

Links, Lags and Ladders

- the subtleties of overlapping tasks -

Logic in a Precedence Network

Precedence diagrams use boxes to represent the basic network elements - the task (or activity). Tasks have durations giving the period of time required to perform the work they represent and may have other descriptive data attached to them. The other key element of precedence networks is the dependency (or link), which defines the logical relationship between the tasks. A link is shown in a precedence network diagram as a line.

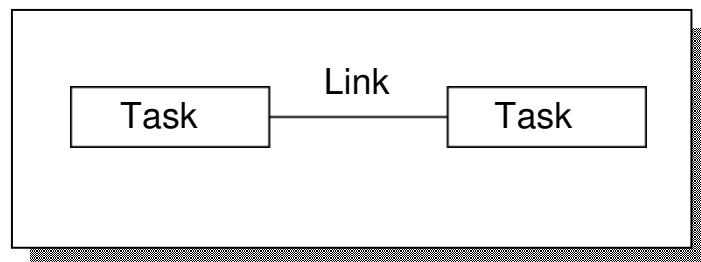


Figure 1 – Tasks and Links

Tasks are identified by a task identifier - for example, A1, A2, A3. Links are usually identified by their preceding task identifier and their succeeding task identifier.

The other element that should be included in every schedule is Milestones. Milestones are 'zero duration' events that mark significant points in the schedule such as its start and finish and are connected to other tasks and milestones with links.

Logic describes the flow of work

The relationships between the tasks define the flow of work through the project. The objective is to organise the tasks into a logical sequence agreed to by the project team. Only real logic should be used to construct the logic diagram (or network) using Finish-to-Start relationships where possible. Real logic can be:

- Dictated by the intrinsic nature of the work
- Mandated by the contract.
- External to the project representing either a deliverable required for the work or something the project has to deliver to a third party.
- A sequence of work that is an express intention of the project team

The first two options above are mandatory logic; the third is an 'external dependency', the last is discretionary logic; but they are all 'real'. Artificial logic inserted to fix a problem should be discouraged as it distorts the schedule and can have unintended consequences as the schedule changes during the life of the project.

Dependency Management

External dependencies require a different management approach to internal logic (discussed in the balance of this paper).

‘Outgoing’ dependencies represent requirements of other projects or an interim deliverable to the client. These are either a contractual requirement which represents a constraint that has to be achieved, or an obligation to assist the overall running of the organisations total project delivery effort. The receiver of the outgoing link is a stakeholder of the project whose needs are important and should be met wherever possible.

‘Incoming’ dependencies are a risk! They represent requirements the project needs to complete its work but the project team does not control the delivery process and the risk needs managing.

Dependency management requires a significant focus, including:

- The identification of the dependencies (at an appropriate level of detail);
- Mapping the dependencies into the schedule (we recommend highlighting each dependency with a milestone);
- Determining the way the dependencies will be technically mapped between projects (there are various software options – fully automated linking is not recommended);
- Determining how the progress on achieving incoming dependencies will be monitored and variances managed;
- Recording key risks in the risk register; and
- The on-going management of the dependencies as work progresses.

External dependencies are similar to the schedule start and finish date in terms of framing the overall project plan.

Developing ‘internal’ Logic

To determine what constitutes a logical relationship within the shedule the key questions to ask are:

- What has to be completed to allow this activity to start?
- What cannot start until this activity is completed?
- What can happen at the same time as this activity?

The resulting logic is a ‘road map’ showing the sequence of work from the beginning to the end of the project.

When this process is complete, every task and milestone should be connected to at least one predecessor and can trace its logical predecessors to the Start Milestone and at least one successor and can trace its logical successors to the Finish Milestone¹. *The Practice Standard for Scheduling*² recommends all activities are preceded by a ‘start’ link (ie, a link that connects to the start of the task) and succeeded by a ‘finish’ link (ie, a link that connects from the end of the task).

¹ See: **Dynamic Scheduling** - www.mosaicprojects.com.au/Planning.html#Core_Papers

² See: <http://www.mosaicprojects.com.au/Books.html#PMI> for details of the Standard

Links in a Precedence Network

As already mentioned, links dictate the flow of work through the project. There are four types of link referred to in the PMBOK. Finish-to-Start (FS), Finish-to-Finish (FF), Start-to-Start (SS) and Start-to-Finish (SF). Of the four standard links, FS links are most common and SF links are rarely used. Using any type of link other than FS can produce unexpected results during schedule analysis as they have not been consistently implemented by project management software developers (ref: 'Logical Inconsistencies').

Finish-to-Start Links

The normal type of link is a Finish-to-Start link (FS). With this type of link, the succeeding task cannot start until after the finish of the preceding task.

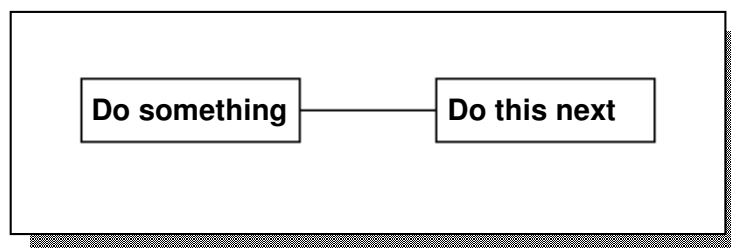


Figure 2 - Finish to Start Link

If a lag time is specified on the link (say 3 days), the succeeding task cannot start until three days after the finish of the preceding task.

