

Chapter 3 – Variation , uncertainty and risk

In the real world , many jobs do not really have a precise fixed times. As I embark on writing this chapter I do not really know how long it is going to take me – although I can estimate. Once I have written it, I could measure the amount of time it took me. But does this mean that the next chapter will take the same amount of time? No - because there is clearly a variation in the amount of work involved in producing different chapters and there is also a learning curve effect. Hopefully I will get faster as I write. Lastly, do I know of anything that could go wrong? I may, in a highly unlikely scenario, discover an Aliquis bug that needs fixing, or I may decide that I need to do some more time-consuming research on a specific subject. So, as things stand I know there is going to be a variation in the amount of time it takes to write a chapter. I also know that I am uncertain about the time it will take me to write this chapter now. So, we have variation and uncertainty and they are intrinsically involved in everything we do. Unfortunately, to our cost, many of us habitually ignore them.

Lets examine these terms more thoroughly.

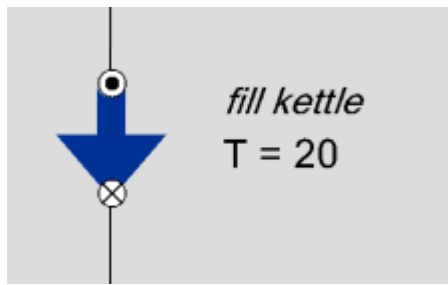
Variation: If you have a mass production process, a key dimension or a time for a specific action might be subject to variation. Every time it is measured it could give a different value. Most repeating processes are subject to at least some variation.

Uncertainty : This is referring to the future. In a project, the time to do a job is always uncertain. You cannot predict exactly how long it will take. In a manufacturing process which is subject to variation, the *next* measurement is uncertain. Uncertainty also refers to future events or condition which may or may not happen. For example a defect in a manufacturing process , or the weather condition in a project, an unsatisfied customer refusing to pay for services.

Risk: In project management, simple definition of risk is “An uncertainty that matters”. Likewise, manufacturing processes are full of uncertainty about their future output. We can define process risks as the process uncertainties that matter. The idea here, is that *you* chose which uncertainties you may wish to monitor, minimise or ameliorate.

Earlier we made a virtual cup of tea – lets think some more about the uncertainty/ variation in this process. We will focus on the simple task of pouring the water into the kettle.

In our previous mode it looked like this :



But what about the real world ?

Variation The time needed to pour water into a kettle is not actually fixed, it varies every time the process is carried out. It is not possible to say “*The time is 20 seconds*” It might be *more* truthful to say “The time can be modelled more accurately as random number which follows a normal distribution. The mean is 20 and the standard deviation is 1. “

Uncertainty : Even if we perfectly understand our variable times we cannot be certain what the *next* pour time will be. You might think of this as level 1 uncertainty.

What if we have not properly researched pouring water into a kettle?

We might *assume* the pouring times follow a standard deviation and we could estimate the mean and the standard deviation based on data and knowledge. These estimates are never perfectly certain. This is level 2 uncertainty. In statistics this is called *Knightian* uncertainty.

Knightian Uncertainty is not quite the same as variation – it reflects a lack of knowledge. Study , research and testing reduces Knightian uncertainty but it does not reduce variation. On the other hand, engineering improvement actions will frequently reduce variation but may actually increase Knightian uncertainty as processes will need further study after changes to establish the new parameters.

Event or condition-based uncertainties : Kettle filling might have an associated event that the tap has run dry. It is often very difficult to be certain about the *probability* that this kind of event will occur. All we really know is that the probability is low . Good data is seldom available. Put another way, the level 2 or Knightian uncertainties are usually quite high for seldom occurring breakdown events. nevertheless, we know this event is possible – we also know that we are uncertain as to whether it will happen today.

Black Swan Events: Prior to Captain Cook finally working out that there was a continent called Australia, there was a basic assumption, in Britain at least, that swans are white. The idea of a black swan was totally inconceivable. So a “black swan event” is an event which can happen but which none of us can conceive. For example, in 2008, the Lehman brothers, a massive American bank, collapsed – The knock-on effect was a ripple that turned into a tsunami. The whole world financial system became completely unstable. To many – this was a completely unpredictable and inconceivable. The point is - you cannot model true black swan events – but it doesn’t mean they will not happen. That said, many so called black swan events are actually perfectly conceivable.. It’s just that people are not willing to think hard about uncertainties.

Risks: As stated earlier, *you* decide which of the uncertainties should be deemed risks. You may decide that the probability of the tap running dry is so low it doesn't matter. If, on the other hand, you consider it does matter, you will want to estimate it more thoroughly by monitoring, and you may want to put plans in place to ameliorate the possibility by planning an alternative process or holding stocks of water near the kettle.

Lean manufacturing and MURA

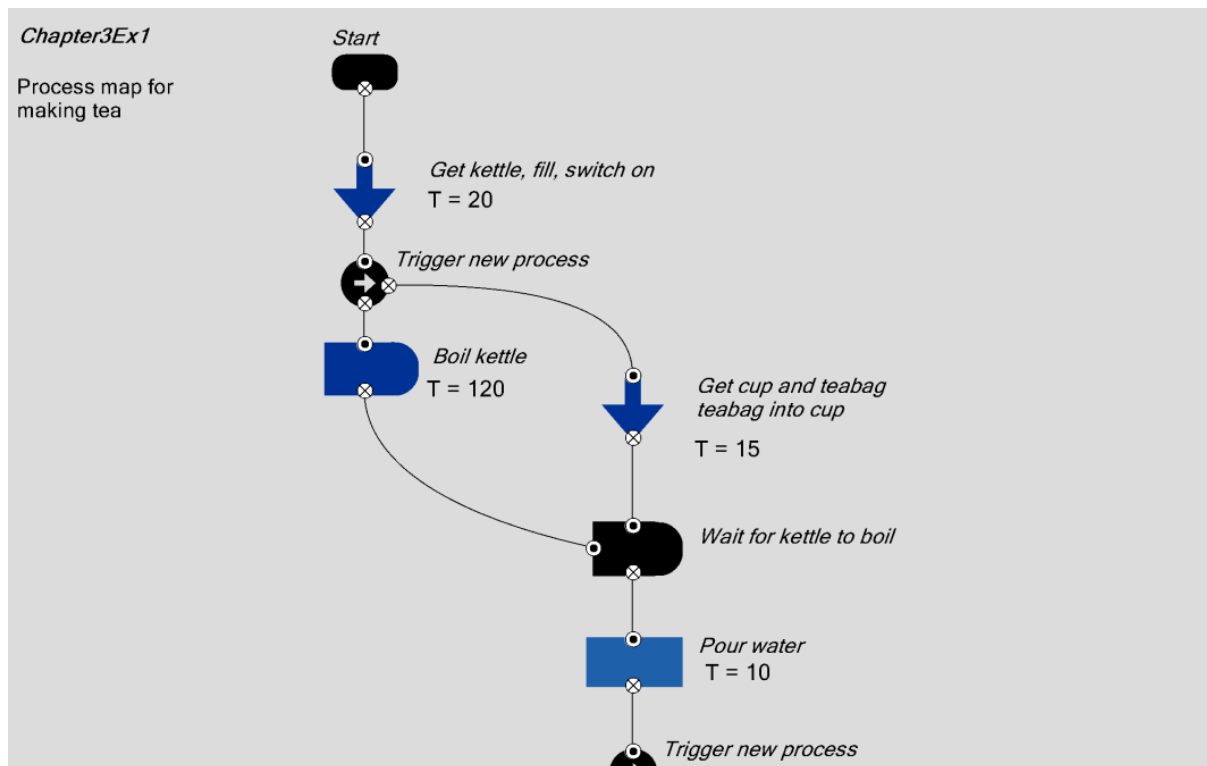
In lean manufacturing, any unevenness, uncertainty or variation is called "MURA". The lean philosophy is very clear – MURA causes waste – which in turn adds costs and reduces profits. The opposite of MURA is "flow" – an idealised state where production flows through a system like a well oiled machine.

In this chapter we will explore how to model variation/uncertainty and risk (MURA) in Aliquis.

A word of warning : This chapter assumes that you have some understanding of probability, probability functions and statistics. It will be a little tricky in places if you do have not. However, we still recommend that you work through the examples as it will give you an elementary understanding of a very important subject in manufacturing systems.

