

Chapter 9 – Stops

Have you ever noticed that google maps – or your equivalent – will always tell you that you will get to a destination earlier than you actually do? It is probably reasonably accurate except it does not account for your stops.

In manufacturing stops make the difference between meeting your production targets, meeting your targets with excessive overtime or not meeting your targets at all.

In this chapter we will consider various types of stops and how to model them, We will also consider how these, along with quality problems determine the overall effectiveness of your equipment.

A stop can be defined as an interruption of the normal processes of a machine where the machine. A stop is not quite the same as the waste process elements we looked at in chapters 3 and 4. If the thing being done is counted in the machine's cycle time it is not a stop.

Broadly, stops can be characterised as planned or unplanned.

Some planned stops :

- Works *shut down* for holidays. Many factories take a couple of weeks off at some point in the year.
- *Machine servicing*. Hopefully this is timed to coincide with the shutdown. However, this is not always possible. Servicing may be required many times a year
- Manually operated machines – coffee and lunch breaks, operator training
- Machining lines – worn tool replacements.
- Shift patterns. A 1 shift system technically has 2 long stops every day. This is a matter of policy and how you treat these stops is a matter for discussion in modelling. In addition, there is frequently a stop for shift changeover.

Some unplanned stops :

- Breakdowns.
- *Blocking* – where the machine's downstream is full and not moving
- *Starving* – where the machine cannot work because there is no work to do
- Power cuts
- Disasters – floods, earthquakes, wars
- Manual machines – comfort breaks.

Unplanned stops are a form of uncertainty and you may remember from chapter 3 that a *risk* is an uncertainty that matters. In the same way it is suggested that when you model a

machine you should be looking at stops which might significantly reduce production or significantly disrupt production flow.

The effect of a stop in *reducing* overall production depends on :

1. The interval from the end of one stop to the beginning of the next and the length of each stop. In particular we are looking at :

$$\frac{\text{stop time}}{\text{stop time} + \text{interval}}$$

2. Whether the stop generally occurs when another stop is already active. So for example, if you can do most of your servicing in the holiday shutdown or at night in a single shift system. production will not be affected.

The more you understand what makes a good production system, the more you will understand how the idea of an even flow is important. Why – because uneven flow is often the root cause of two important unplanned stops – blocking and starving. It also leads to a requirement for buffer stocks if these problems are to be avoided. It is a good idea to consider how much a stop will disrupt the flow of a production line. Here , the length of the stop becomes crucially important..

OEE - Overall equipment effectiveness.

You may remember in chapter 2 we talked about yield and scrap rate We lose parts to scrap but we also lose parts if we lose production time to stoppages. There are two definitions of OEE.

Definition 1

$$OEE = \frac{\text{Actual production with stoppages and scrap}}{\text{theoretical maximum production}}$$

Definition 2

$$OEE = \text{performance} * \text{planned utilisation} * \text{unplanned utilisation} * \text{yield}$$

Where :

$$\text{Performance} = \frac{\text{cycle time}}{\text{takt time}}$$

$$\text{Planned utilisation} = 1 - \frac{\text{planned stop times in a period}}{\text{period}}$$

$$\text{Unplanned utilisation} = 1 - \frac{\text{unplanned stop times in a period}}{\text{period}}$$

And :

$$\text{yield} = 1 - \frac{\text{scrap parts}}{\text{scrap parts} + \text{good parts}}$$

The first method of defining OEE is actually more useful for modelling – but the second may be a good way to focus ideas for productivity improvement. If you want to improve OEE, you need to either :

- Increase performance by balancing your process to the takt time
- Increase planned utilisation by reducing planned stops
- Increase unplanned utilisation by reducing unplanned stops
- Increase the yield by reducing scrappage.

Modelling stoppages

Modelling planned stoppages

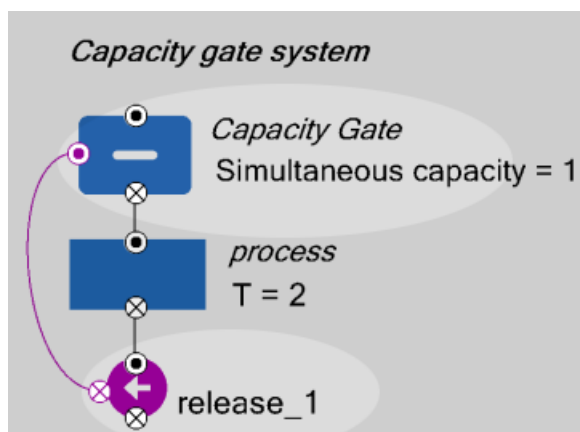
We will start with a very simple example and move on from there.

Open *Chapter9Ex1*:

Animate this for a while. It is the simplest of manufacturing systems. Production is pushed with a takt time of 150 seconds. The machine has a cycle time of 120 seconds.

In our first exercise we are going to model a simple planned stoppage, eg a 15m tea break at 10am (assume production starts at 8 am).

To do this we will simply stop the source for this stop. This is done in Aliquis by means of *messages*. All Aliquis blocks can send or receive these messages to change their behaviour or trigger an action. You have already seen messages in action with the capacity gate...



When a token reaches the purple message sender – it sends the message “release 1” to its output message port. This message is routed to the gate which releases the next token.

Actually Aliquis supports a number of block messages :

Message	Used on block types	Action
Stop	Source /delay/workstation	In a source issuing a stream - the stream will stop until a restart is received. In a delay/workstation , the time spent in the block will not be counted while the stop is active (this ceases on a restart)
Restart	Source/delay/ workstation	
open	Gate	Opens a closed gate – any tokens in the gate are ejected. Subsequent tokens are allowed through immediately
close	Gate	Closes an open gate – self explanatory
Release1	Gate	In a gate containing 1 or more tokens, one is released immediately. If there are no tokens in the gate, the subsequent arrival will be let through
Release quantity	Gate	As above – but a quantity – usually more than 1. If there are not enough tokens in the block then subsequent arrivals will be allowed through
Release quantity or flush	Gate	As above – but if there are not enough, subsequent tokens will not be let through
Flush then close	Gate	Release all the tokens currently in the gate. Subsequent tokens are not let through in any circumstances
Reject	Delays/workstations	If a token is in a delay or workstation it will be rejected through the secondary destination (usually another port)
Issue	Source	The source will issue the quantity of tokens set up immediately.
Simulation start	Solving subsystem	Used for nested solving subsystems (we will cover this later)

In this case, we will use *stop* and *restart*.

Lets build a new set of blocks to model the tea break.

First you will need to build a connected general source and sink.

My workstream

Process map/model



Basic process map

Process time modelling



Standard Process times



Variable process times

Result:

Solve Finish time = 28920

Start



Stop

Add two delays with times

My workstream

Process map/model

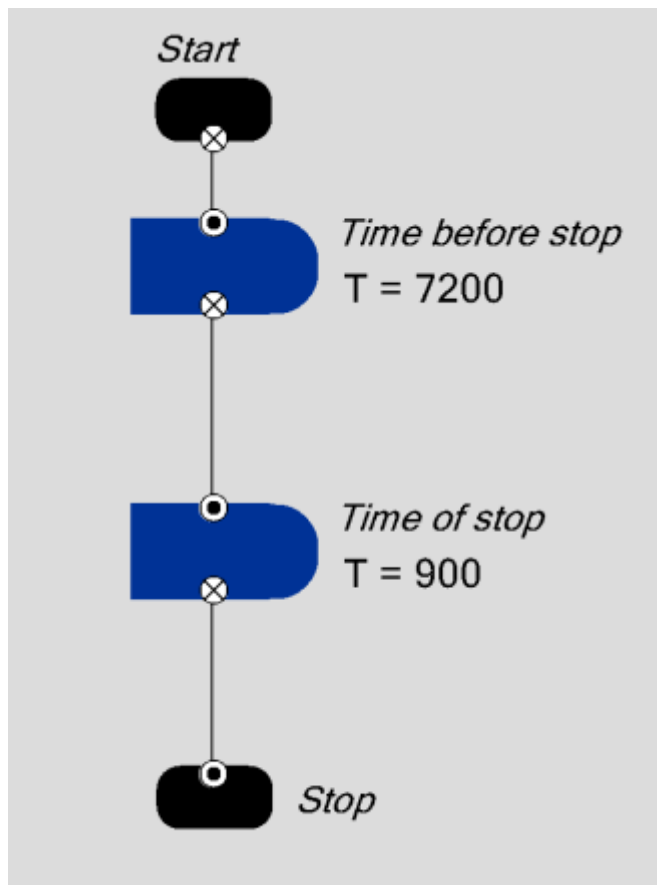
  *Basic process map*

Process time modelling

  *Standard Process times*

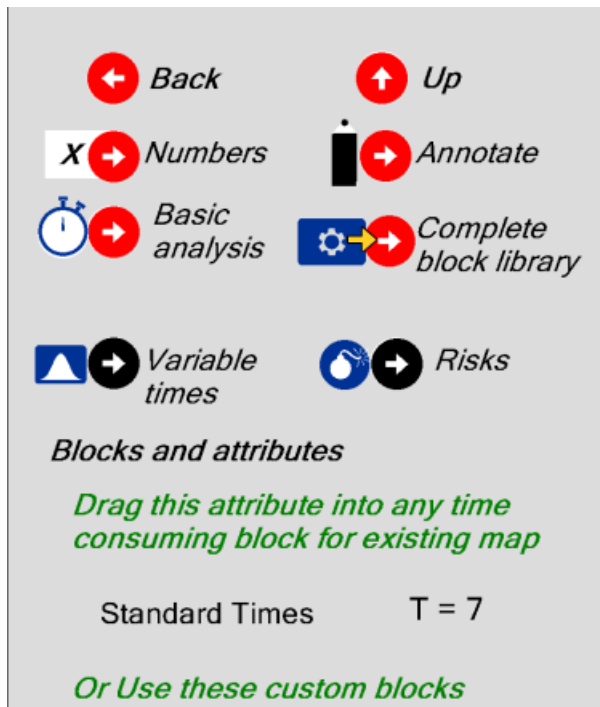
  *Variable process times*

Result :