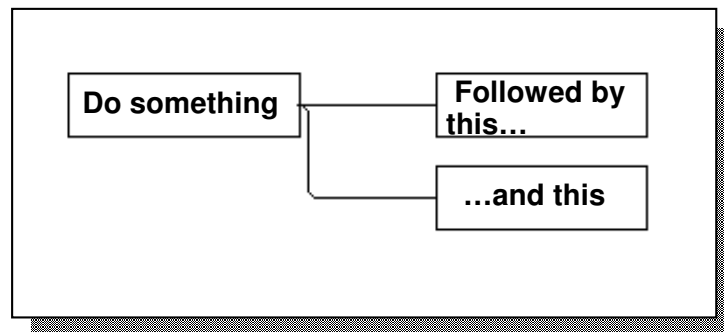


Figure 3 – Succeeding Tasks

Links work independently. In Figure 3, neither of the following tasks can start until after the leading task is finished BUT they do not have to start at the same time and they do not have to proceed together.

Finish-to-Finish Links

Finish-to-Finish links (FF) constrain the completion of a task. The completion of the succeeding task is delayed until after the completion of the preceding task. If a lag is nominated (say three days), the finish of the succeeding task is delayed until three days after the finish of the preceding task.

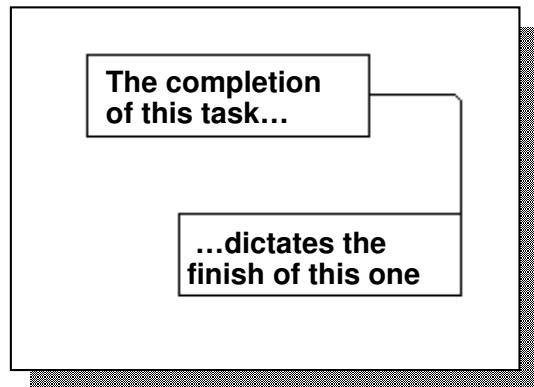


Figure 4 – Finish-to-Finish Link

This type of dependency primarily controls the finish of tasks (not the start). A typical example would be writing and editing a book. The editor does not have to wait until the writing is finished to start the editing process; editing could start as soon as the first chapter is finished. BUT, it is impossible to finish editing until after the writing is complete. The editor may require a week to complete the editing once the book is finished and this is represented by creating a Finish-to-Finish link with a lag of 5 days.

Start-to-Start Links

Start-to-Start links (SS) constrain the start of a task. The start of the succeeding task is delayed until after the start of the preceding task. If a lag is nominated (say three days), the start of the succeeding task is delayed until three days after the start of the preceding task.

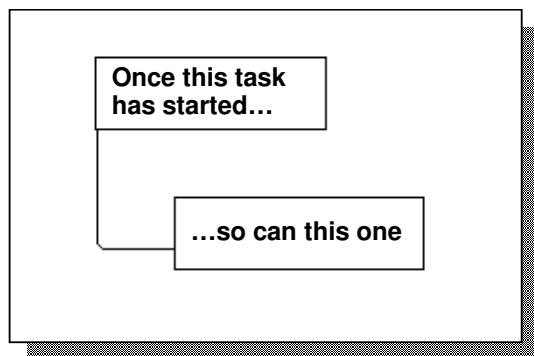


Figure 5 – Start-to-Start Link

This type of dependency primarily controls the start of tasks (not the finish). Staying with the writing and editing of a book, it is also impossible for the editor to start editing until some of the writing is complete (maybe the first chapter). The author may require two weeks to format the overall plan for the book and write the first chapter. This is represented by creating a Start-to-Start link with a lag of 10 days.

If you need to control both the start and the finish of the relationship between two tasks (as would be the case with writing and editing), it is best to insert both links between the tasks (SS and FF). If this is not possible (some software will only allow one link), then you must decide which link is most important (see: *Managing the Overlap* below).

Start-to-Finish Links

Start-to-Finish links (SF) constrain the finish of a task based on the predecessor starting. The finish of the succeeding task is delayed until after the start of the preceding task. If a lag is nominated (say three days), the finish of the succeeding task is delayed until three days after the start of the preceding task.

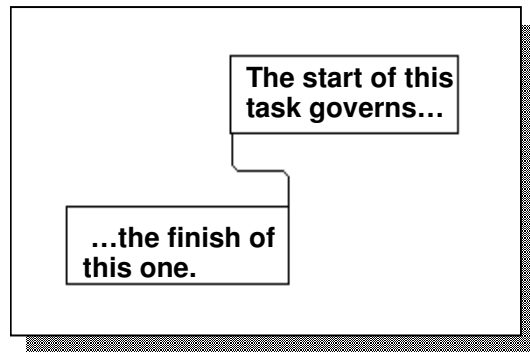


Figure 6 Start-to-Finish Link

This type of link is used to control the change over between two processes, if a business is changing from a security system that uses key cards for access to one that uses bio-metrics, the use of the key card system cannot finish until after the start of the bio-metric system. If both systems are required to run in parallel for a time, a lag is added to the S-F link.

Leads and Lags

As described above, a 'positive lag' has the effect of delaying the succeeding task by the number of time units specified.