Modelling process variation/level 1 uncertainty in Aliquis

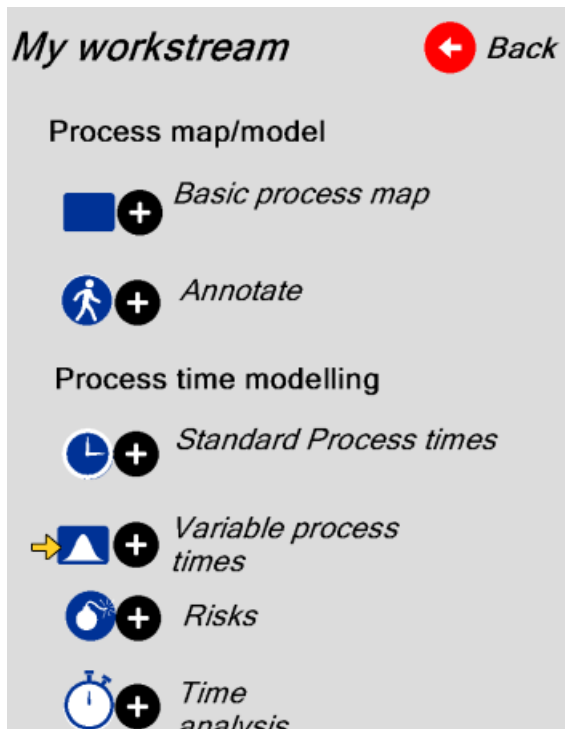
We will use our tea making process - you can find it in the exercises directory:
chapter3ex1.inp



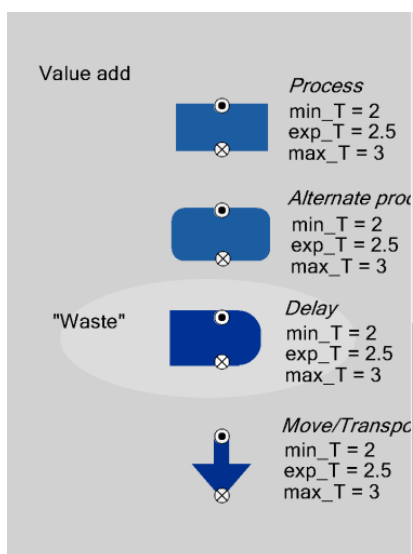
To be accurate – all the processes in this model will have some variation in their completion times – but some more than others. However it is a law of statistics that the items with the most variation contribute disproportionately to the variation in the whole system . When you are modelling variation it is generally acceptable to leave the small items with standard times. It makes the model run a bit faster. In this example we will model time variations in just three of the blocks – kettle boiling, brewing and drinking. Other blocks will be left with standard times. In each case we will model the variation with a different method.

Method 1 : Use a pre-prepared block in the library.

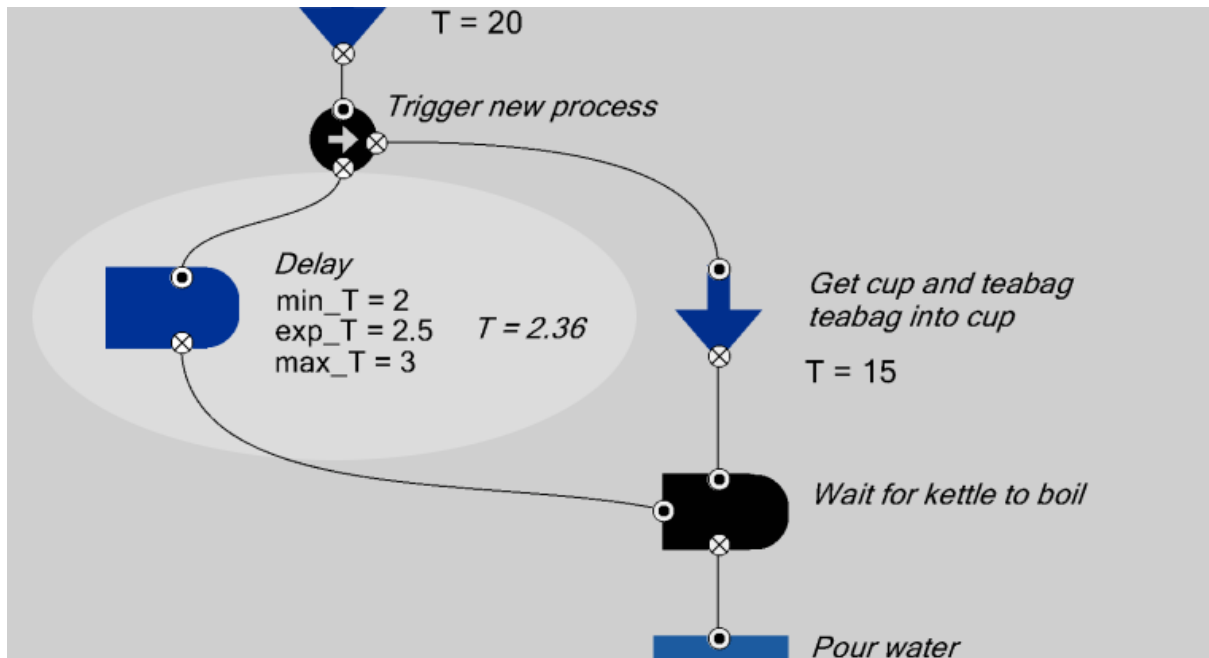
Go into the variable times sub library



Scroll down if necessary to find this block :



Use this to replace the kettle boil drop :



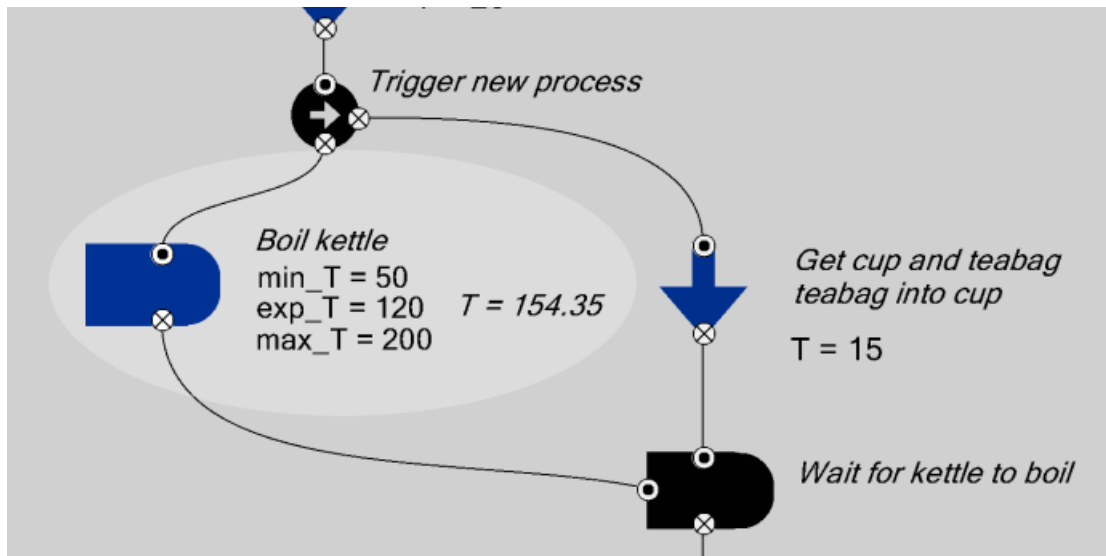
And make the following adjustments :

Description -> Boil kettle

$\min_T$ (minimum time) -> 50

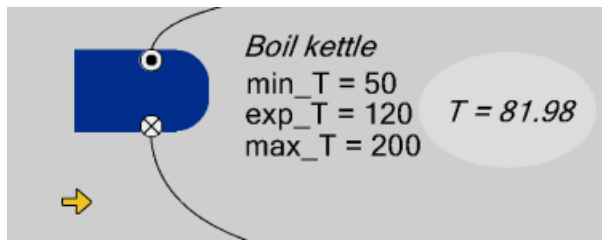
$\exp_T$ (expected time) -> 120

$\max_T$ (Maximum time) -> 200



IF you click in space repeatedly you will get examples of the calculations that this block makes for T .