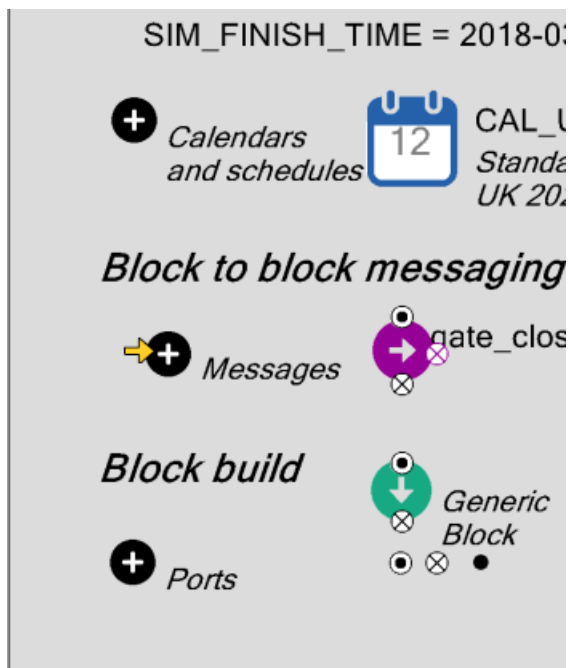


We will now find the blocks to create the stop and restart message.

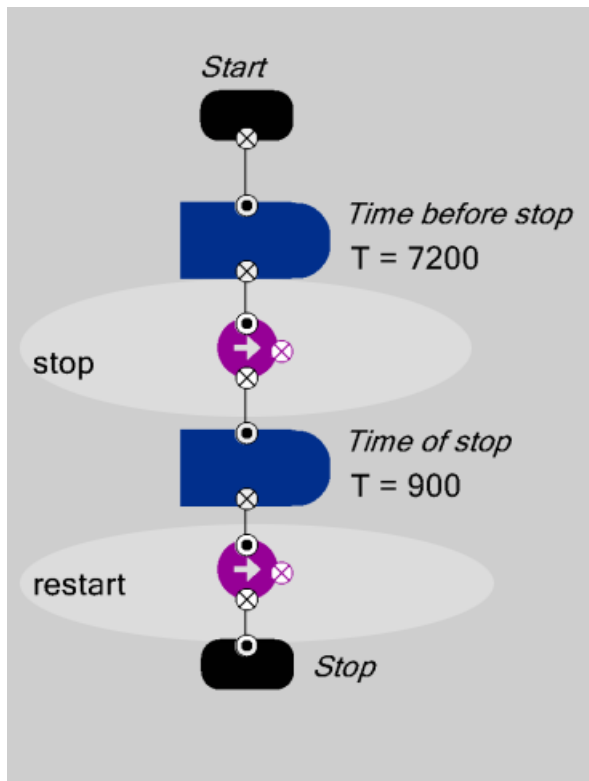
These can be found in the “complete block library”



Scroll down to messages :



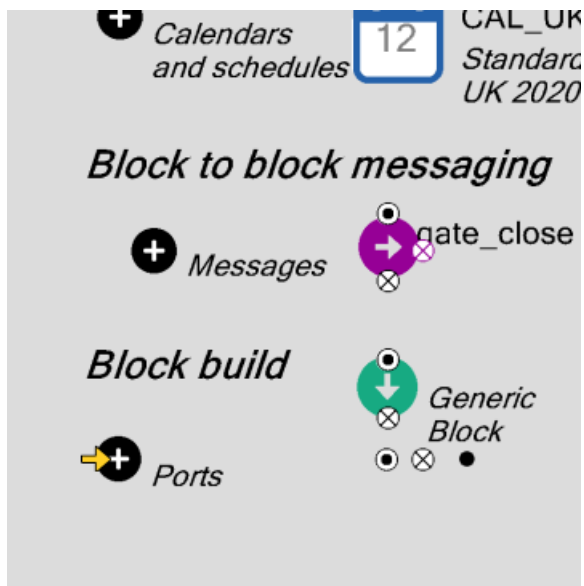
Find and add these blocks :



We now have the right messages firing at the right time – but how do we link them to the source?

The problem is : The source needs a message port.

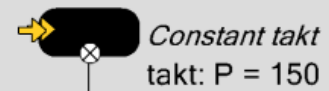
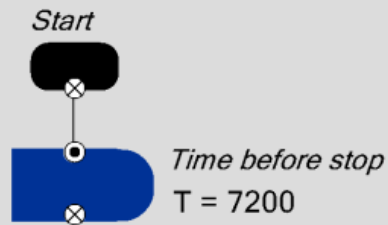
Go back up one level and find “Ports”



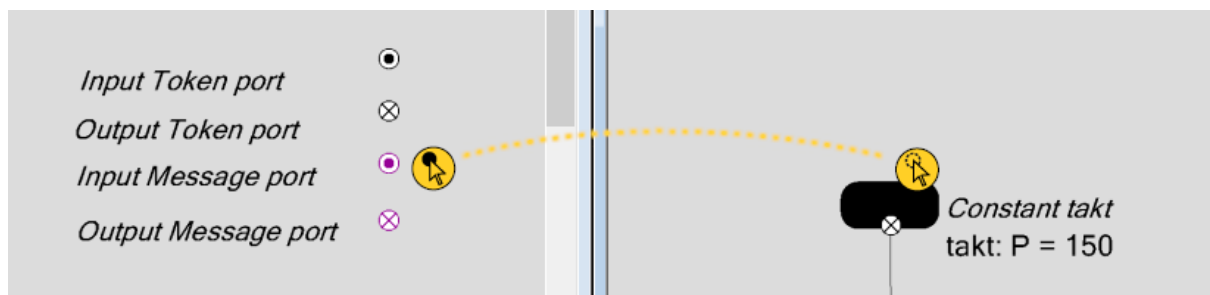
In your model Double click your source block :

Chapter 9 Ex 1

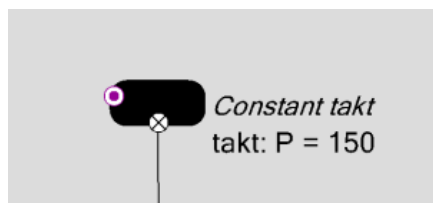
Solve Finish time = 28920



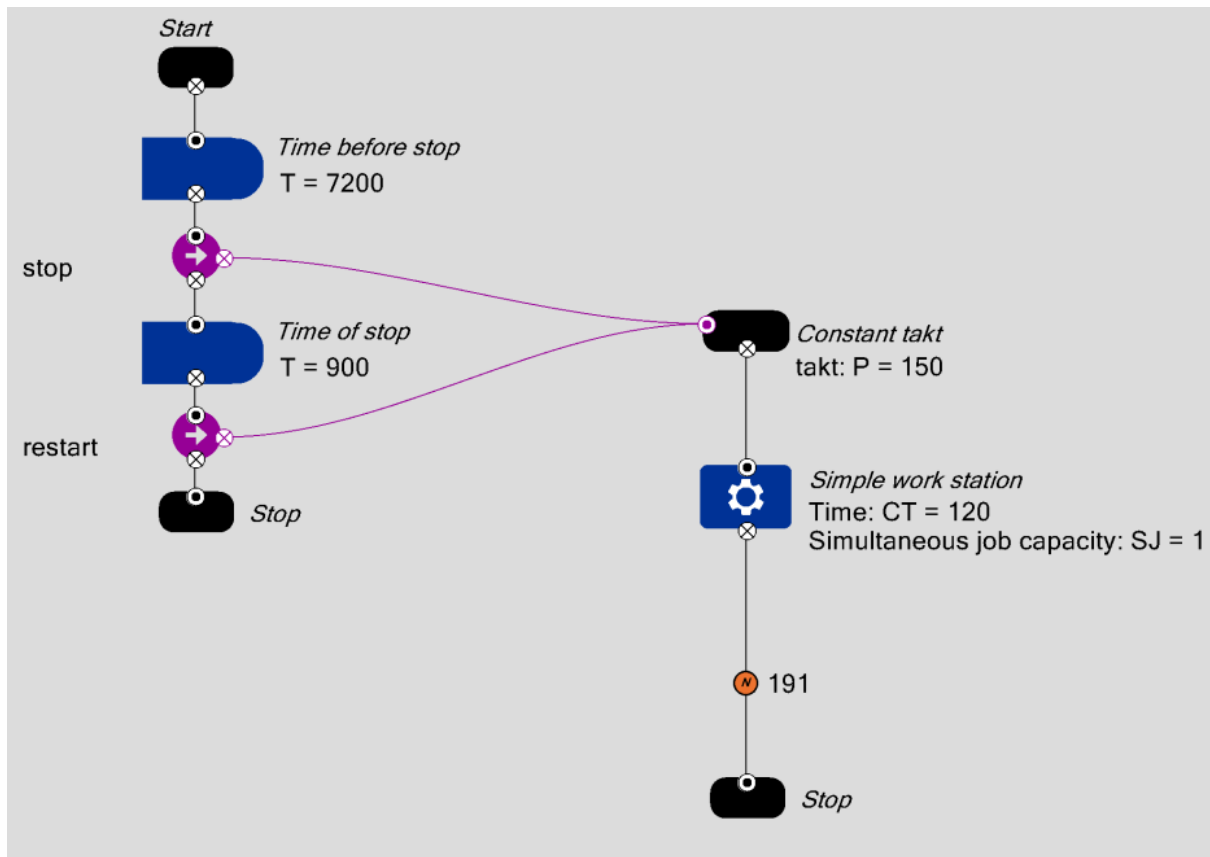
Now drop a message input port in :



And leave the source block by double clicking in space :