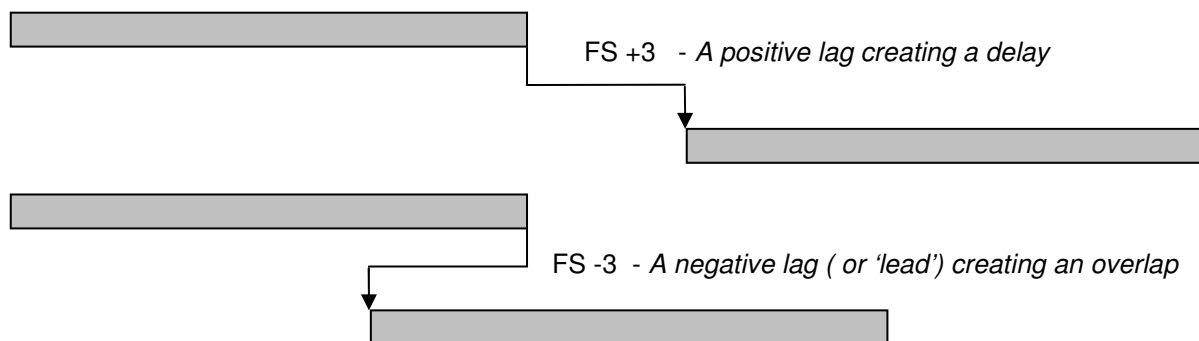


Figure 7 Leads and Lags

Negative lags (or 'leads') have the effect of accelerating the succeeding task by the number of time units specified. Consequently, if the lag value is specified as a negative number, it has the effect of overlapping the tasks. A lag of - 3 days on a F-S link would mean the succeeding task can start 3

days before the end of the preceding task (ref: Fig. 7). Negative lags (or leads) are allowed in some software packages but need to be used with care³.

Lags should not replace work (logic). Even where work is to be performed by others, this work should be included as a task. For example, if the contract allows one week for the review of a drawing by the client; do not insert a lag of 5 days on the link between the task for creating the drawing and the task for using the drawing (both your work). Rather, insert a 5 day task for the client review; this task can then be coded and reported upon during status updates of the schedule⁴ and any delays properly attributed to the responsible party.

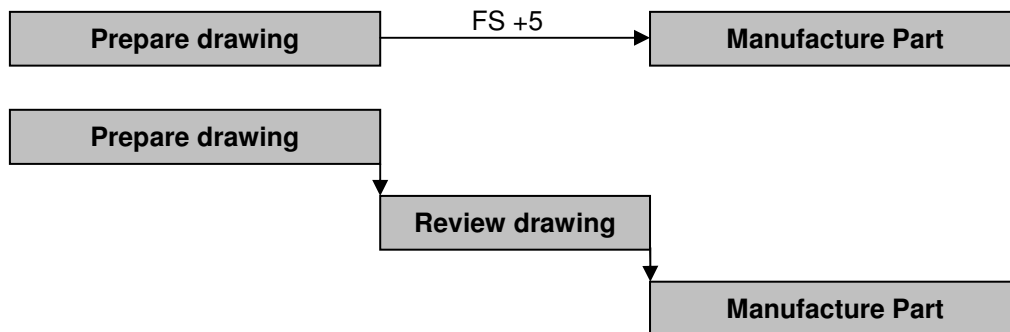


Figure 9 - Lags should not replace logic

If the time between the activities is needed for a purpose, but no work is happening (eg, concrete setting time or paint drying time) a FS lag is appropriate and the 'space' has a purpose. However, Lags should not be used simply to create a space between two activities 'for convenience' or to make the schedule look correct. These '*leaps of logic*'⁵ bypass true network logic by linking tasks with inherent gaps in time between the activities and can be misleading and may cause computational errors when used; the effect is similar to putting artificial constraints in the schedule and should be discouraged.

Managing the Overlap

Where inserting an additional task is not appropriate and the gap is 'real', the nature of the gap needs to be clearly understood⁶: Why is this lag needed?

- Does the time represent an imposed delay to create a sensible flow of work allowing the leading task to clear sufficient work space for the succeeding task to commence within?
- Does the time represent administrative works needed to prepare for the succeeding task?

³ From a logical perspective a negative lag is difficult to justify and its use is discouraged by most professional schedulers. In most circumstances the combination of SS and FF lags can achieve a more sensible overlapping of activities. However, as a number of limited tools only allow a single link between activities, the concept of a 'Lead' (or negative lag) is retained in this paper.

⁴ See: **A Guide to Scheduling Good Practice** - www.mosaicprojects.com.au/Planning.html#Core_Papers

⁵ Term developed by Jim Peter and Kelvin Murray to describe this effect.

⁶ See: **Faster Construction Projects with CPM Scheduling**, 'Anatomy of a relationship' page 177. Details of book at <http://www.mosaicprojects.com.au/Books.html#books>

- Does the time represent a productive work segment (Ref: Fig. 10 & 11) where a certain amount of work has to be completed on Task A before Task B can start to use the handed over work?

