

ELIMINATING POTENTIAL PROCESS HAZARDS

It is impossible to eliminate all hazards in a plant, but for safety's sake all potential ones must be determined. Then we must establish the most important so as to know which to attack first. Here are ways to do this.

Trevor A. Kletz, University of Technology, Loughborough, U.K.

The techniques for identifying hazards (finding out which hazards are present in a plant or process) and the techniques for assessing those hazards (for deciding how far one ought to go in removing the hazards or protecting people from them) are often confused. Fig. 1 may help to make the differences clear.

The top shows some of the methods used for *identifying* hazards and problems that make operation difficult.

Some problems are obvious. If we manufacture ethylene oxide by mixing oxygen and ethylene close to the explosive limit, we do not need a special technique to tell us that if we get the proportions wrong there may be a big bang.

The traditional method of identifying hazards, in use from the dawn of technology until the present day, was to build the plant and see what happened. The old adage says, "every dog is allowed one bite"; until the dog bites someone, we can say that we did not know it would. This was not a bad method when the size of an incident was limited, but is no longer satisfactory now that we keep "dogs" that may kill many people at one bite.

Checklists are often used to identify hazards, but their disadvantage is that any missing items are not brought up, so that our minds are closed to them. Checklists may be satisfactory if there is little or no innovation, and all the hazards have been met before; they are least satisfactory when the design is new. For this reason, the chemical process industries have come to prefer the more creative or open-ended technique, known as a hazard and operability study, or hazop. It is described later.

After we have identified the hazards, we have to decide how far to go in removing them or in protecting people and property. Some of the methods used for assessing hazards are listed at the bottom of Fig. 1. Sometimes, there is a cheap and obvious way of eliminating the hazard; sometimes our experience or a code of practice tells us what to do. At other

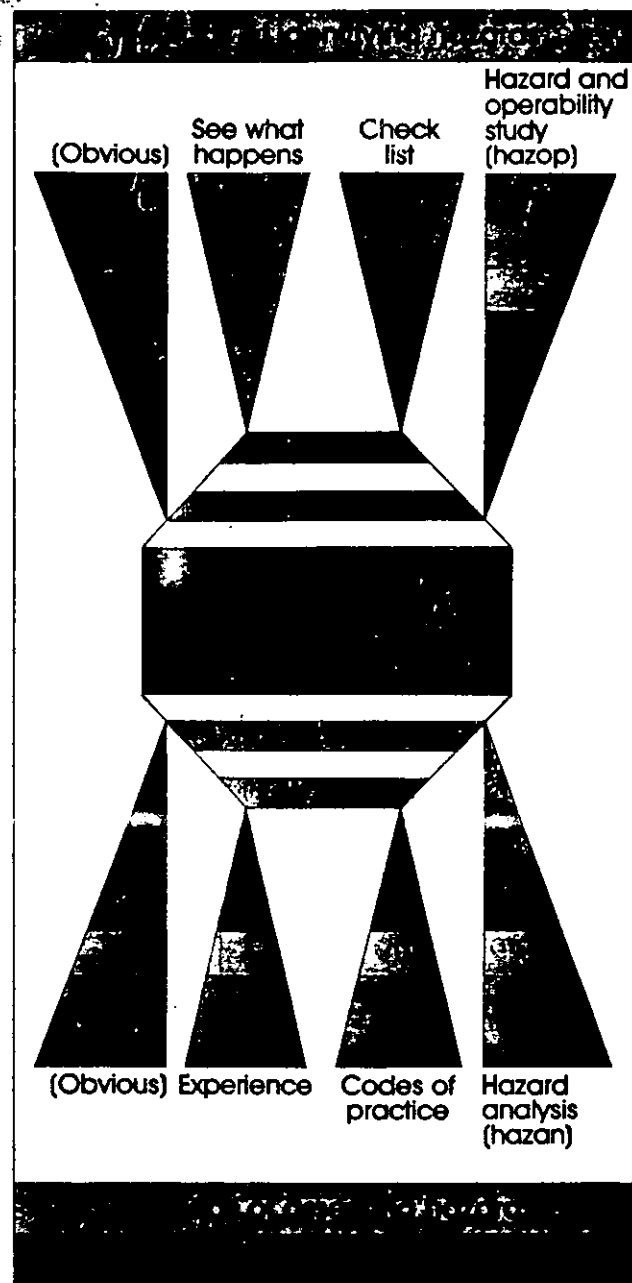


Figure 1 — Methods for identifying and assessing plant hazards

...it is less easy to decide. We can then try to work out the probability of an accident and the extent of the consequences, and compare them with a target or criterion. This method is known as hazard analysis, or hazan. Sometimes, a five minute estimation is sufficient. On other occasions, detailed studies take many weeks.