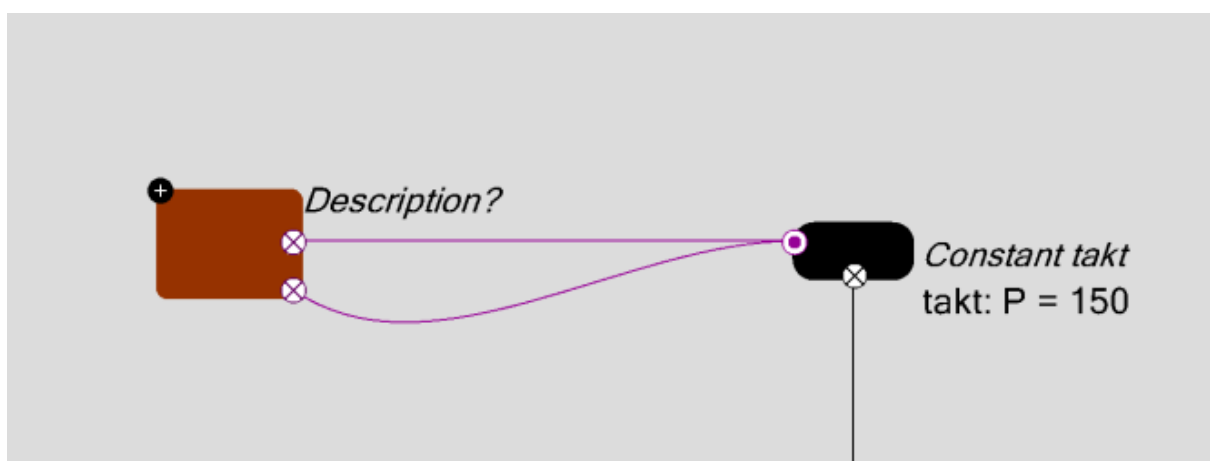


Finally connect up the message ports as follows :



Subsystems:

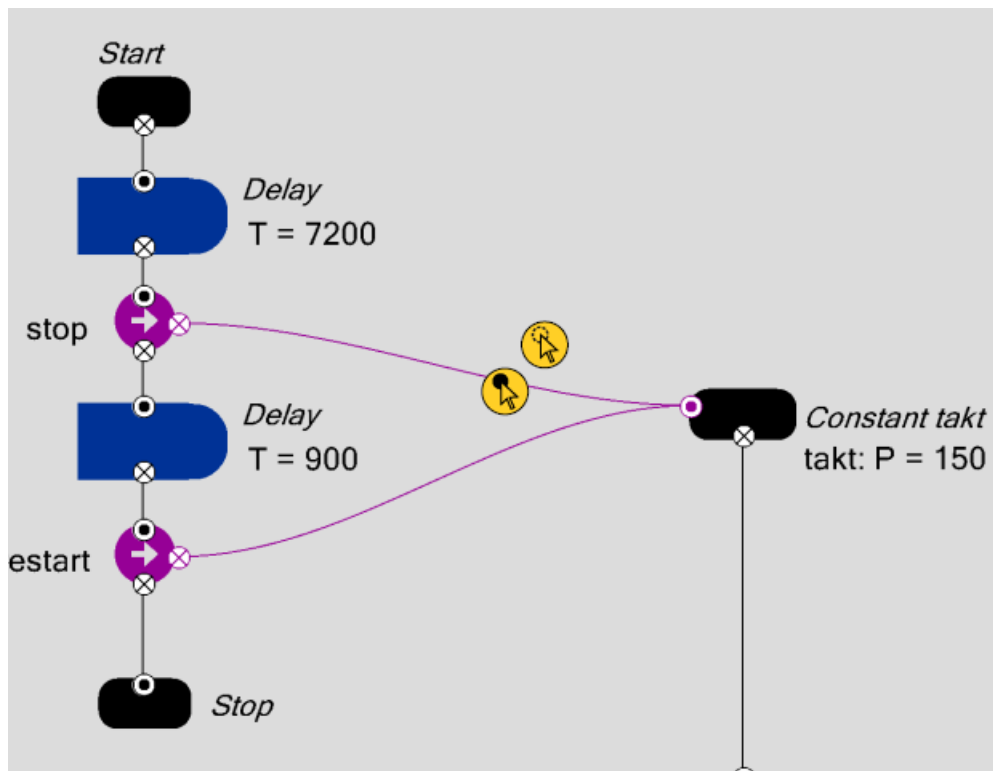
It makes perfect sense to turn the stoppage model into a subsystem. However, if you try it you may end up with this :



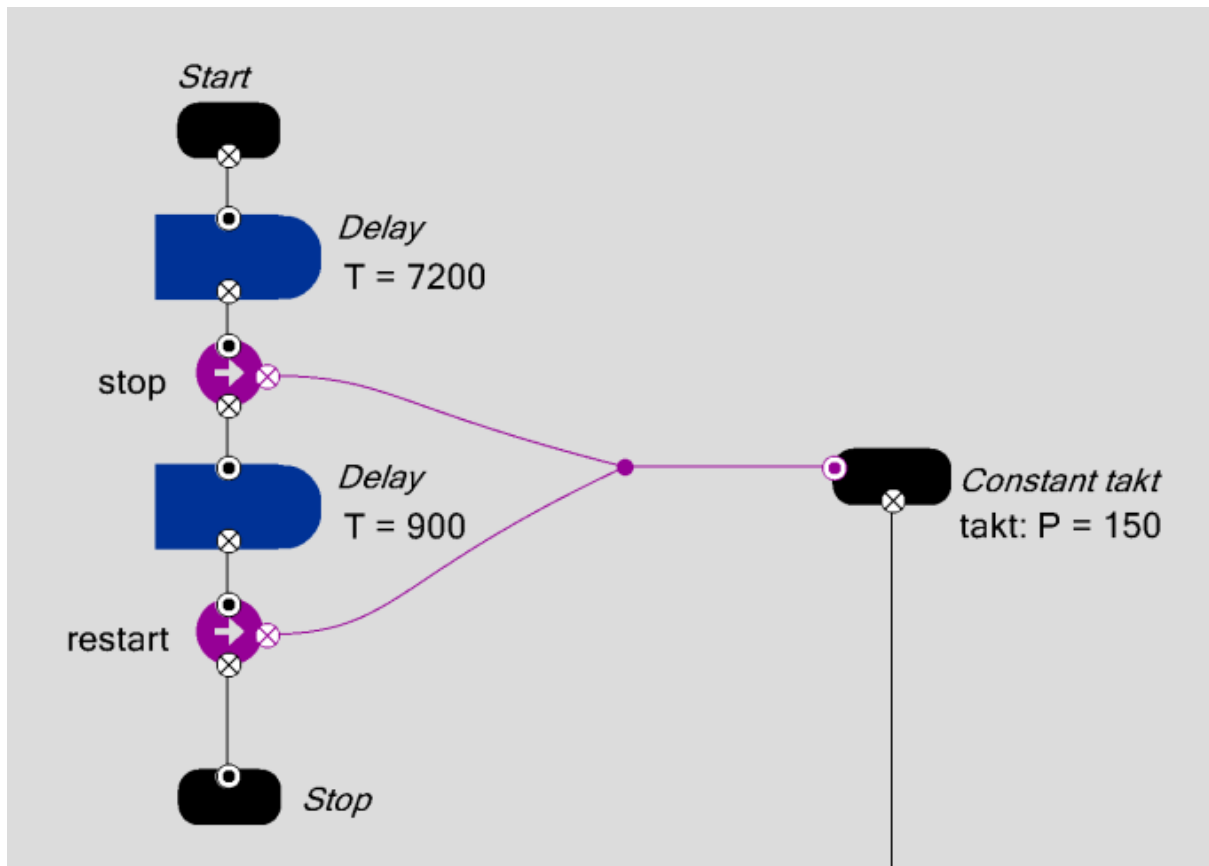
Because there are two lines cut by the boundary, aliquis produces two ports on the subsystem.

Solution :

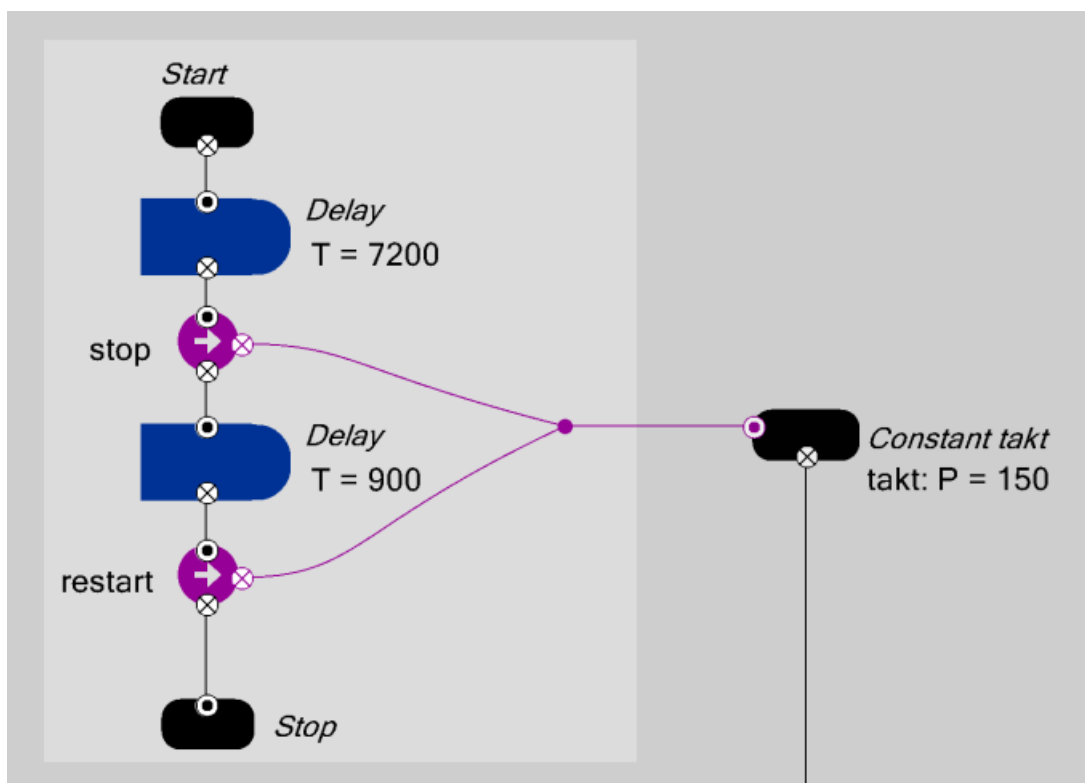
Drag one of the lines to create a connector :