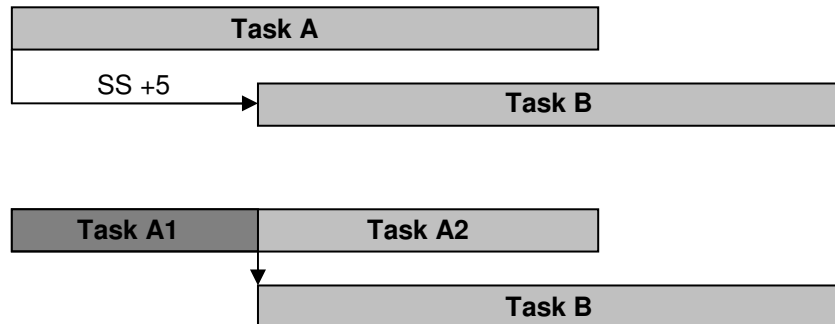


Figure 10 SS Link = Productive work segment

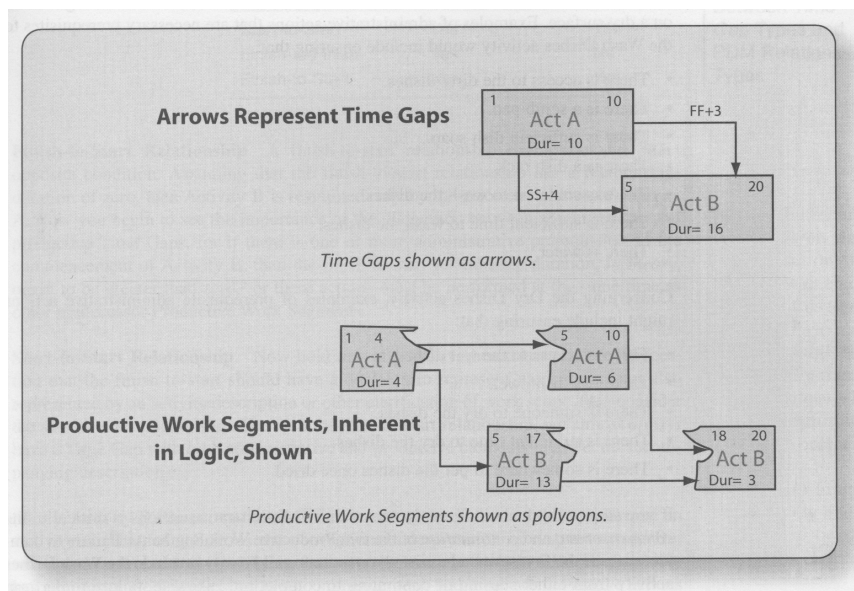


Figure 11 - Extract from Woolf's Book⁷

Understanding the nature of the relationship is critical to effectively managing the schedule; anecdotal evidence suggests most of the minor delays that are the responsibility of the project team (ie the contractor) occur in the gaps between tasks represented by lags. In aggregate these delay can have a major impact on the momentum of the project and cause delays to completion.

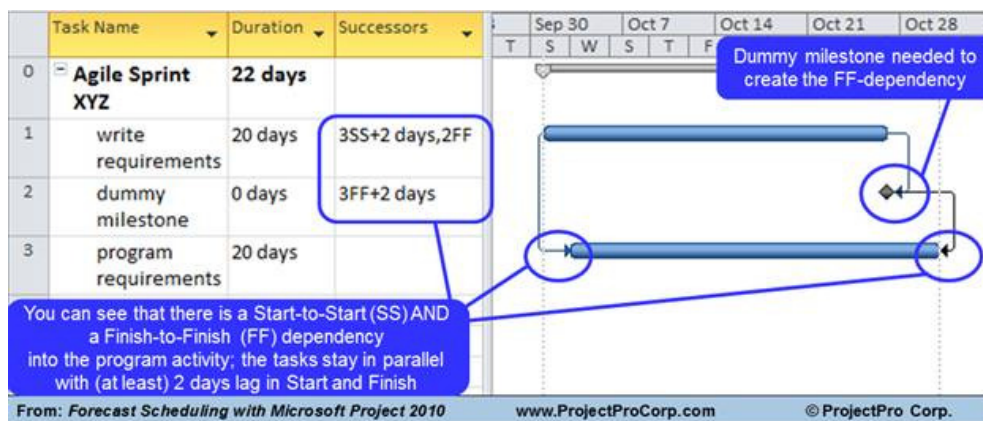
Where only one link is used the next question is does the remaining part of Task A have any influence on Task B? In the case depicted in Fig. 10, there is a high probability that all of the work in Task A has to be completed to allow Task B to finish, but this is not necessarily the case. However, if there is a need for Task A to continue to feed work to Task B our strong recommendation is to either:

⁷ Woolf, M.B. (2007) *Faster Construction Projects with CPM Scheduling*. McGraw-Hill, New York.

- Set the link type to 'progressive feed'; a number of tools have this feature. Progressive feed only allows B to progress proportionally to A.
- Use both a SS and a FF link to at least constrain the start and finish of B in relation to A.

If your current tool is incapable of either and you want to develop useful schedules that produce predictable results during the progress of the works either, stick exclusively to Finish-to-Start links, buy a software tool that works or add some additional logic to simulate the effect.

The problem with inserting dummy logic (as per the example below) is the tool cannot manage the dummy relationship and milestone – you have to do the work. Failing to remember the 'dummy milestone' will sooner or later cause an error in your updating.



Dummy logic is necessary in some unsophisticated tools.

Logical Inconsistencies

As previously mentioned, the use of links other than Finish-to-Start can cause unexpected problems. Fig. 12 represents the dry walling work on Level 5 of a high rise block of units (one complete floor):

- Task A is the erection of the framing. This 10 day activity involves 2 days to set out the walls and fix the head and floor tracks and 8 days to fix the rest of the studs and frames
- Task B is the in-wall services rough-in. This involves a total of 3 days work by electricians, plumbers and others to run their pipes and cables inside the wall ready to connect to fixtures and fittings at a later date. This task can start 4 days after Task A has started (this allows time for the framers to have installed around 25% of the studwork) but cannot finish until 1 day after all of the framing is installed. By its nature this work is intermittent requiring several short visits to the floor by each of the services trades.
- Task C is the fixing of the wall sheeting. This can start one day after the 'in-wall services rough-in' has started and needs 3 days to finish after the last of the services are installed in the wall. The three days allows sufficient time to fix the last sheets, finish setting the joints and on the final day complete the sanding of the joints. However, fixing, setting and sanding the wall sheeting will take 12 days overall. Progress on the wall sheeting is only partly dependent on the in-wall services because not every wall has services inside it and as long as

the service trades have access to one side of the walls where there are internal services, the sheeting can be installed on the other. The sheeting also needs at least 3 days after the completion of the framing (Task A) before it can finish.