Such a study is called *hazard* analysis rather than *risk* analysis, because risk analysis is used to describe methods of estimating economic risks; and hazard *analysis* because, as we shall see, an essential step is breaking down the events leading to the hazard into their constituent parts.

Hazop can and should be applied to all new designs (unless we are making an exact copy of an existing plant that has been proved satisfactory), as we need to know all the hazards and problems that can prevent efficient operation. *Hazan* on the other hand, should be used selectively — there are not the need, the data or the resources to attempt to quantify every problem of every plant.

In the development of a design, the hazard and operability study comes first. We identify the hazards and the problems that prevent efficient operation and then decide what to do about them. However, if there is an obvious major hazard, we may start on the hazard analysis before the hazard and operability study is carried out.

It is hoped that Fig. 1 has made the difference between hazop and hazan clear. However, if someone asks you to carry out a hazop or hazan on a design, you should first make sure that the person who asks is clear on the difference.

Note that in a hazard and operability study the operability part is as important as the hazard part. In most studies, more operating problems are identified than hazards.

The techniques described below are sophisticated ones that enable companies to use their resources more effectively. These techniques assume that management is competent, and that the plant will be operated and maintained in the manner assumed by the design team and in accordance with good management and engineering practice. In particular, it is assumed that protective systems will be tested regularly and, when necessary, repaired promptly.

If these assumptions are not true, then hazop and/or hazan are a waste of time. It is no use identifying hazards or estimating their probability if no one wants to do anything about them; it is no use installing trips and alarms if they are going to be ignored or not maintained. The time spent on a hazop and hazan would be better spent on bringing management up to standard.

If you wish to introduce hazop or hazan into an organization in which they have not been used before, you should start small. Do not try to set up a large team capable of studying all new and existing designs. Instead, apply the methods to one or two cases. If your colleagues find that the methods are useful they will ask for more and the use of the techniques will grow. If, on the other hand, the methods do not suit your organization, little has been lost.

Despite all efforts, we shall fail to foresee every hazard and some will cause accidents. We can learn from accidents, not only from those that cause serious injury or damage but also from those that do not — e.g., leaks of flammable fluids that do not ignite. It is essential that these "near misses" are investigated and the lessons made known to those concerned, or the next time, injury or damage may result.

HAZARD AND OPERABILITY STUDY (HAZOP)

As already explained, a hazard and operability study is the recommended method for identifying hazards and problems that prevent efficient operation. In the following pages, the technique is described mainly as it would be applied to a continuous plant. Modifications of the technique, so that it can be applied to batch plants, are described in [1] and [2].

Hazop is now widely used on designs for new plants and plant extensions but, because of the effort involved, has been less widely used on existing plants.

The technique provides opportunities for people to let their

Table 1 — Deviations generated by the various guide words

Guide word	Deviations
None	No forward flow when there should be, i.e., no flow or reverse flow.
More of	More of any relevant physical property than there should be, e.g., higher flow (rate or total quantity), higher temperature, higher pressure, higher viscosity, etc.
Less of	Less of any relevant physical property than there should be, e.g., lower flow (rate or total quantity), lower temperature, lower pressure, etc.
Part of	Composition of system different from what it should be, e.g., change in ratio of components, component missing, etc.
More than	More components present in the system than there should be, e.g., extra phase present (vapour, solid), impurities (air, water, acids, corrosion products), etc.
Other than	What else can happen apart from normal operation, e.g., startup, shutdown, uprating, low running, alternative operation mode, failure of plant services, maintenance, catalyst change, etc.

imaginations go free and think of all possible ways in which hazards or operating problems might arise. But, to reduce the chance that something will be missed, it is done in a systematic way, each pipeline and each sort of hazard being considered in turn.

A pipeline for our purposes here is one that joins two main equipment items; for example, we might start with the line leading from the feed tank through the feed pump to the first feed heater.

A series of guide words are applied, in turn, to this line:

- None.
- More of.
- Less of.
- Part of
- More than.
- Other than.

None, for example, means no forward flow, or reverse flow when there should be forward flow.