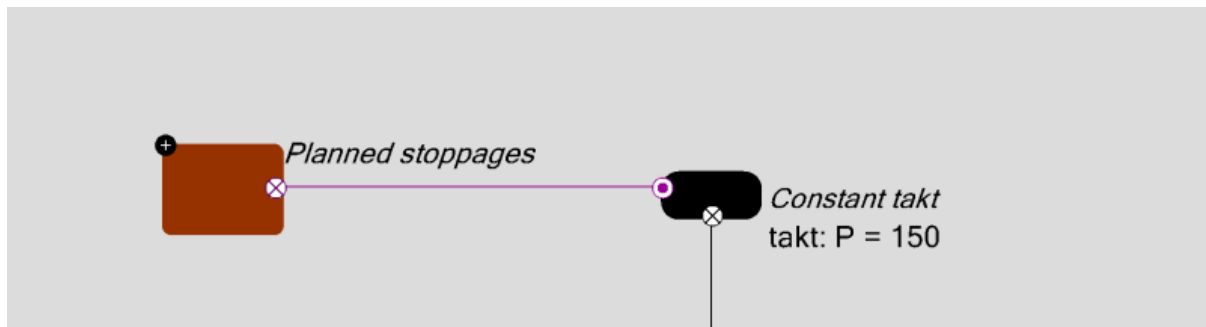
Reconnect like this :



Make the subsystem : (Select, Right click->"Make into subsystem"



For reference, take a look at *Chapter9Ex1Done*



Exercise 1– extend your planned stoppage model to include a lunch break for 40 minutes at midday and another afternoon break at 2.30 for 15 minutes. What is your expected production per shift?

Modelling unplanned stoppages.

In this example we will model a simple specific machine breakdown. Data on maintenance actions shows that this breakdown occurs completely randomly. Last week 5 times. Previous weeks 0, 6, 2 and 1 time respectively. This particular breakdown is well known and it takes 50 minutes to repair the machine when it occurs. The production line is not stopped and a small queue of work will be allowed to develop.

Time for a bit of statistics. We have data for 5* 8 hour shifts or 144 000 seconds. We also have 14 breakdowns. So our *best estimate* of our MTBF (mean time between failures) is

$$MTBF = \frac{\text{total time}}{\text{number of breakdowns}} = \frac{144000}{14} = 10286 \text{ s}$$

An aside :

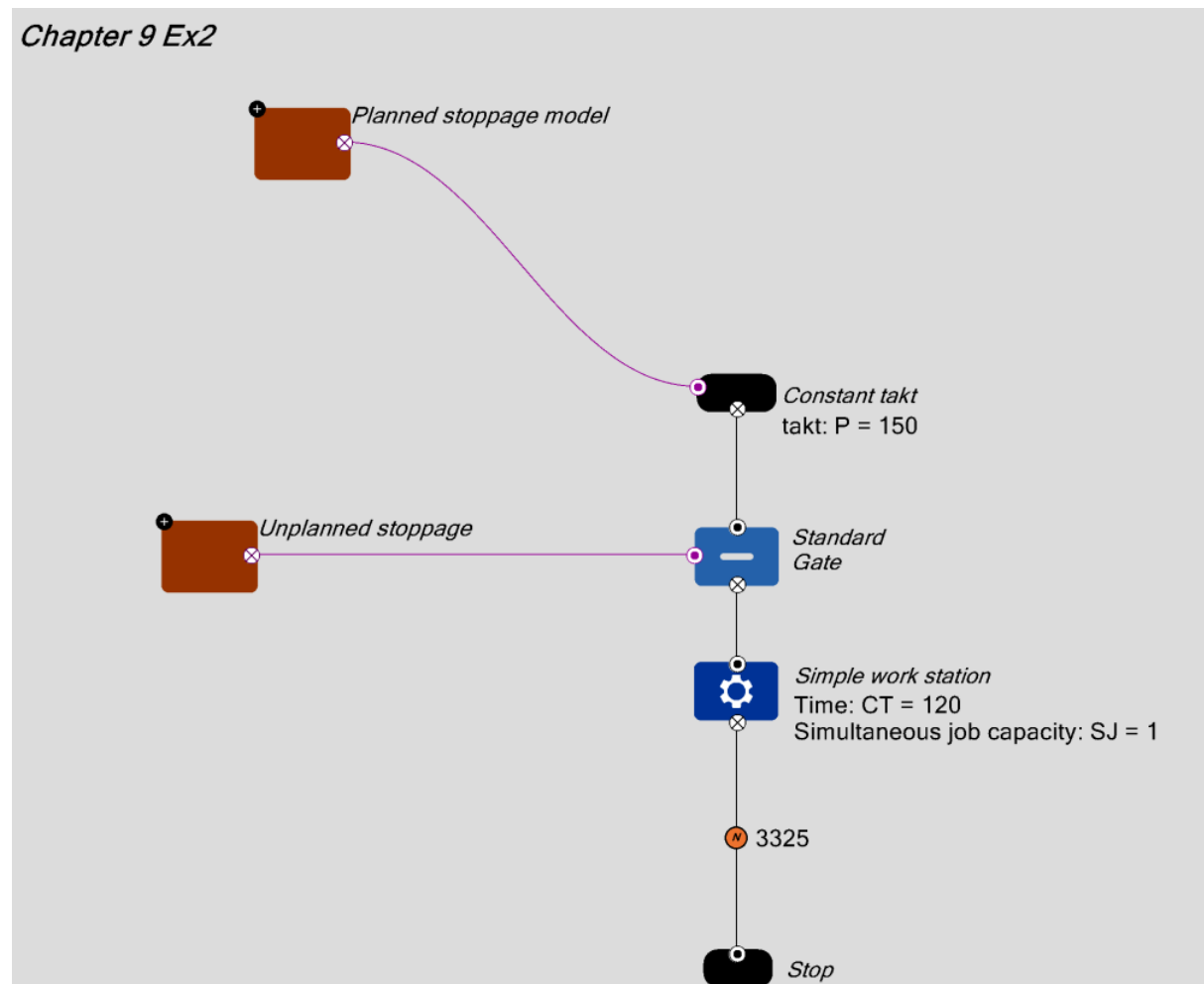
What do the words “completely randomly” and “best estimate” mean to you?

“Completely Randomly” is not a technical term. In reliability engineering theory, it might be interpreted as “Constant hazard”. That means that in any given minute that the machine is operating, there is a constant probability that it might fail. According to the established theory, constant hazard is equivalent to times between failures being modelled as an exponentially distributed random number with the mean being the MTBF

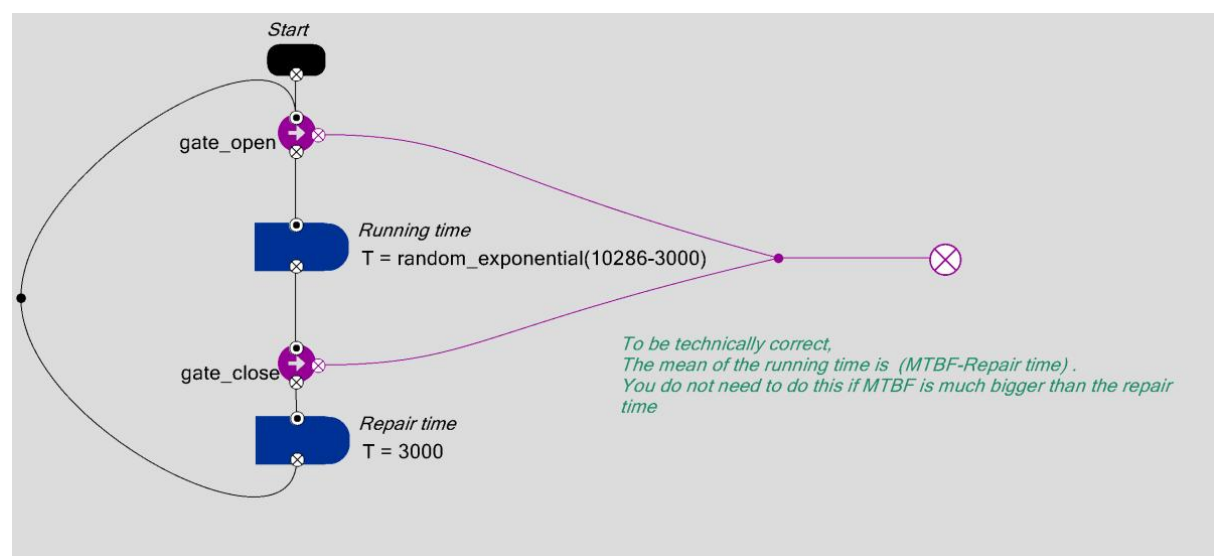
“Best Estimate”, however, is a technical term. According to statistical theory there is a “population” in this case the population does not really exist. But you can imagine it as a system which has undergone an infinite number of breakdowns over an infinite time. This population will have a true MTBF. However, we just have a sample of 14 breakdowns. Our calculation of MTBF based on this data is obviously not the population mean. Imagine the next break down coming in. If the interval since the last breakdown is highly unlikely to be exactly the same as our current 14 datapoint based number. So we have to accept that *our* MTBF is going to change every time we get a new breakdown. It can only be an estimate of the population MTBF. But the point is – it is the best estimate of the population MTBF we can make, given the data we have.