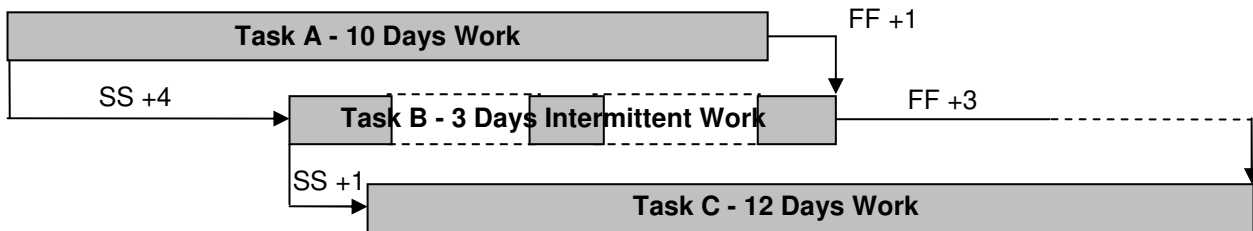


Figure 12 - Wall Framing Level 5

The situation in Figure 12 represents the optimum situation. Task B starts 4 days after Task A allowing Task C to start one day later. Task B finishes 1 day after Task A allowing Task C to complete 12 days after it started. The overall duration of this work is 4 days at the start of Task A, plus 1 day at the start of Task B plus the full 12 days for task C equalling 17 day work.

The calculation of Float in this situation is interesting! Only the first 4 days work of Task A are actually critical, and only the first day's work of Task B is critical. Looking at the completions, Task B can finish on Day 11 (10 days work on Task A plus one day to finish off Task B). However, Task B has a Finish-to-Finish relationship to Task C of FF+3. This means Task B does not have to finish until Day 14, which would still allow the 3 days (day 15, 16 and 17) needed to complete the wall sheeting. Given Task B can finish on Day 11, but its finish could be delayed until Day 14, and this delay will have no effect on any other work, arguably the completion of Task B has 3 days Free Float (but not the whole task). A similar conundrum exists with Task A; it can finish up to 3 days late and will only delay the finish of Task B which has 3 days float.

From the 1960s through to the 1980s, (and particularly with Activity-on-Arrow networks) float was dealt with in a far more sophisticated manner than today's simple calculation of Free Float and Total Float. The range of float options is set out in Fig. 13 and many of these ideas have been incorporated in the new scheduling methodology, RD-CPM™, the Relationship Diagramming variation of the Critical Path Method⁸.

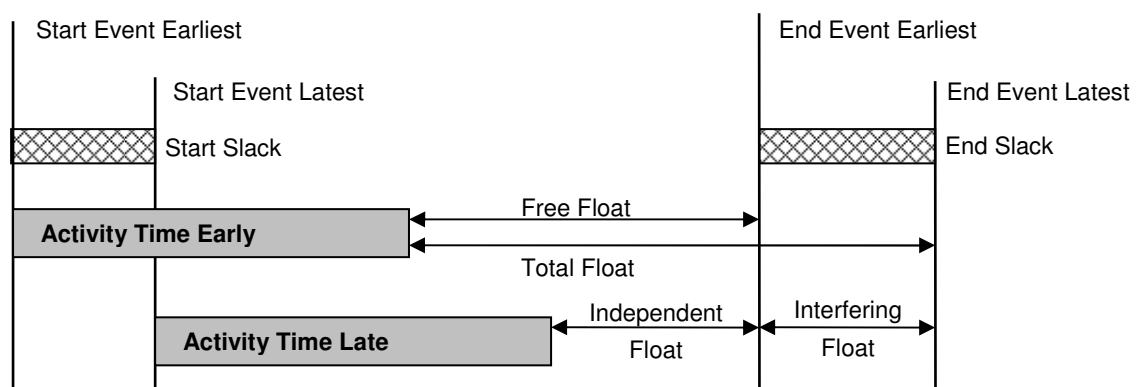
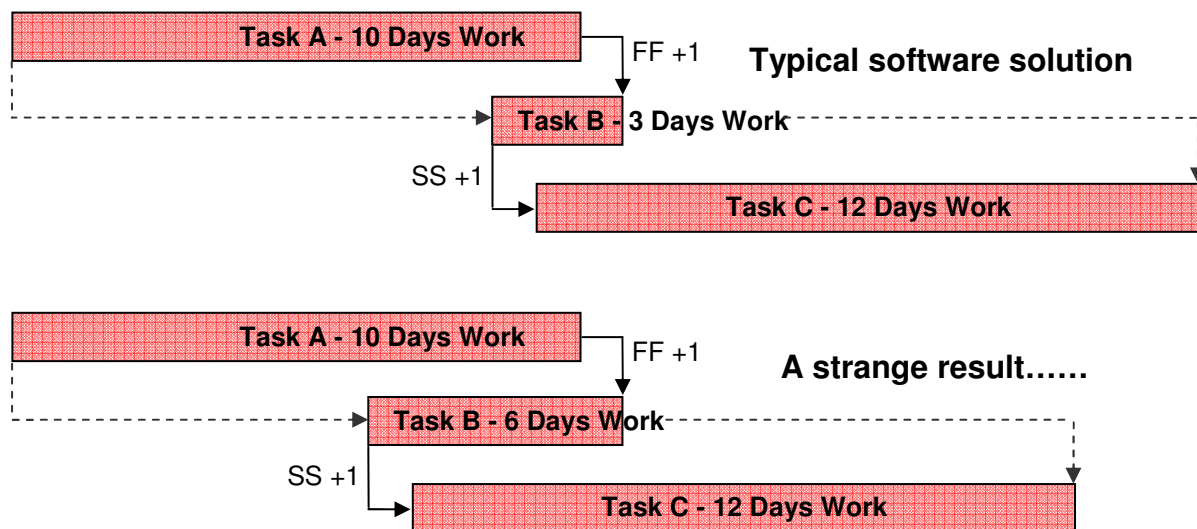