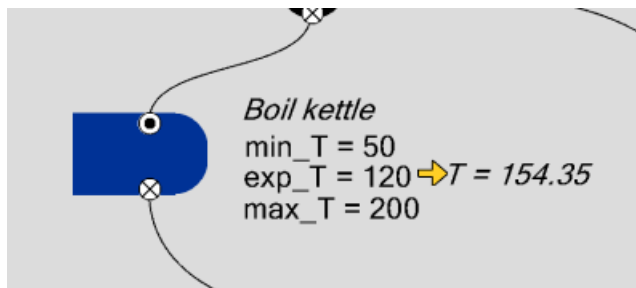
Eg



This value will change for every time the model runs.

Method 2: Using probability distributions defined directly by formulas.

In the last example, if you open up the variable T (time) you will see the underlying formula that determines how T is calculated :

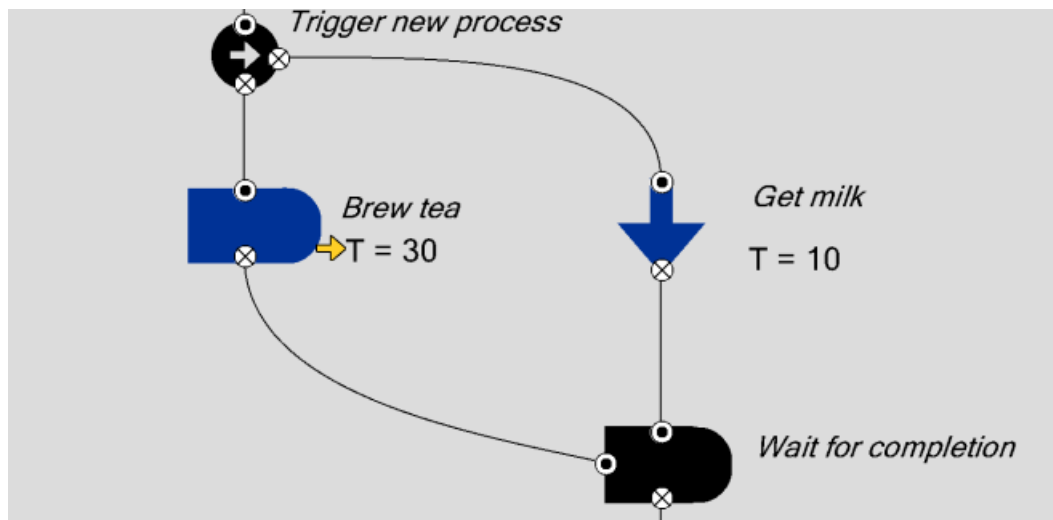


T <i>Time</i>	=	<code>random_weibull_range(min_T , exp_T , max_T)</code>				
--------------------	---	--	--	--	--	--

This is an Aliquis special formula which produces a random variable according to a Weibull distribution – it internally converts the parameters given to the standard Weibull parameters (locations, shape and scale) and then calculates a value accordingly. In this formula, the maximum value $\max_T$ is not the absolute maximum – rather the 99.9 percent confidence limit and the expected value $\exp_T$ is the median, not the mean. But it's a great tool for approximating skewed distributions quickly and in a way that most users can understand.

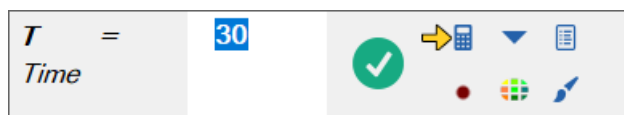
You might prefer to control the formula yourself. So for example, the time between customers arriving at a sales point is commonly modelled as an exponentially distributed random variable. We will do this now on the brew time block :

Find the *brew* block and click on the T (time) attribute.

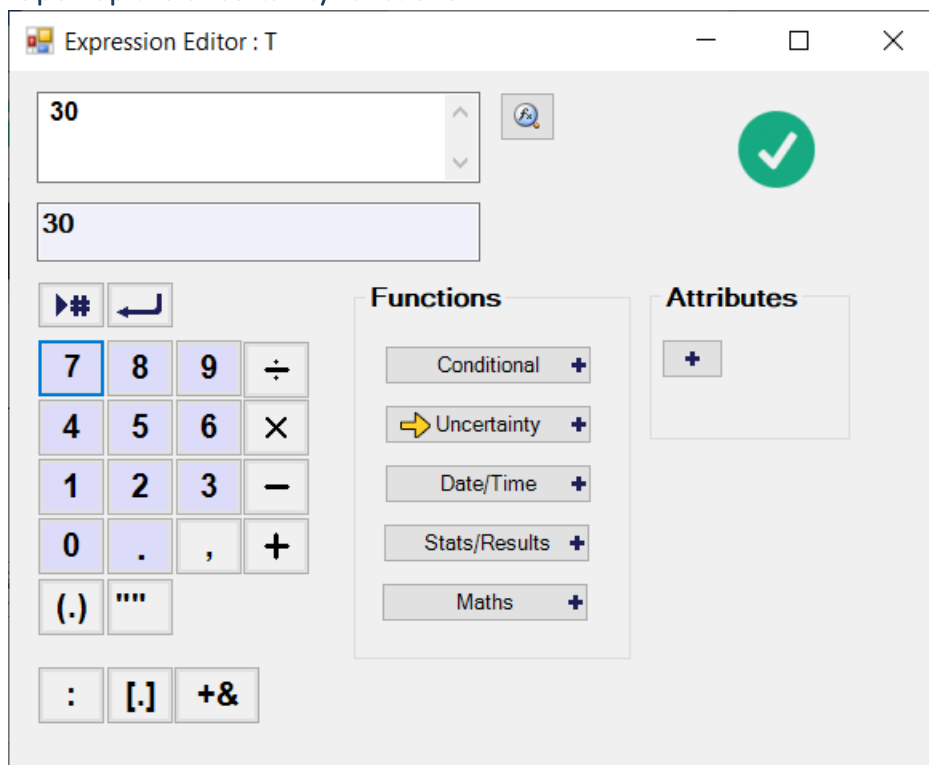


We will now use the Aliquis Expression editor to edit this function to produce a normally distributed random variable.

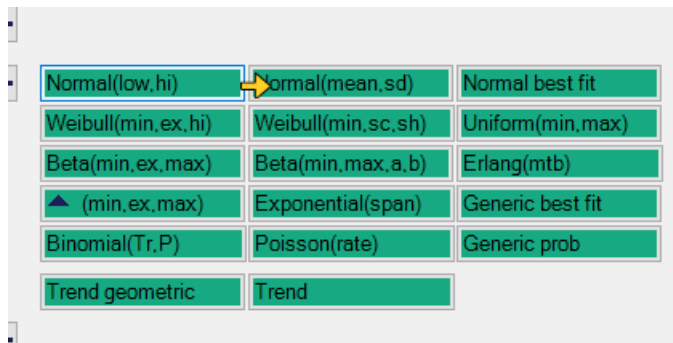
Open up the expression editor



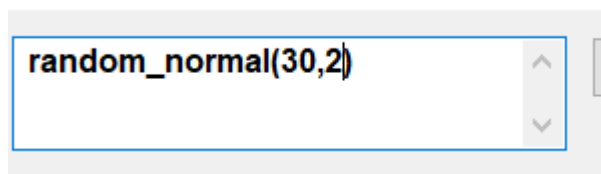
Open up the uncertainty functions :



Pick *normal (mean,sd)*

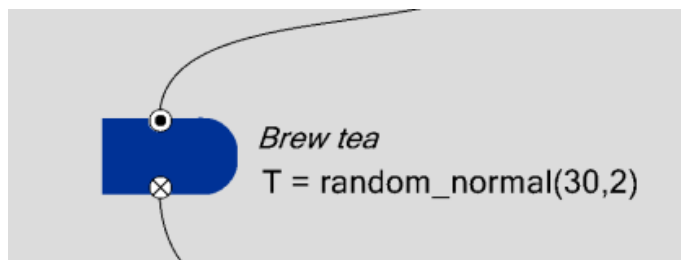


Now adjust the formula as follows :



You can use the keyboard or the buttons on the expression editor.