Back to modelling :

Starting with Chapter9 Ex2, on your own, model the breakdown like this :

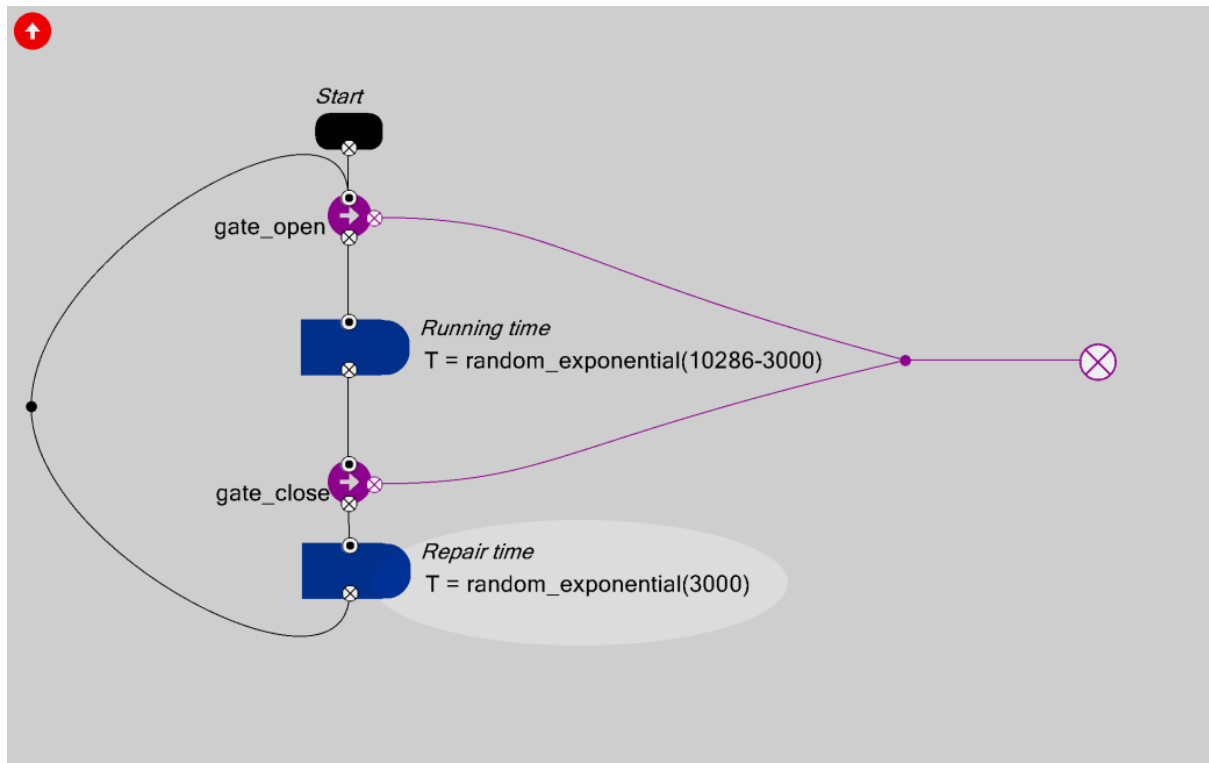


Inside the unplanned stoppage subsystem :



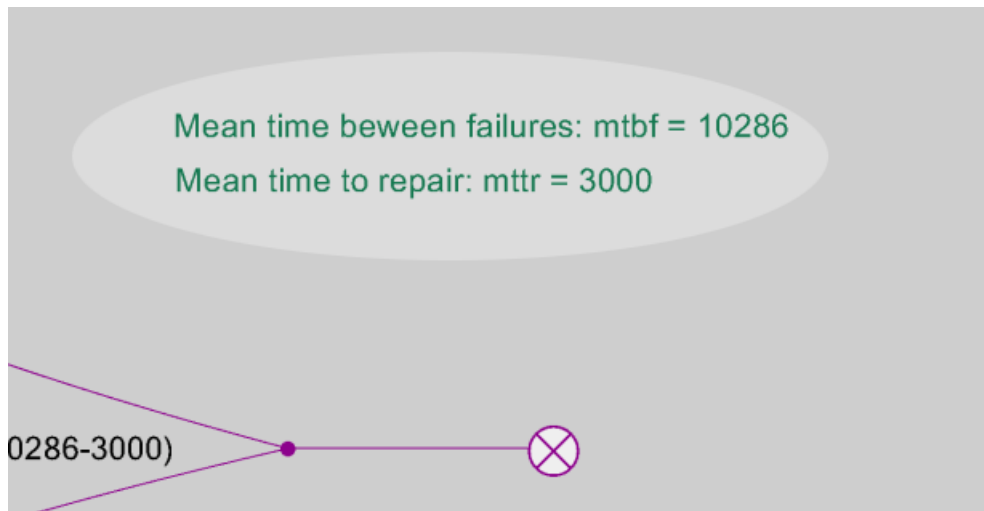
Adding random repair times:

Let's extend this. Let's also assume that the time to repair the fault is also subject to variation. (Spare parts or fitters might not be available, The repair might need testing and reviewing etc). It is common to log the stoppage times and get another reliability parameter from a real system. This is Mean Time to Repair (MTTR) . Like the breakdown, it can be modelled in different ways. For a generic system, it is often the case that we would model a breakdown like this :

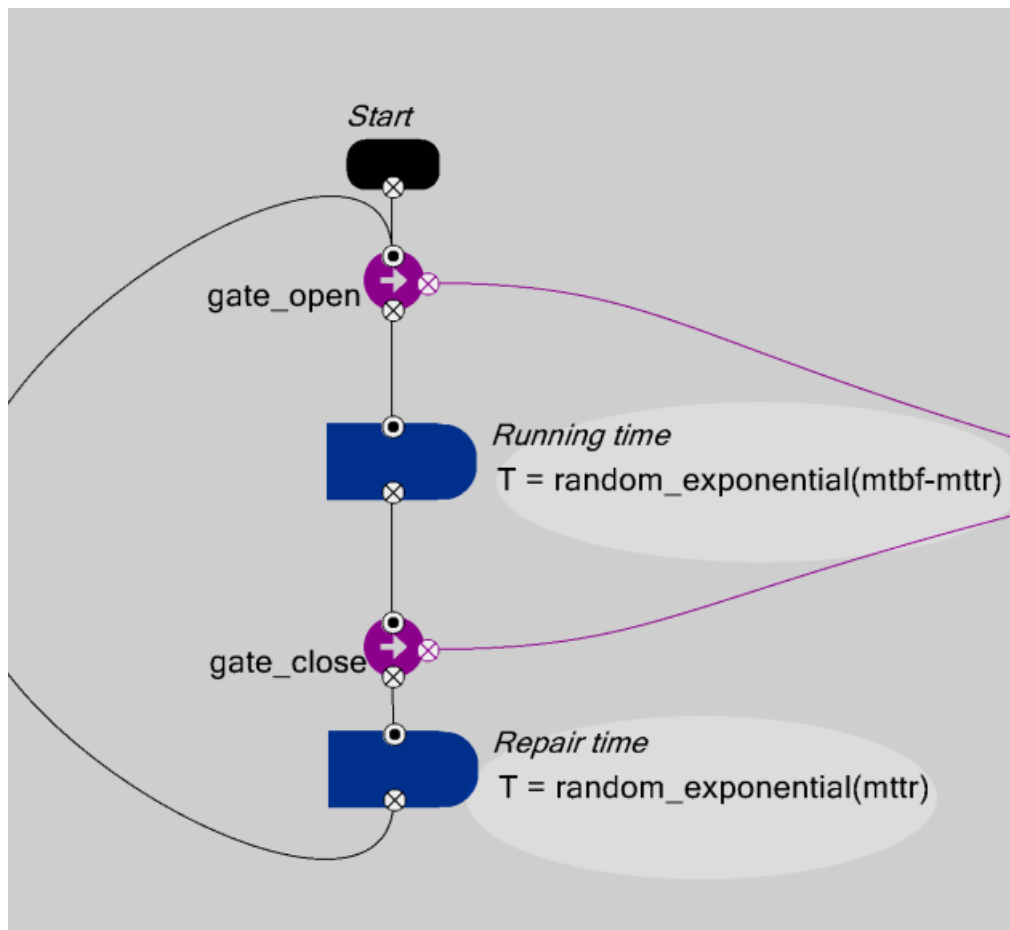


Paramaterisation of the breakdown model :

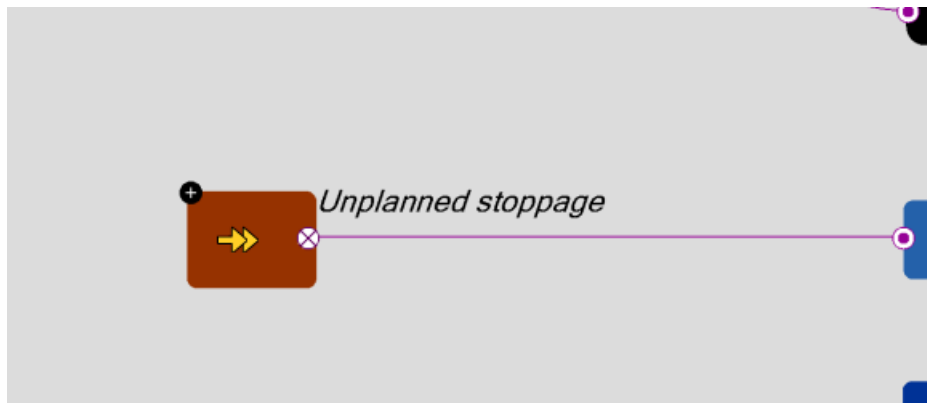
Lets make our breakdown model much more generic. There are 2 key parameters – MTBF and MTRR. Add variables to represent them :