Figure 13 - Types of Float⁷

⁸ For more on RD CPM™ see: <http://www.fplotnick.com/rdcpm/>

The calculations in a standard Precedence network should assess the situation at the start of the activity (the Start Event) and the completion of the activity (End Event). All of the above ‘floats’ have relevance in efficient resource levelling algorithms, unfortunately they are rarely considered⁹.

Unfortunately, very few of today’s software tools will resolve the situation in Fig. 12 satisfactorily. Most will resort to the solution in Figure 14; delaying Task B to comply with its finish link and schedule ‘B’ from Day 9 to Day 11. The consequence of this is to push the start of Task C to Day 10 and the end of the three tasks to Day 21. This effect is described as ‘lag drag’. Paradoxically, in this situation the whole of Task B is critical, but increasing the duration of Task B actually reduces the overall time for the three tasks to complete.



Increasing the duration of ‘critical’ Task B reduces the overall duration of the work!

Figure 14 - Some typical software induced problems

Ladders

The ladder technique was invented in the UK by ICL in the early 1960s¹⁰ (now Fujitsu), and gained wide acceptance in scheduling tools developed in the UK, the concept is still a key part of the scheduling algorithms used in the Micro Planner range of software¹¹.

Activity-on-Arrow diagramming became complicated when projects had multiple resource types and multiple identical activities usually differing only in their physical location. To keep the correct logical relationships most of the nodes had to be split by using ‘dummy’ arrows. In a ‘ladder-feed’ diagram for a pipeline or roadway segmented into discreet sections, there could be as many

⁹ For more on **Schedule Float** see: http://www.mosaicprojects.com.au/PDF/Schedule_Float.pdf

¹⁰ ICL 1500/4 PERT included Ladders on its release in May 1963. The documentation suggests Ladders were part of the 1500/3 PERT program (1962) with only minor improvements in the /4 release.

¹¹ For more on **Micro Planner** see: <http://www.microplanning.co.uk> or <http://www.microplanning.com.au>
The assistance of Micro Planning International’s Raf M. Dua in providing information on Ladders is acknowledged.

‘dummy’ arrows as work activity arrows. The use of the logic-splitting ‘dummy’ arrows had to be precise. Figure 15 is an edited version of this type of schedule and for each double node [OO] there is also a logic-splitter ‘dummy’ arrow, [O→O] that is not drawn.

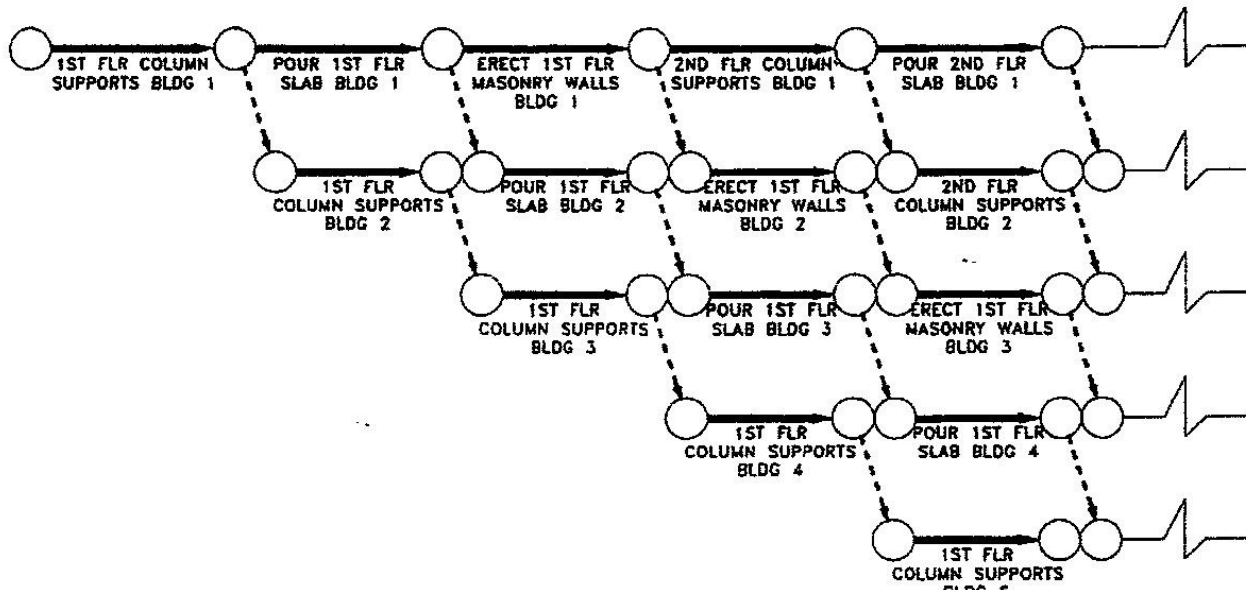


Figure 15 – A typical progressive feed problem

Precedence diagrams are not much better; using normal links, SS only controls the start relationship, FF only controls the finish relationship and whilst combining SS and FF provides the best control, only the ends (or start and finish events) of the tasks are linked and problems similar to the one defined in Fig. 14 above can easily occur.