Click done twice and you should end up with this :

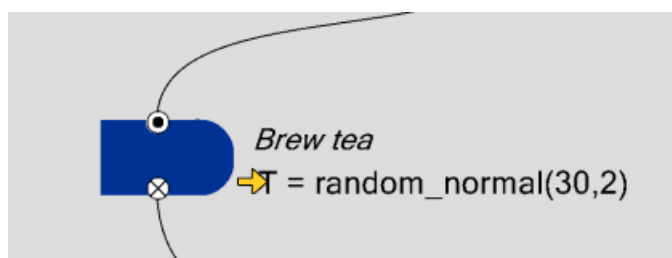


The time to brew the tea is a random number distributed according to the normal distribution with a mean of 0 and a standard deviation of 2.

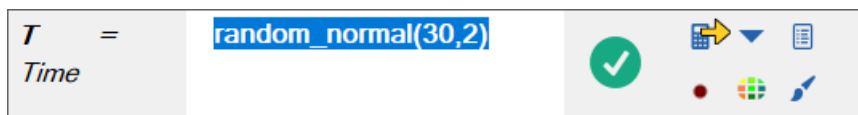
At this point we will discuss how to alter what is shown by an attribute display.

In this case we would probably want to see a value as well as a formula:

Click the attribute :



Click the expand button :



You now have the choice as to what you will see in the attribute display.

A screenshot of the 'Display' settings panel. It has a title 'Display:' on the left. To the right are several checkboxes: 'Name' (checked), 'Description' (unchecked), 'Expression' (checked), 'Choices' (unchecked), 'Timing' (unchecked), and 'Value' (unchecked). There is also a 'Suffix' label followed by an empty text input box.

The **name** is the short form name of the attribute – it must be unique within the block. In this case it is “T”

The **description** of the other hand could be “process time” . It does not have to be unique. Many of the library block attributes have description turned on and name turned off.

The **expression** is generally the driving mathematics - in this case “random_normal(30,2)”. Aliquis also allows text expressions.

Numerous attributes in library blocks are set up as a list of choices The **choices** button shows the text for the choice selected..

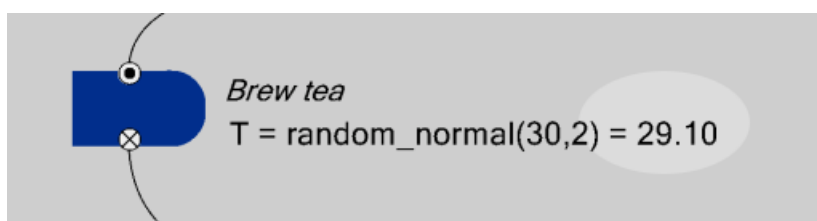
You can define when a variable is set within the solve (actually this is particularly useful for random numbers) . This choice can be displayed if **timing** is ticked.

The **suffix** is for you to add some text. So for example you may want to write the unit in here.

Finally you can display the **value** and the decimal accuracy of the value. This is what we want to do here:

A screenshot of the 'Display' settings panel, similar to the one above. In this version, the 'Value' checkbox is checked. The 'Precision' field, which is a spinner box currently showing the number '2', is highlighted with a light blue oval. The 'Timing' checkbox is also highlighted with a light blue oval.

Example Result :

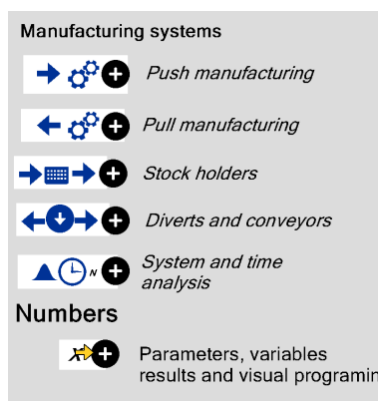


Method 3 using data arrays and array functions.

Suppose we are in a situation where we do not know how to model the length of time it takes to drink the tea but we do have a stop watch and we can see a number of people about to start drinking. We can get some measurements:

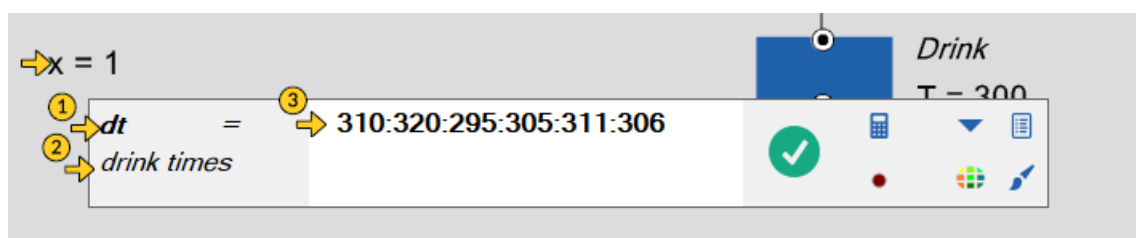
In Aliquis we can hold a group of numbers in an *array*:

Go into numbers



Drag a variable *x* into your model :

Edit the name, description and value of the variable as follows :

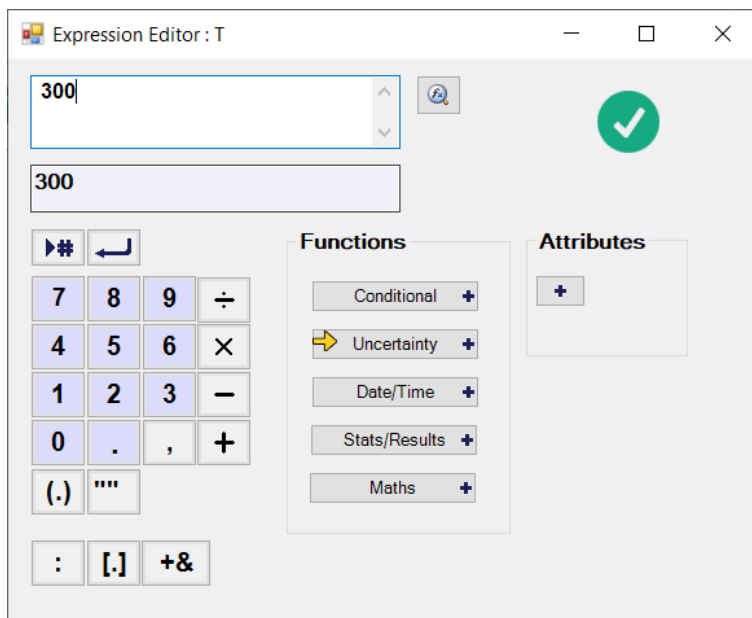
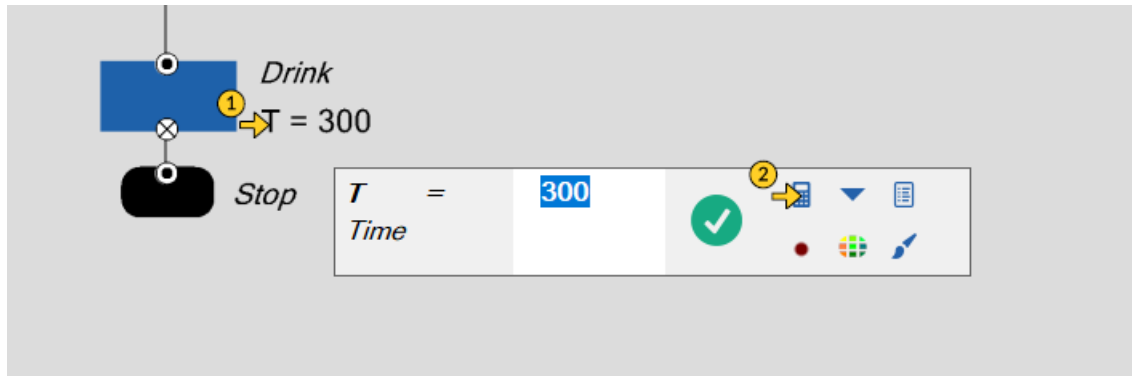