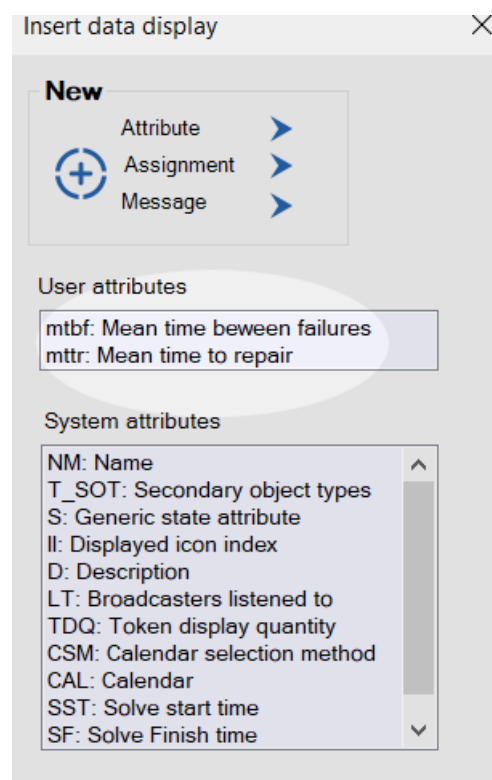
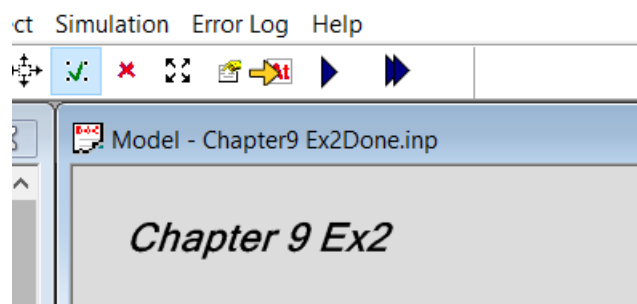
Set the blocks to refer to these parameters:



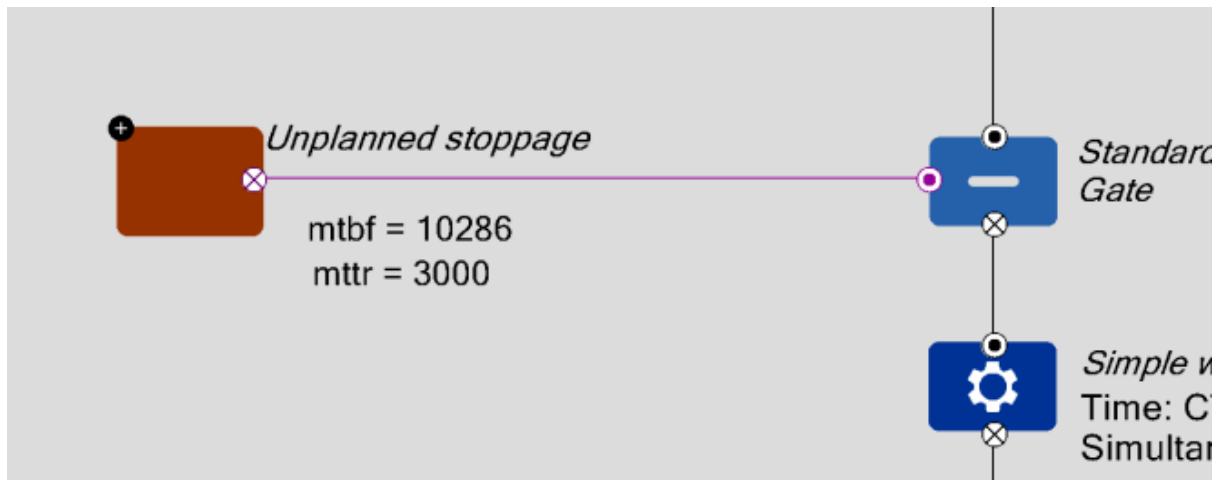
Now leave the subsystem (double click in space) and enter its external view (double click on the subsystem)



You can now expose these two parameters :

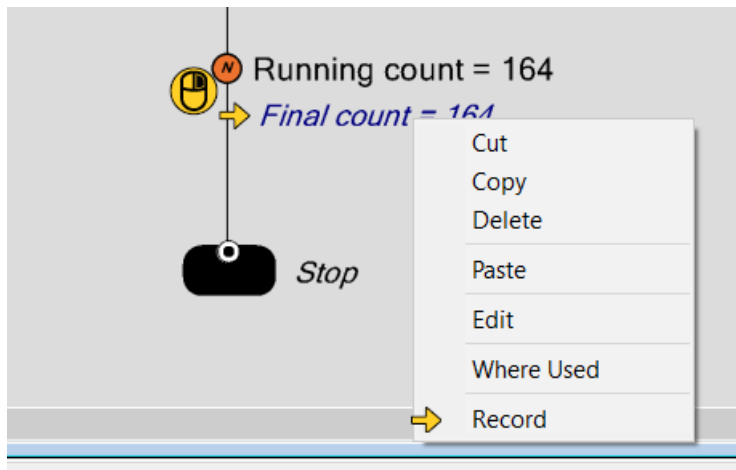


To get this result :



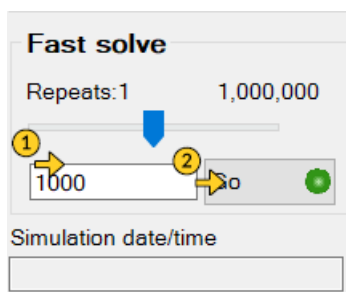
Reference model : Chapter9 Ex2 done

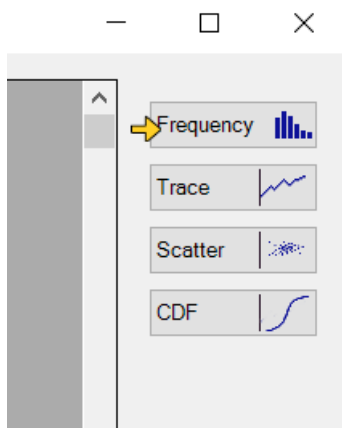
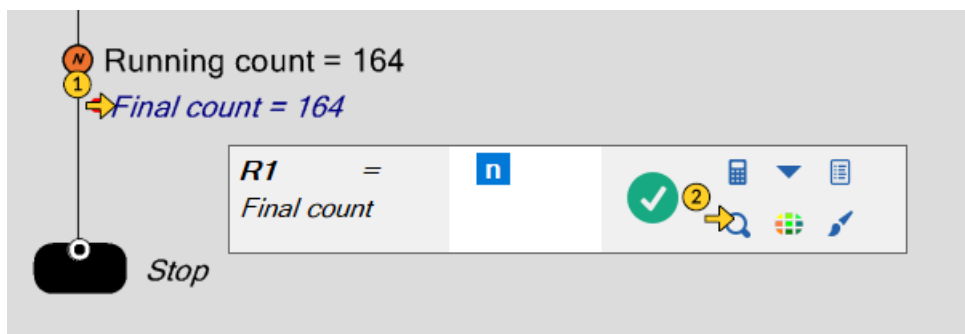
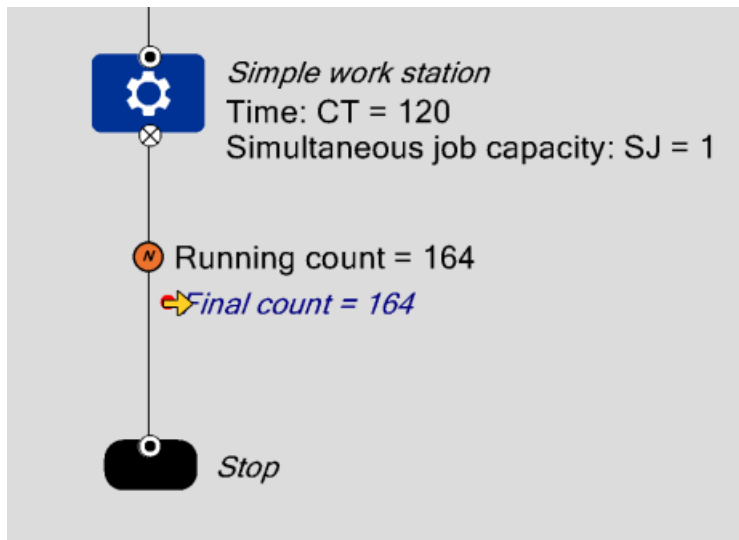
Use the RH button to save this variable :



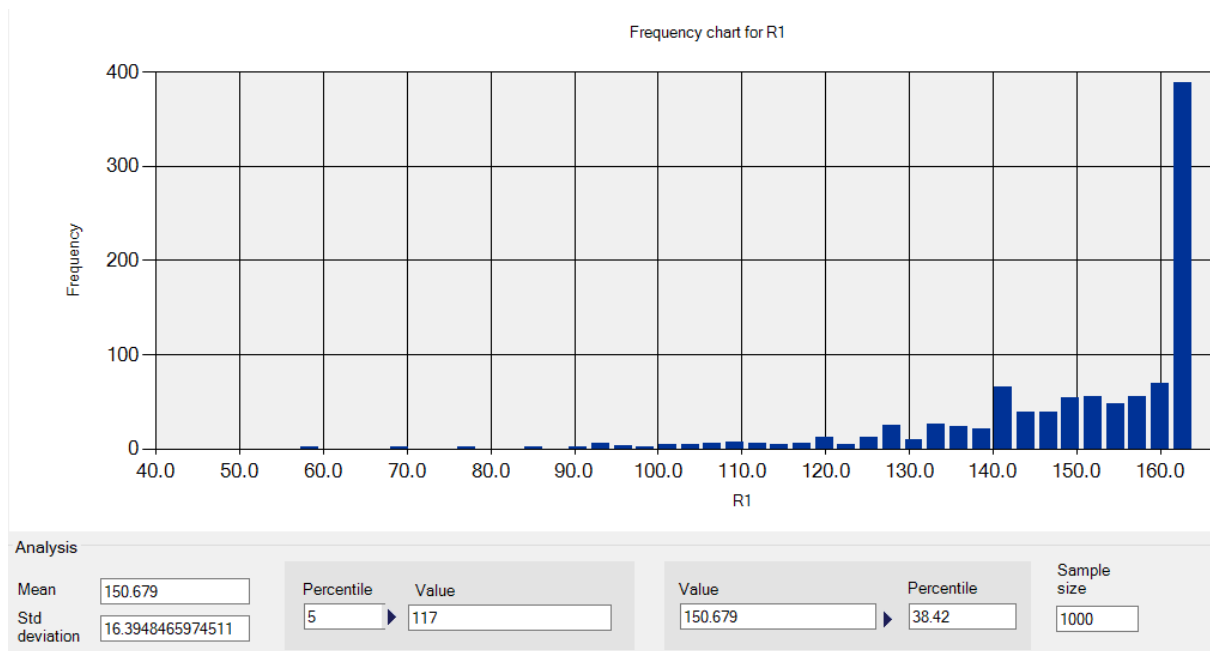
Now solve the model 1000 times to see the variability of daily production.