Ladders are different! The concept of a ‘Ladder’ moves the management of overlapping activities forward to incorporate the idea of ‘progressive feed’.

Ladder activities were developed as a special group of activities that are used to represent progressive feed tasks. An example of a progressive feed task occurs in the manufacture of a number of identical components, each component having to go through several processes such as manufacturing, assembly and testing. To represent these processes in a network in the normal way would require one activity for the manufacture of each component, another to assemble the unit, probably another for inspection, etc. The same sequence of activities would have to be repeated for each unit. The resulting network could be extremely complex; ladders simplify the representation of the work.

Rung activities are the various tasks to be undertaken with defined durations, resource requirements, etc but designated as a ‘rung’ type of activity. The leads and lags are special activities specified with reference to the rung activity from which they originate. Before the second task in such a progressive feed process can start, the first task must have been in progress for a given time to ensure a supply of components for the second task. The time that must elapse before the second task starts is called lead time. Similarly, there is a lag time after the completion of the first time before the second task can be completed.

A ladder in an ADM network, and a representation of the 3 components in a time scaled network:

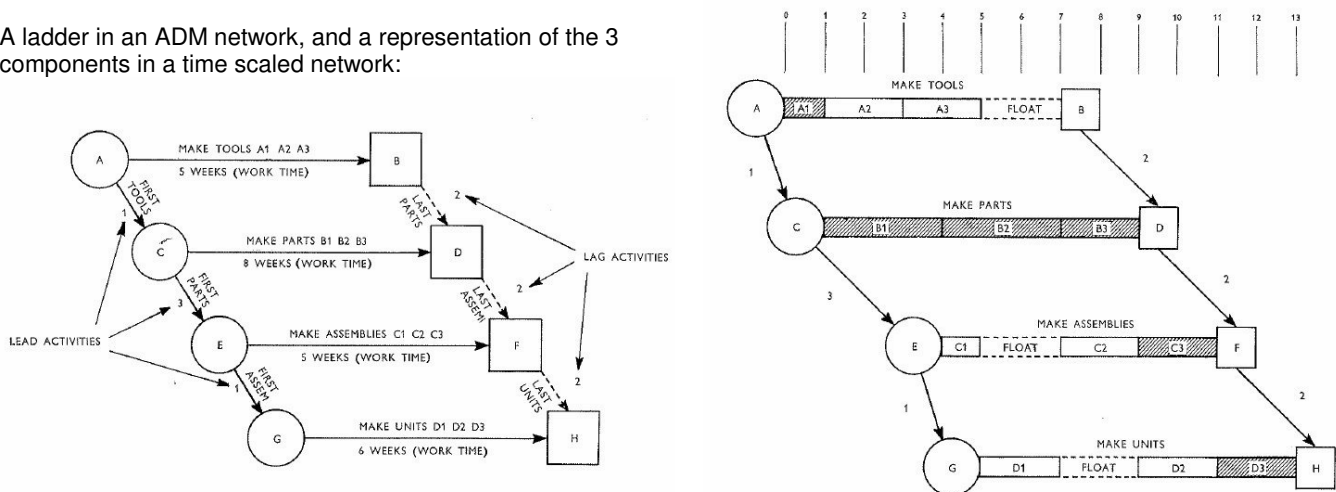


Figure 16 - A Ladder

This is similar to the operation of SS and FF links, however, from an analytical viewpoint, the major advantage of a ladder is if work stops on one rung, the delay is automatically flowed through to the work on all of the dependent rungs, not just the end event.

Progressive Feed Links

The Metra Potential Method (MPM) introduced a number of additional link types that can now be seen in some advanced PDM network tools. These links use the concept of progressive feed in the same way the ADM 'ladder' described above. Depending on the tool, the degree of overlap between two activities can be managed based on either a percentage complete or a set duration. In both cases, the leading activity must stay the designated amount in front of the succeeding activity and if the lead activity stops (eg, as a consequence of resource analysis), the succeeding activity stops as well.

- ACOS+1¹² uses the AP link type, AP=3 means the succeeding task cannot start until 3 days after the start, and cannot finish until 3 days after the completion of the predecessor.
- Deltek Open Plan¹³ allows percentage lags on all link types. The leading task needs to maintain the specified percentage completion ahead of the successor. A 20% lag means that if the predecessor is 60% complete, the maximum completion on the successor is 40% (it may be less but cannot be greater).
- Spider allows the concept of a 'Volume Lag', in pipeline construction trench excavation shall be done before lowering pipes but these activities can be done in parallel as long as the trenching crew and the lowering crews work at certain distance from one another. This is typical laddering relationship - both SS and FF are needed. This relationship is physical: the distance between crews shall be no less than 500 meters and so the SS and FF links both have 500 meters lags.