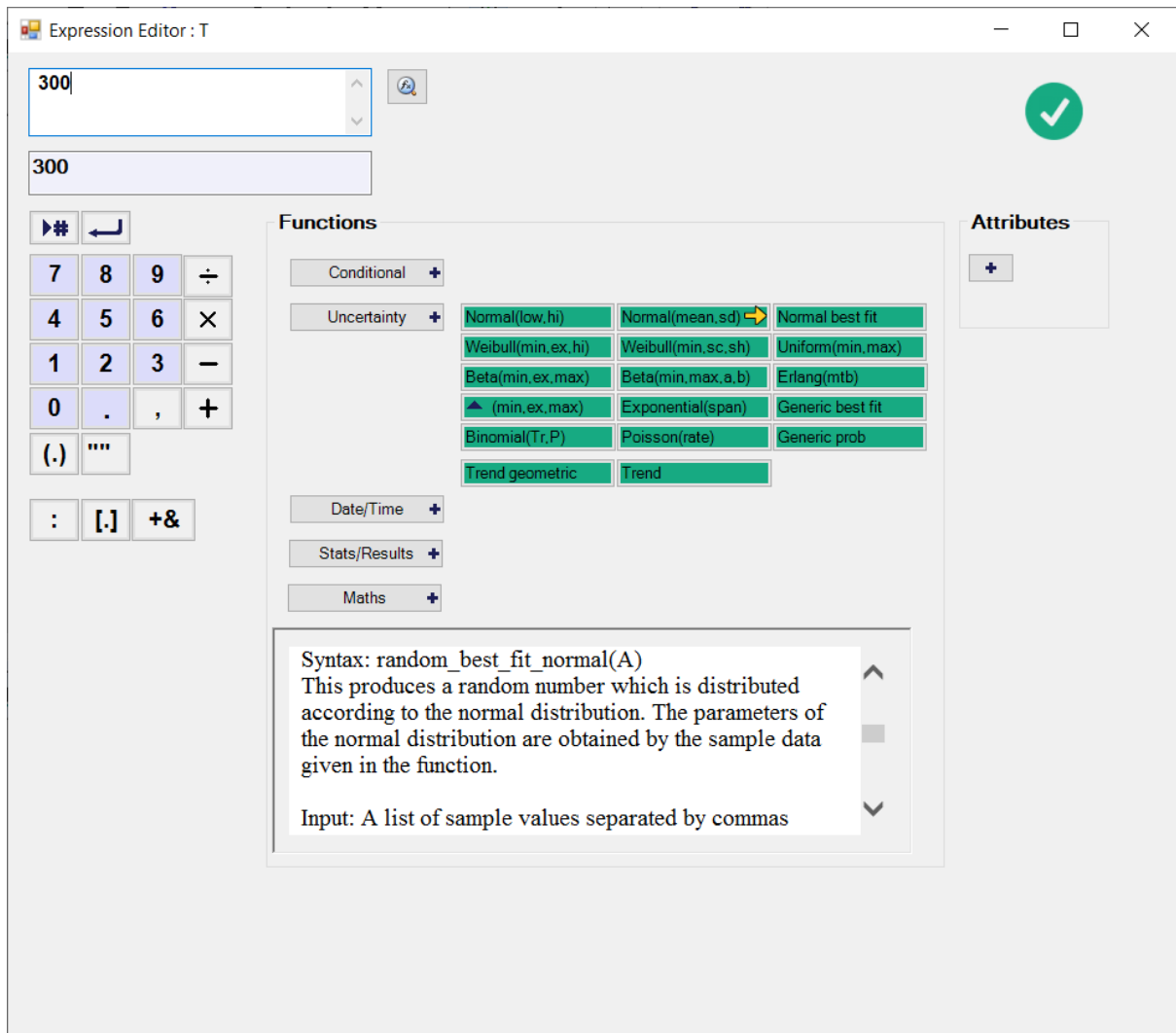
The numbers represent a sample of 6 measurements. In Aliquis numbers in an array are separated by ':'.

Now we will set the time to be a normal distribution based on these numbers :

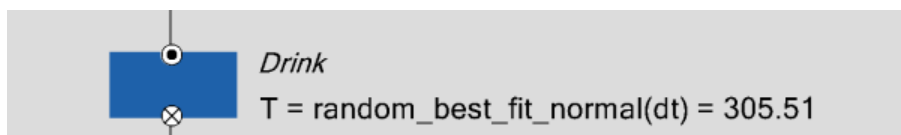
Edit the time parameter on the drink block as to

random_best_fit_normal(dt)



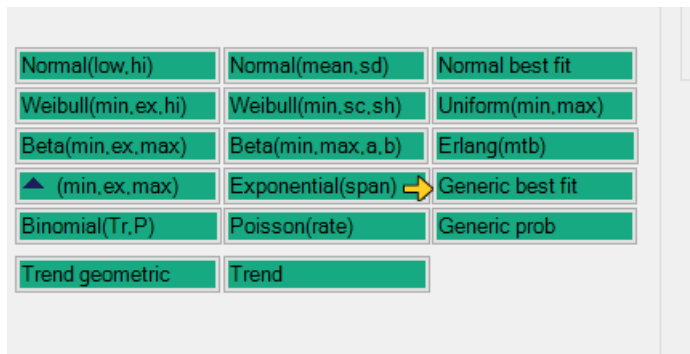


On your own, make sure the variable also displays its result :



This function interrogates the values in the variable *dt* (*drink times*) and internally works out the best estimate for the mean and standard deviation. It then computes a normally distributed random number. It is the best function to use for small data sets when we can safely assume the number is normally distributed.

It may be that your variable is not normally distributed. In this case we can use the Generic_best_fit function :



This function should be used if the data clearly shows a non-symmetrical characteristic. It loses accuracy when there are only a few data points as in this example.

Modelling Event based uncertainty

As we have already discussed, bad things happen. Sometimes good things happen too. You can model these *if* you can conceive them. You can't model them if you are totally unaware. So, it's clearly in every process designers interest to research and think about the possible events as much as possible and as early as possible.

It is then the job of a process designer to define which of these event-based uncertainties should be regarded as risks.

For these, we would need an estimate of the probability and some plan to deal with the consequences should the event occur

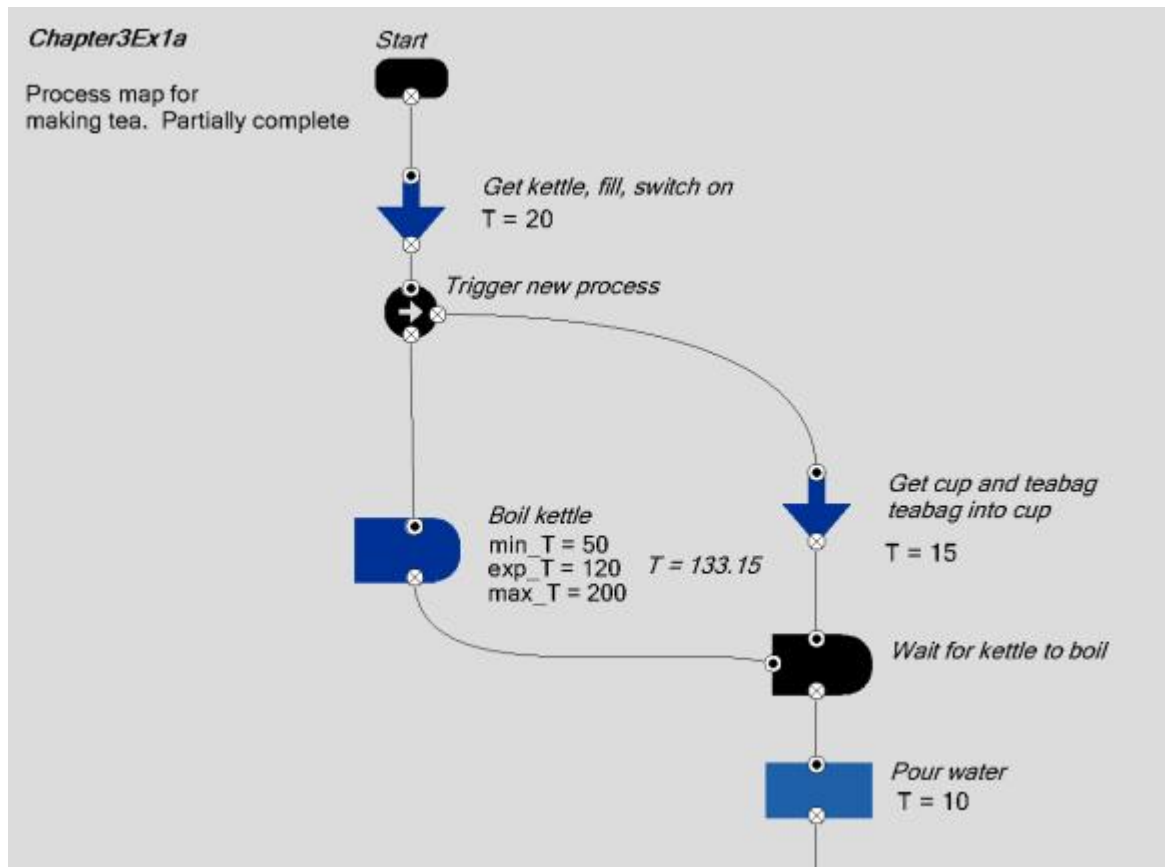
So, in our example let's suppose that one of the events while pouring the hot water into the cup is spillage. We could decide to view this as a risk. The risk *event* would be "spillage", the *probability* is the chance that it occurs, the risk **consequence** might be the need to get a cloth, wiping up, and make another cup of tea.

