opportuno incastrare "al secondo" proprio perché i tempi non sono sempre uguali. Come regola generale dobbiamo evitare ai lavoratori un alienante avanti-dietro dentro uno stesso ciclo. Un fattore da tener presente è anche il valore assoluto della durata del tempo ciclo: anche se matematicamente riuscirei a fare l'abbinamento di cicli di pochi secondi non posso chiedere ad un operatore di fare una serie di operazioni in un tempo strettissimo.
- 3) Le macchine devono essere sufficientemente automatizzate da poter essere lasciate senza supervisione, devono quindi prevedere almeno dei sistemi di arresto automatico a fine lavori, in caso di guasti, ecc.
- 4) Le macchine devono essere sufficientemente vicine tra loro, altrimenti deteriorerei la vita dell'operatore.

- 5) Se le macchine non sono uguali le operazioni devono essere compatibili: non posso pretendere che un operatore alzi 30 kg e poi effettui una saldatura di precisione per esempio.

Potrebbe capitare che alcune condizioni non sono verificate e ci venga chiesto comunque di fare abbinamento, per esempio se non fosse verificato $TP'_1 \geq TA'_2 + T_{Spost}$ o dovesse risultare troppo preciso.

Per prima cosa si potrebbe pensare a modificare il layout (oggi c'è flessibilità nella movimentazione delle macchine) per cercare di rendere più immediato il passaggio da una macchina all'altra; in un secondo momento, se non dovesse bastare, si dovrebbe pensare ad allungare il tempo macchina di una delle due, questa opzione è preferibile solo se:

1. Nessuna delle due macchine è una macchina BottleNeck per l'impianto;
2. Tale modifica non ne diminuisca la capacità produttiva rispetto alla domanda;
3. Le macchine devono lavorare in modo da fermarsi in automatico, altrimenti si potrebbero apportare alcune modifiche per cercare di renderle ad auto-fermante.

Non forzerò l'abbinamento su una macchina BottleNeck di una linea che non è troppo distante dalla domanda in termini di capacità produttiva (sia in eccesso, che, a maggior ragione, in difetto).

Bisogna sempre tener conto di come fino a questo momento ci si sia sempre riferiti alla teoria dei tempi e metodi che opera con tempi medi e non con distribuzioni temporali, si dovrà quindi evitare l'abbinamento sulla macchina BottleNeck pur quando, avendo un margine lo stesso non è troppo largo. Ovviamente sulle altre macchine posso osare fin tanto che non diventano collo di bottiglia (in generale, con questa tecnica, si tende ad aumentare il tempo ciclo).

Il problema principale dell'abbinamento è il tempo passivo eccessivamente frazionato.

Riusciamo ad automatizzare in maniera da ridurre alcuni frazionamenti del ML?

Bisogna cercare soluzioni analizzando la natura del processo, e.g. *avere tanti tempi ML frazionati potrebbe essere dovuto a controlli che si ripetono lungo l'operazione, in questo caso si cerca di ingegnerizzarli.*

Come tratto i tempi ausiliari?

I tempi ausiliari sono tempi particolari che dipendono sia dalla natura della macchina (affidabilità) ma anche dalle competenze dell'operatore, e.g. un operatore che è anche manutentore quando interviene su una macchina per guasto si fermerà di conseguenza anche l'altra. Risultando quasi al pari di una connessione rigida per quanto visto nell'OEE. I tempi ausiliari nell'esempio fatto si andranno a sommare.

Se invece l'operatore non è manutentore allora ogni macchina ha i propri tempi ausiliari.

ATTENZIONE: se aumenta la saturazione dell'operatore aumentano anche i bisogni fisiologici, per questo generalmente ci fermiamo ad una saturazione massima dell'80%.

Abbiamo quindi capito che l'abbinamento si usa per incrementare la produttività dell'azienda, ricordando che:

$$Produttività = \frac{Fatturato}{Addetto}$$

Quindi al pari di fatturato diminuiscono il numero di dipendenti di cui l'azienda ha bisogno.

Ma sarebbe davvero questa la mossa migliore? Cercare di aumentare la produttività aziendale licenziando e agendo quindi sul denominatore di quel rapporto? E se non dovessi licenziare che faccio fare alle risorse che avanzano?

Se licenzio passa l'idea che l'abbinamento faccia da apripista ai licenziamenti, perdendo feeling, collaborazione e rapporto umano con i dipendenti.

Bisogna sempre cercare di mettersi nei panni dell'operaio che arrivato ormai ad una certa età potrebbe trovare molte difficoltà nel mercato del lavoro una volta licenziato.

Cerco quindi di non licenziare. Ma allora come aumento la produttività?

Come detto prima l'abbinamento lo si usa per incrementare la produttività, diminuire il denominatore della relazione prima riportato è solo uno dei modi (sbagliato) di incrementare tale rapporto.

Aumentare il numeratore andando a cercare tutta una serie di attività di cui c'è bisogno nell'impianto significa fare insourcing, ovvero aumentare la produttività. Riportando dentro alcune attività che potrebbero incrementare il fatturato avendo anche maggior controllo dell'operazione.

Se mi viene detto che ogni 4 setup (ognuno di 2 ore) c'è un fermo di 2 ore allora per affrontare il problema imputo un pezzo di quel fermo a ogni setup, dunque è come se ogni setup durasse 2,5 ore. Un problema simile potrebbe avvenire se dopo ciascun setup la macchina riparte ad una velocità rallentata, e.g. per la prima ora la velocità è dimezzata: è come se ogni setup durasse mezz'ora in più. *Se dicessi che dopo il setup i primi 20 minuti di produzione vengono completamente scartati: è come se il setup durasse di più, perché per 20 minuti è come se non producessi, un problema saranno i costi imputati a tale setup, perché questa produzione scartata diventerà costo di setup (perché la causa è il setup).*

Imputare i costi è importante per capire le cause e trovare le giuste soluzioni (ma anche per il calcolo del lotto economico).

Andiamo ora ad affrontare un discorso analogo ma nel caso dei guasti, ad esempio: quando riparo la macchina per la prima ora la velocità è dimezzata, come si affrontano tali problemi?

Per prima cosa bisogna ricordare che i cicli di guasto sono composti dal periodo di MTTR e dal periodo di MTBF, di conseguenza se allungo l'MTTR, inevitabilmente devo diminuire l'MTBF del tempo opportuno. E.g. *se scarto la prima mezz'ora di produzione allora c'è una mezz'ora che va dall'MTBF all'MTTR.*

Se dico che la difettosità della prima ora dopo la riparazione è del 10% anziché del consueto 2%: allora il tempo speso per fare l'8% dei pezzi di un'ora (capacità produttiva) è tempo che passa dall'MTBF all'MTTR, oltre che un costo da associare alla manutenzione (oppure l'8% di un'ora), materie prime e lavorazioni.

(10/04/2018) Lezione sulla manutenzione tenuta dal prof. Vito Introna

(16/04/2018) Esercitazione

(17/04/2018) Lezione sulla Sicurezza sul lavoro

(19/04/2018) Lezione sulla DFMA

(23/04/2018) Esercitazione

(03/05/2018) Lezione sulla manutenzione 2 tenuta dal prof. Vito Introna