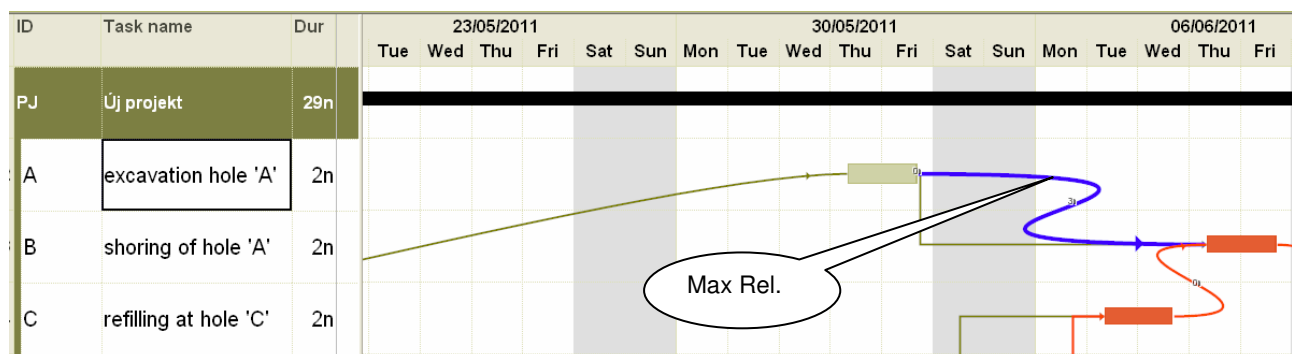
¹² ACOS+1 see: <http://www.mosaicprojects.com.au/Tools.html#ACOS>

¹³ Open Plan see: <http://www.mosaicprojects.com.au/Tools.html#OpenPlan>

The precise way these capabilities are incorporated into various tools differs. Planners and schedulers to be fully aware of precisely how the options function before using them.

Maximum Links

The Metra Potential Method (MPM) also allows the concept of a 'Maximum' relationship. Maximum relationships maxSS, maxFS, max SF, and max FF. Force the following activity to start within a defined period of time after the predecessor. An example would be responding to the people who contributed to a customer survey. After the thank you mail out cannot be sent until after the completion of the survey, by using a maxFS 5day link, the 'thank you' can be sent as soon as the survey is completed or at any time up to 5 days after the survey. But if it has not already started, the 'Send thank you' activity will be forced to start on the 6th day. These links are included in the ACOS⁹ system and other European tools based on MPM.



The effect of the blue 'maximum' in the network above is to pull the start of the 'excavation' activity back nearer to the availability of the shoring which is being transferred from 'Hole C'¹⁴

¹⁴ Example provided by Hajdu Miklós, Faculty of Civil Engineering and Architecture, Budapest University.

Hammock Activities

The 'Hammock Activity' is a cross between a link and an activity. The duration of the 'Hammock' is derived from the time between its start connection and its finish connection (it has no predetermined duration) but the hammock can have descriptions, codes and other attributes of a normal activity. Hammocks are very useful for carrying time related costs and determining the duration of supporting activities and equipment needed for a project.

The example I use when teaching is the time the tower crane is needed on a high rise construction project. The start of the crane working on-site is driven by the concreting of the foundations and erection of the crane. It is then required through to the time the last heavy lifting to the roof is finished (typically roof mounted plant and equipment) once this activity is finished it can be removed. The duration of the hammock is derived from the timing of these two events and is calculated automatically by scheduling tools that implement hammocks correctly.

Many software tools that do not have the capability to implement Hammocks and to hide the deficiency confuse a 'hammock' with either a 'Level of Effort' or a 'Summary' task.

Summary tasks are part of the logic structure and summarise lower level tasks within a coding system. Hammocks are not dependent on any coding structure.

The benefit of a 'Hammock' over a Level of Effort (LOE) task is the Hammock's duration is flexible and automatically adjusts to changes to the underlying logic in the schedule, whereas LOE activities have a set duration that requires manual adjustment if the project changes.

Conclusions

The developer of the PDM networking methodology, Dr. John Fondahl, was always of the view the only safe link to use in a precedence schedule was the Finish-to-Start link. Similar warnings are contained in the *PMBOK® Guide* and the *PMI Practice Standard for Scheduling*.

The issues raised in this paper clearly demonstrate the inconsistencies and problems that can develop using S-S and F-F links. However, it is highly unlikely their use will diminish significantly. Therefore, the responsibility must fall to the managers of schedulers, and the schedulers themselves to make sure the logical constructs used in their schedules are both sensible and mathematically correct.

This paper is a core reference in our PMI-SP and PTMC credential courses.

For more information on these courses see:

http://www.mosaicprojects.com.au/Training-Planning_One-on-One.html



The papers in this series:

- A Guide to Scheduling Good Practice: http://www.mosaicprojects.com.au/PDF/Good_Scheduling_Practice.pdf
- Attributes of a Scheduler: http://www.mosaicprojects.com.au/PDF/Attributes_of_a_Scheduler.pdf
- Dynamic Scheduling: http://www.mosaicprojects.com.au/PDF/Dynamic_Scheduling.pdf
- Links, Lags & Ladders: http://www.mosaicprojects.com.au/PDF/Links_Lags_Ladders.pdf
- Schedule Float: http://www.mosaicprojects.com.au/PDF/Schedule_Float.pdf
- Schedule Levels: http://www.mosaicprojects.com.au/PDF/Schedule_Levels.pdf
- Schedule Calculations: http://www.mosaicprojects.com.au/PDF/Schedule_Calculations.pdf

Additional information; see Mosaic's Scheduling Home page at: <http://www.mosaicprojects.com.au/Planning.html>