One of the problems in modelling these seldom occurring events is that it is extremely difficult to quantify the probabilities. In the real world you would have to run processes a disproportionate amount of time to get an accurate value for risk probability especially when they are low. Moreover, The probability value in most real world scenarios will probably vary over time and according to conditions. So, in summary this kind of analysis is nearly always done in the context of high levels of level 2 uncertainty. That said, this does not mean we should give up on modelling these events. It almost always yields insights which in turn can reduce unforeseen costs in time and money.

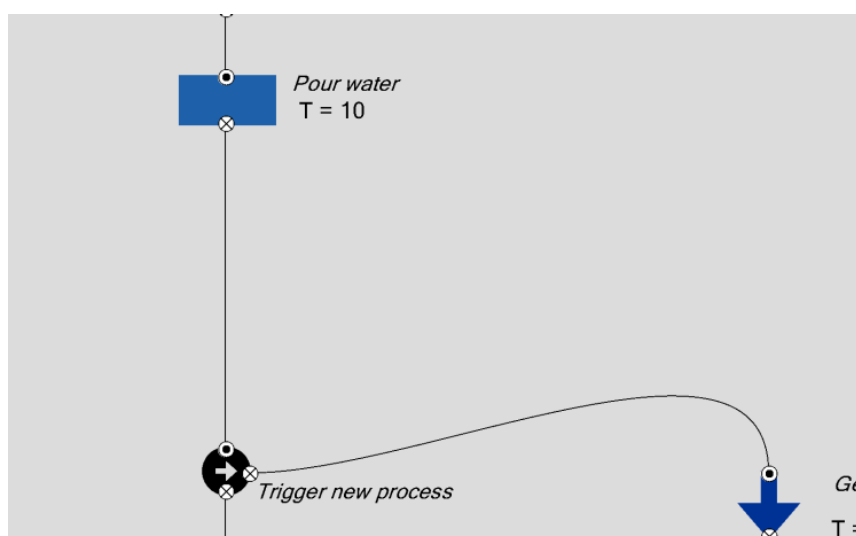
In manufacturing, this means reduced scrap, reduced accidents and reduced hidden costs!

Event uncertainty example

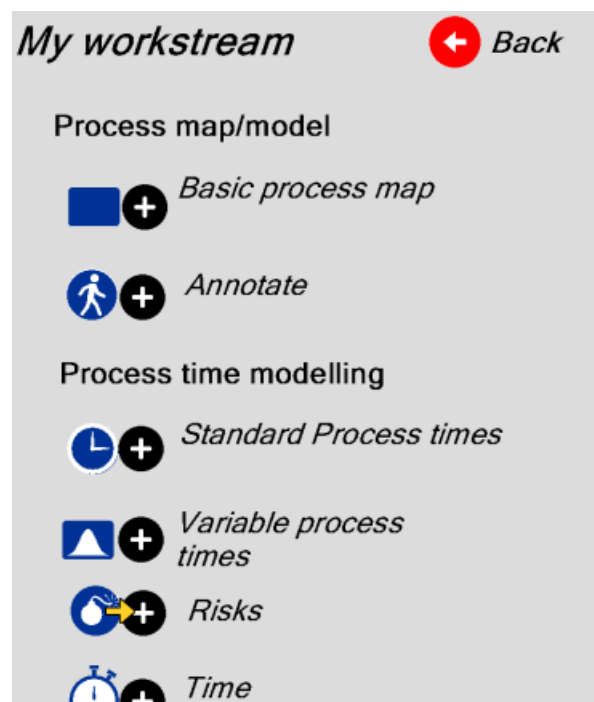
We will continue with our current model. We are going to develop a model for the spillage risk discussed earlier. The work you should have just done It can be found in *Chapter3Ex1a* in your model store if you need it.



We will model the risk here :



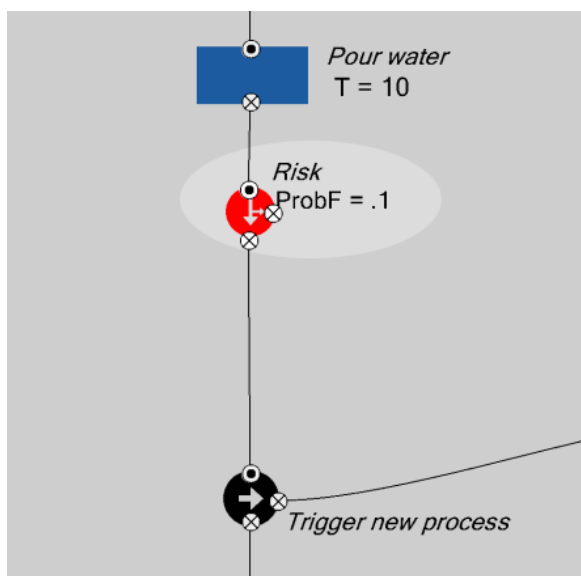
Go into the risks sub-library



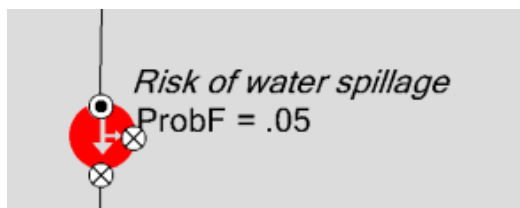
This is the block you need :



Drop it into the model here :



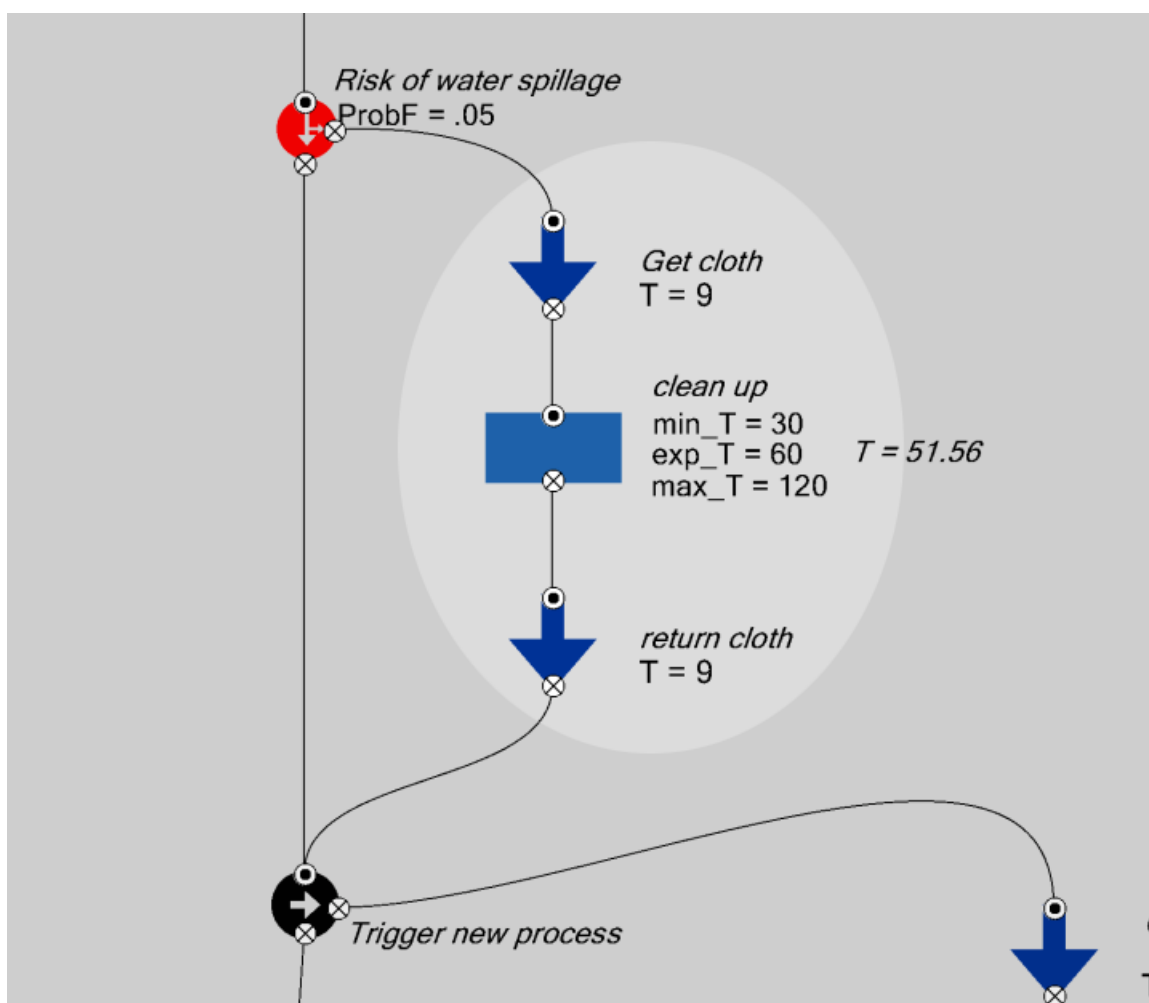
And set its name and the probability that the risk event will occur :