Error Log Help





Result :

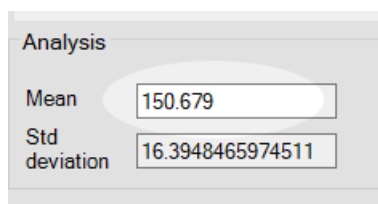


This shows that, with random failures and random repair times, production per shift is subject to considerable variation. In the seven working hours we have and with a takt time of 150 s, the maximum production would be :

$$\frac{7 \cdot 60 \cdot 60}{150} = 168.$$

This is frequently reached. In many shifts there is no unplanned stoppage. However, when these stoppages occur, the loss in production can be severe. In one test, production was down to 60 units.

The graph, also gives us an average production of 151 units :



It is worth reflecting that production line capacity is often related to contractual obligations. What would you need to do in this situation if the production line is obliged to produce a minimum 130 units a day? Based on this information, this production target cannot be guaranteed without stock holding (inventory).

Calculations vs modelling :

Lets go through the standard textbook approach to predict the average production of this facility.

This production line contains just one machine with the following parameters:

Cycle time	120
MTBF	10286
MTTR	3000

And the takt time is 150 seconds.

First calculate OEE remember :

$$OEE = performance * planned utilisation * unplanned utilisation * yield$$

Where :

$$Performance = \frac{cycle\ time}{takt\ time}$$

$$Planned\ utilisation = 1 - \frac{planned\ stop\ times\ in\ a\ period}{period}$$

$$Unplanned\ utilisation = 1 - \frac{MTTR}{MTTR + MTBF} =$$

And :

$$yield = 1 - \frac{scrap\ parts}{scrap\ parts + good\ parts}$$

In this case :

$$Performance = \frac{120}{150} = 0.8$$

$$Planned\ utilisation = 1 - \frac{60}{480} = 0.875$$

$$\text{We know : } Unplanned\ utilisation = 1 - \frac{unplanned\ stop\ times\ in\ a\ period}{period}$$

But this can be reformulated to :

$$Unplanned\ utilisation = 1 - \frac{MTTR}{MTTR+MTBF} = 1 - \frac{3000}{3000+10286} = 0.774$$

And in this case we have no scrap :

$$yield = 1$$

Therefore

$$OEE = 0.8 * 0.875 * 0.774 * 1 = 0.54$$

$$\text{Maximum theoretical production per shift} = 28800/120 = 240$$

$$\text{Expected production per shift} = \text{Max theoretical} * OEE = 240 * 0.54 = 130$$

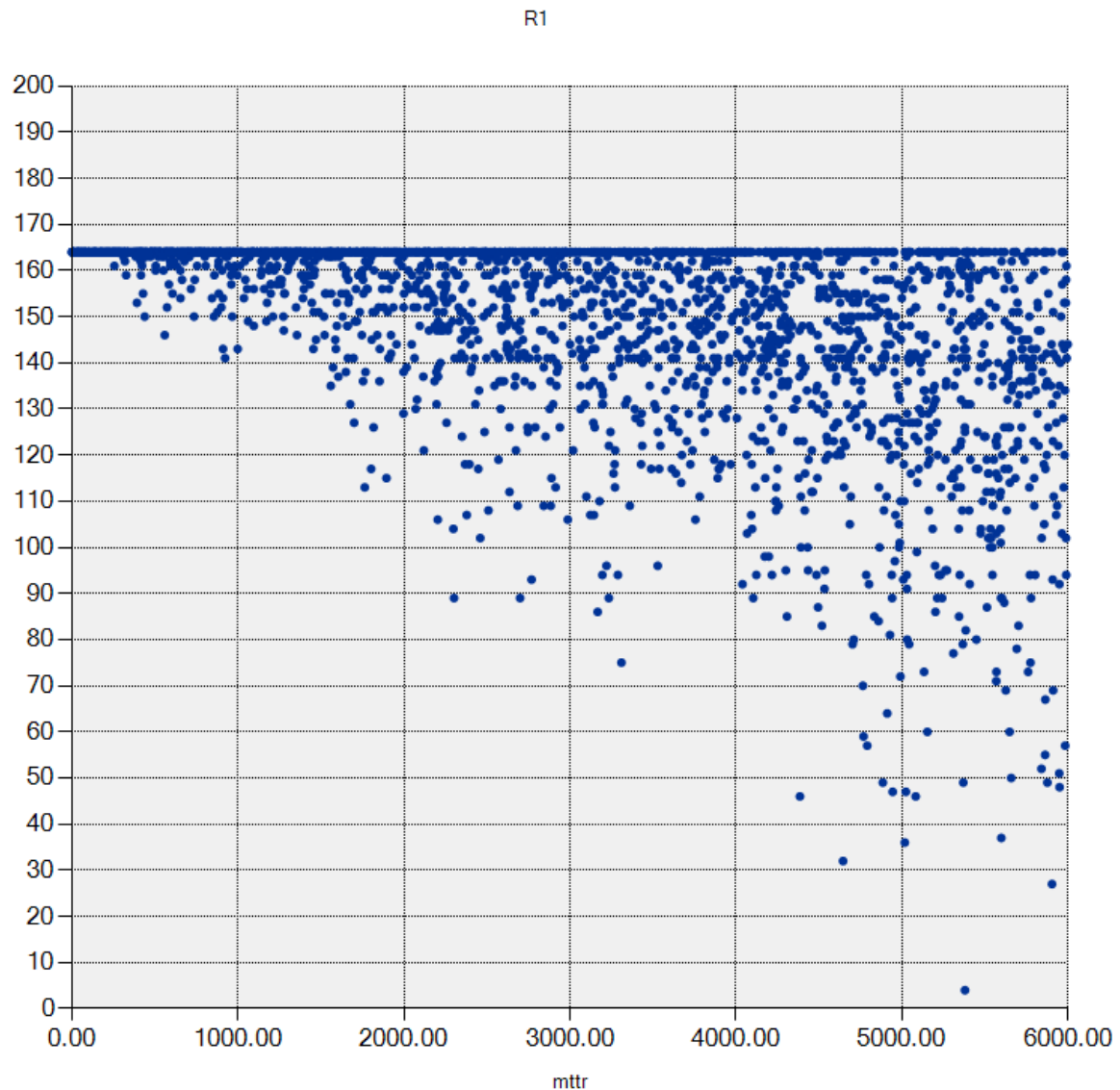
Interestingly, the model produces a higher average output than the calculation. There is a good reason for this. In the model, if there is a break down, a queue will form. Once the queue is there and when the machine is working again, each part will take 1 cycle time (120) rather than one takt time (150) to be processed and the queue will slowly be reduced.

Put another way, the machine plays catchup with the line after a breakdown.

Of course this is for one machine/workstation and still theoretical. Real production lines are made of multiple machines, with or without buffers . If the buffers are there, they are finite rather than infinite and whole production lines may need to be shut down when the buffers are through. This will be dealt with in chapter XXX

Exercise A

Using the example model, Obtain a graph which shows the sensitivity of daily production to MTTR:



Reference : Chapter9 Exercise A Solution

How do you interpret this scatter graph?

Multiple breakdowns, one workstation

In this model, there is one machine and we have considered one type of unplanned stoppages. In reality, we may log various types of unplanned stoppages and each type may have a different MTBF and MTTR.

Lets assume that our machine has two types of breakdown. Sumarised as follows :

Cycle time	120
MTBF (A)	10286
MTTR (A)	3000

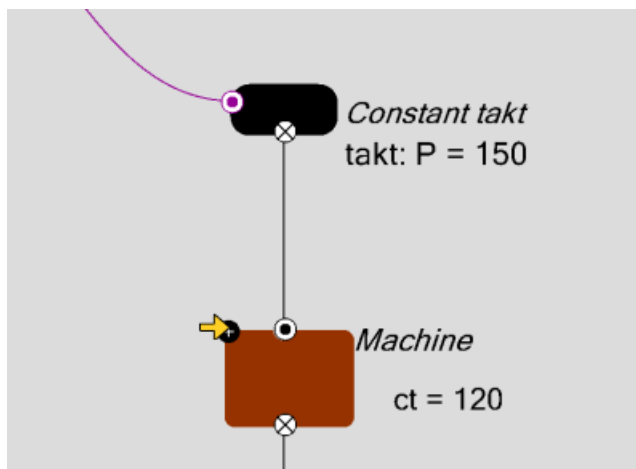
MTBF (B)	20000
MTTR(B)	2500

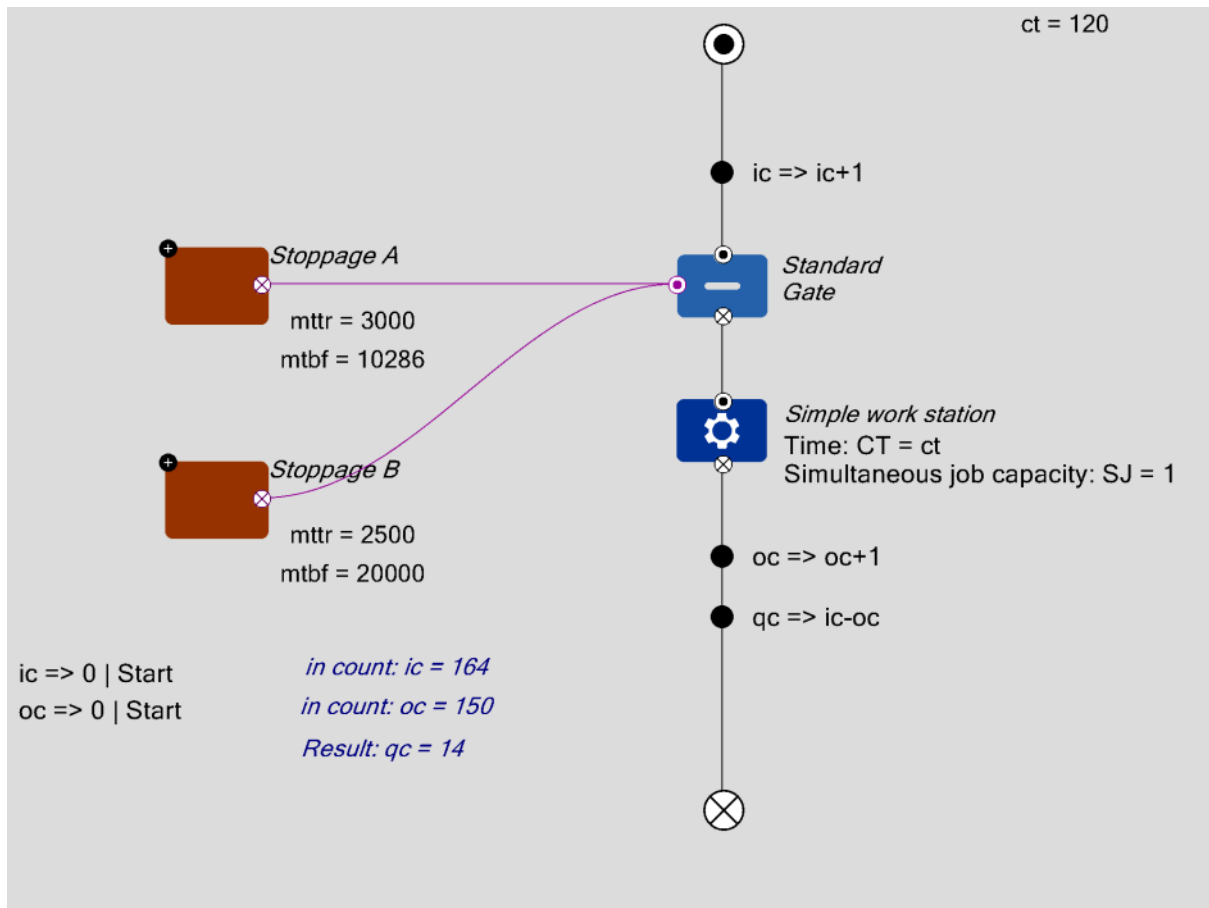
Modelling approach :

This can be modelled reasonably accurately by copying the existing model and adjusting the parameters:

Create a different model for each kind of breakdown :

Chapter 9 Sample 2





Note two stoppage models and parameters changed.

A More Mathematical approach

Alternatively we can do some maths using the standard reliability theory :

$$\text{Combined MTBF} = \frac{\text{MTBF}(A) * \text{MTBF}(B)}{\text{MTBF}(A) + \text{MTBF}(B)}$$

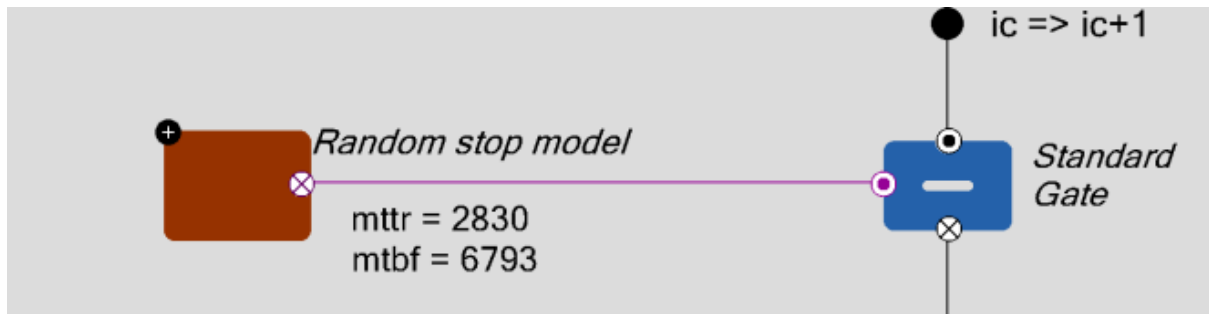
$$\text{Combined MTTR} = \frac{\text{MTBF}(A) * \text{MTTR}(B) + \text{MTBF}(B) * \text{MTTR}(A)}{\text{MTBF}(A) + \text{MTBF}(B)}$$

Putting the numbers into these formulas gives us

Combined MTBF = 6793s

Combined MTTR = 2830s

In *Chapter 9 sample 3* we can see the result of this.



It gives a similar but slightly lower average daily production than in Sample 2 (around 150 units) . This is because repairs can occur in parallel in sample 2 but not in sample 3.

Assuming that machines do not frequently break down and that repair times are relatively small, both solutions will tend towards the same results.

Multiple breakdowns Multiple machines

In *Chapter 9 sample 4* we look at a production line with three identical machines in series.

Because they do not break down at the same time, sometimes a lower machine will be starved (ie not have anything to come in) . This will result in a lower production level.

Running the model 10,000 times gives an average production of 125 – lower than previously.

Multiple breakdowns Multiple machines no buffer stocks

In this case, when an unplanned stoppage occurs anywhere on the line, the machines above the breakdown must be prevented from delivering anything through their exit if there is a breakdown below.

Chapter9 Sample 4 captures this . The average production per shift is 61. – Illustrating the need for buffer stocks (inventory) if you have unreliable systems.

Summary of results – Sample 1-4

	Model	Mean	Standard deviation	Conclusion
No breakdown- Single machine (BaseLine)	Sample0	168	0	This is the base line

Single machine/Single Failure mode	Sample1	153	16	Any breakdown will lead to loss of production and variability in production levels
Single Machine/Multiple failure mode	Sample2/3	146	19.4	Increased complexity generally reduces reliability which in turn decreases average production and increases variability in production levels
Multiple machine/ multiple failure mode buffer stocks	Sample 4	114	24	Multiple unreliable machines in a production line (with starving but not blocking) further decreases average production and increases variance
Multiple machine/ multiple failure mode no buffer stocks	Sample 5	36	19	Buffer stocks between locations increase overall production and decrease variability. This is especially important when reliability is a problem. But don't forget – this is a waste in lean (inventory)
Multiple machine/ multiple failure modes. No buffer/ scrap	Sample 6	31	17	Problem exacerbated by scrap

Put another way, This clearly shows why machine reliability (and scrappage)matters in manufacturing systems if you also want to be lean and eliminate waste. The ideal lean situation is to have :

- very large MTBFs - ie reliable equipment
- very small MTTRs - ie quick repair system (This can also be seen as time to get back on line – which could be improved with spare machines in critical areas)
- Small or zero buffer stocks between workstations.
- High quality – low scrap

To get there requires a lot of work !

Exercise

Using sample 6 , what results do you get if you double MTBFs, half MTTRs and half the scrappage rate.

Summary

In this chapter we have introduced the very real subject of workstation stoppages and downtime into our equation. We have examined some important concepts :

- Performance
- Planned stops
- Unplanned stops
- OEE
- MTBF and MTTR

We have also looked at how modelling can produce unexpected and perhaps more realistic results due to “performance catch up”

Finally we looked at a more probabilistic approach to daily production quotas. – Which shows that risks of not meeting production levels are still significant even on generally reliable systems – Highlighting the importance of reliability, maintenance and reliability data collection in any manufacturing system.

Questions

1. Going back to chapter 4. Which of the 8 wastes (Defects, Over Production, Waiting, Non-utilised talent, Transportation, Inventory, motion, Extra processing) are increased by machine breakdowns and why?
2. In your own words, describe some reasons why it is important to understand system reliability in general .
3. What is your critique (good and bad) of OEE as defined by $\text{Performance} * \text{Utilisation (planned and unplanned)} * \text{Yield}$.

Answers

1. Mainly Waiting and Inventory. Arguably – non utilised talent. Waiting – because production elements below the stoppage will be waiting for parts during the stop. Alternatively, the company could build buffer stocks between machines – but this gives us inventory. Non utilised talents= - because the guys above and below the stop will be idle.
2. It says – in your own words! However, briefly Poor reliability and high difficulty in repair both decrease average production levels and increase the variation of production levels in any given period. Decreased production can lead to decreased sales. Increased variability increases the probability that production quotas will not be reached within any given period leading to loss of sales and customer dissatisfaction plus increased “fire fighting” style management – which frequently impacts quality and morale.
3. This measure of OEE is good because it can be easily measured for any single machine and allows you to identify the areas for improvement. However, It is not actually a good measure for predicting production volumes and projects a single

number solution – when actually production volumes really should be treated in a statistical way – made possible, either by more advanced mathematical techniques and/or by modelling.