At this stage we have to think about the consequences if this risk event occurs.

Lets assume a simple amelioration plan can work the tea maker gets a cloth and clears up then returns the cloth.

Model the process as follows:

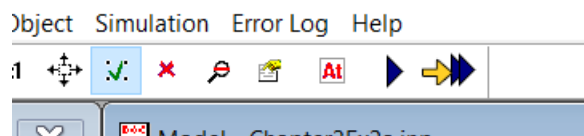


Working with a stochastic model – Monte Carlo simulations

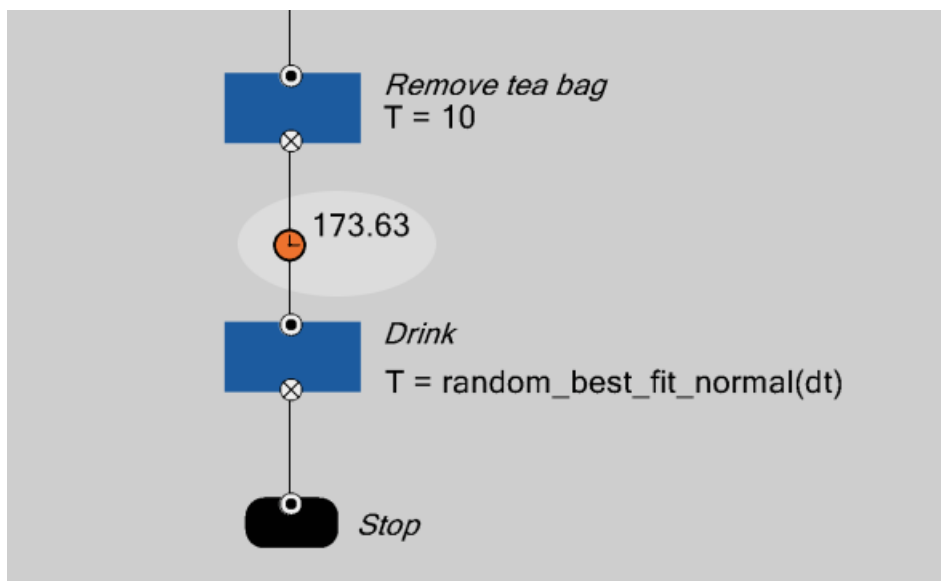
A *stochastic* model is a model that includes random behaviour (like most real world systems). You have just created one. Our aim now is to use it to gain insight into the real world process.

If you are not sure you have created the model correctly you can use *Chapter3Ex1b* from your model library as a start point.

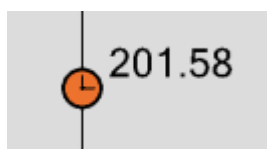
Solve the model once



Note the time to make the cup of tea:



If you repeat this you will get another number :



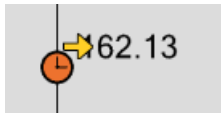
In fact its never the same!



So how do we make sense of these results? Well, if we solve it many times and record the results for each we can build up an idea of the statistical distribution of this result. Here's how to do it in Aliquis :

To record the values :

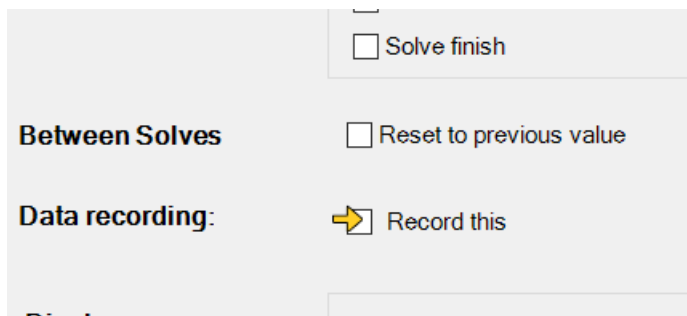
Click the value you see



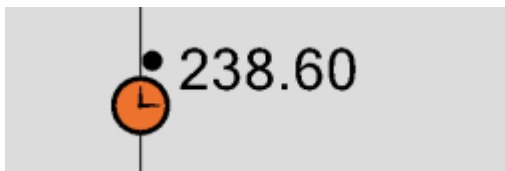
Expand the interface



Click "record this" :



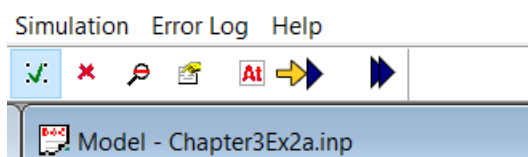
Result :



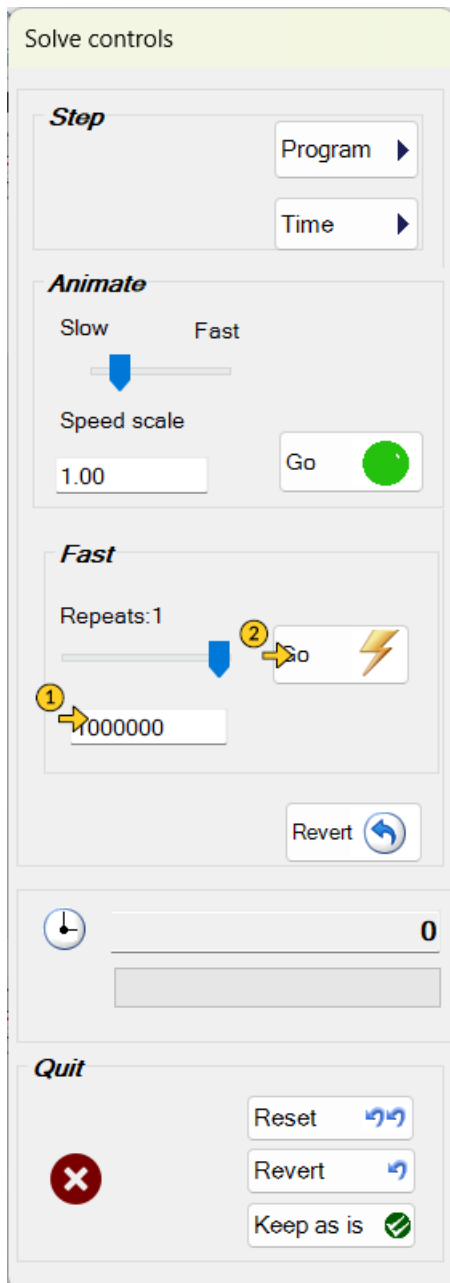
The black dot by the variable shows it is being recorded.

Solving the model multiple times :

Bring up the solver

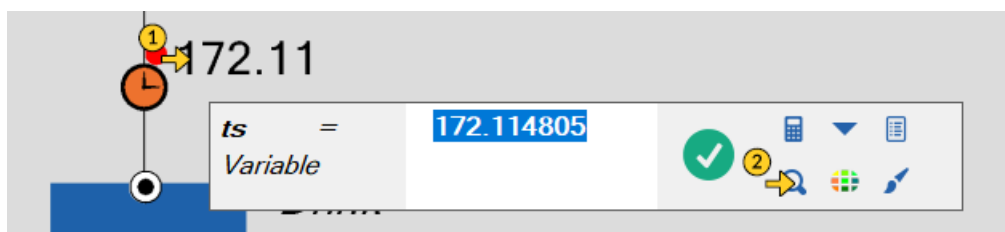


Under "Fast Solve" enter the number of solves (eg 10000) and "go"



When the solve has completed, close the solver

The red dot shows that there is now data associated with the variable. Pick the variable again and the magnifier :

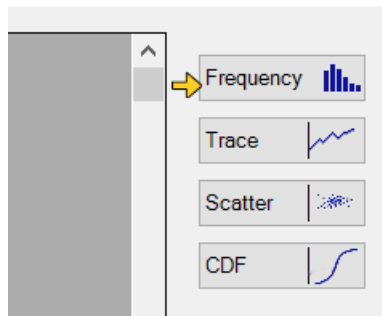


Result:

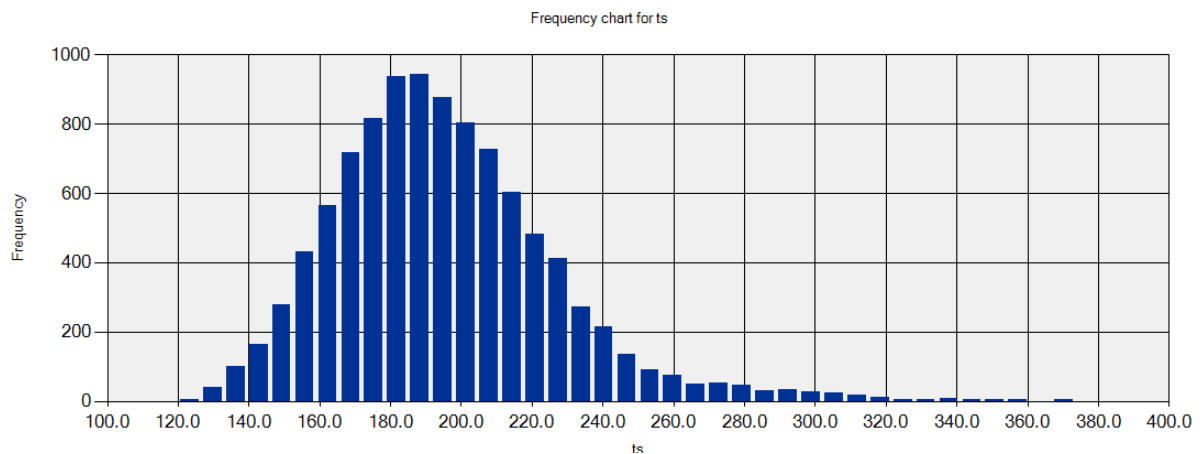
▶	167.158917...
	220.034088...
	245.972631...
	236.306391...
	190.794470...
	200.679410...
	199.114586...
	222.828918...
	175.602710...
	194.616396...
	138.749225...

These are the individual data items created by our 10000 solves.

To plot a histogram:

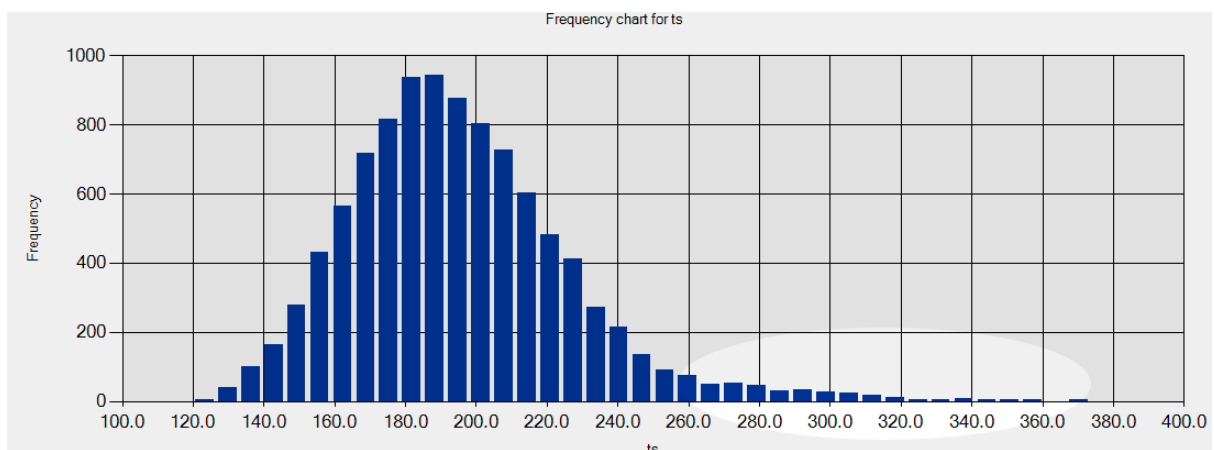


You should get something resembling this :



This is a histogram of our 10,000 results. It immediately shows that the best time was 120 seconds and the worst was 370. Moreover, it shows that the vast majority of times, we can expect to get our tea between about 130 and 320 seconds.

On close inspection, this is not a classic bell curve. It is skewed because of the risk of spillage.



The graph tool has a number of features –

You can see the average time here :

Analysis			
Mean	194.863204448838	Percentile	Value
Std deviation	31.0968894042494	5	150.630187777124
		Value	Percentile
		194.863204448838	54.52
			Sample size
			10000

But, for example, you may want to find out the estimated probability that the time result will be over 300:

Enter 300 here and note the percentile

The image shows a user interface for a simulation. It has three input fields: 'Value' with a yellow arrow pointing to it and the number '300' inside; 'Percentile' with the number '99.15' inside; and 'Sample size' with the number '10000' inside. A blue arrow points from the 'Value' field to the 'Percentile' field.

In this case we can see that there is 99.15 percent probability that the result will be less than 300. Put another way 0.85 percent probability that it will exceed 300. To get more accuracy you would probably need to put more trials into your monte carlo simulation.

Summary

This concludes our initial look into systems where variation and risk in a process are modelled. We have also defined uncertainty.

This approach to modelling gives results which are much more like the real world where times are rarely fixed and risks are active.

We can examine the results in terms of means, variation and probabilities – again all useful in the real world.

Additional notes

Before we move on there are a couple of statistics concepts to discuss....

Another definition : **Stochastic**. A stochastic model is one where elements of uncertainty and variation have been incorporated into the model. You, therefore, cannot predict the exact answers that the model will give.

For now, we will not explore deeply the statistical concept of **independent** and **dependent** variables. However, it is worth saying that the value of a parameter may include variation as discussed *and* some causal link with another parameter. So, for example, the mean time and the time variation of a process may both reduce over time as operators learn the process. The parameters are not independent of time

Key Engineering concepts :

- Variability, uncertainty and risk
- Monte-carlo simulations

Key Aliquis concepts:

- Various ways of adding uncertainty/variation in the model. (Library blocks, random number distributions and data based)
- How to execute a monte carlo simulation
- Reporting results – histograms and scatter graphs
- Reading an aliquis histogram

In chapter 4 we will look at analysing waste in processes. Chapter 4 will also give you a good understanding of how to splice basic maths into your simulations.

Questions :

1. What is the difference between a risk and an uncertainty?
2. Does the technical word “risk” used in this chapter correspond to the way it is used in every day speech?

Answers

1. Uncertainty means you don't know exactly what or going to happen next, or how long it will take, or how much it will cost etc. All processes are subject to uncertainty. A risk is “an uncertainty that matters”. This means that planners or engineers might want to study the uncertainty in detail and plan for events or conditions that could occur.
2. Risks in every day speech focus on negative outcome uncertainties. Risks in technical speech can include uncertainties with positive outcomes. In day to day speech we would call these opportunities.