We ask:

- Could there be no flow?
- If so, how could it happen?
- What are the consequences of no flow?
- Are the consequences hazardous or do they prevent efficient operation?
- If so, can we prevent no flow (or protect against the consequences) by changing the design or operating method?

many accidents have occurred because process materials flowed in the direction opposite to that expected, and because nobody foresaw that this could occur. For example: Ethylene oxide and ammonia were reacted to make ethanolamine. Some ammonia flowed from the reactor, in the wrong direction, along the ethylene oxide transfer line into the ethylene oxide tank — past several check valves and a positive-displacement pump (it got past the pump through the relief valve that discharged into the pump suction line). The ammonia reacted with 30 m³ of ethylene oxide in the tank, which ruptured violently. The released ethylene oxide vapor exploded, causing damage and destruction over a wide area [22].

A hazard and operability study would have discovered that backflow could occur, and the methods described in this article would have enabled a rough estimate of that probability to be made.

On another occasion, some paraffin passed from a reactor up a chlorine transfer line and reacted with liquid chlorine in a catchpot. Bits of the catchpot were found 30 m away.

On many occasions, process materials have entered service lines, either because the service pressure was lower than usual or the process pressure higher than usual. The contamination then spread via the service lines (steam, air, nitrogen, water) to other parts of the plant.

On one occasion, ethylene entered a steam main through a leaking heat exchanger. Another branch of the steam main supplied a space heater in the basement of the control room, and the condensate was discharged to an open drain inside the building. Ethylene accumulated in the basement and was ignited (probably by electrical equipment that was not protected), destroying the building. Again, a hazard and operability study would have disclosed the route taken by the ethylene.

The Bhopal disaster

At the time of writing, the full technical causes of the disaster at Bhopal, India, in which about 2,500 people were killed and perhaps 10 times that number injured, are not known. On Dec. 3, 1984, there was a leak of methyl isocyanate from a storage tank, and the vapor spread beyond the plant boundary to a shanty town that had grown up around the plant.

According to press reports, the tank became contaminated with water and other materials, a reaction occurred, and the contents rose in temperature. The refrigeration system on the tank was not working. The relief valve lifted, but the scrubbing system, which should have absorbed the vapor, and the flare system, which should have burned any vapor that got past the scrubbing system, were underdesigned or not in working order.

If the reports are correct, and the tank did indeed become contaminated, then a hazard and operability study would probably have shown up the ways in which contamination could have occurred.

A more important cause of the Bhopal disaster was the failure to keep safety equipment and instrumentation in working order. As stated in the text, hazop and hazan are a waste of time if basic safety requirements are ignored (as was apparently the case in the Bhopal plant).

• If so, does the size of the hazard or problem (that is, the severity of the consequence, times the probability of occurrence) justify the extra expense?

The same questions are then applied to reverse flow and we then move on to the next guide words, *more of*. Could there be "more flow" than that designed? If so, how could it arise? And so on. The same questions are asked about "more pressure" and "more temperature," and if they are important, about other parameters such as "more radioactivity" or "more viscosity." Table I summarizes the meanings of the guide words, while Fig. 2 summarizes the whole process.

When all the lines leading into a vessel have been studied, the term *other than* is applied to the vessel. It is not essential to apply the other guide words, as any problems should come to light when the inlet and exit lines are studied. However, to reduce the chance that something will be missed, the guide words should be applied to any operation carried out in the vessel. For example, if settling takes place, we ask whether it is possible to have no settling, reverse settling (i.e., mixing), more settling or less settling (and similarly for stirring, heating, cooling, and any other operations).

Startup, shutdown and other less frequent conditions, such as catalyst regeneration, should be considered, as well as normal operation.

Table II describes in detail the results of a hazop on the plant shown in Fig. 3. More details follow. The procedure will become clearer as each item in the table is gone through in turn. When groups are being trained to do hazop, in order to get the most out of the table, they should be shown Fig. 3 on a screen, or be given copies, and should be asked to carry out a hazop on it, the instructor acting as team leader. Their results can then be compared with those in Table II.

However, the material in Table II should not be considered as the correct answer. Those taking part in the discussion may feel that the authors of the table went too far, or did not go far enough, and they could be right.

Table II is based on a real study of an actual design. It is not a synthetic exercise. However, it is written up in more detail than would be necessary in a real-life situation.

In studying a batch plant, one must apply the guide words to the instructions (whether written or on a computer) as well as to the pipelines. For example, if an instruction states that 1 metric ton of *A* has to be charged to a reactor, the team should consider such deviations as:

- Don't charge *A*.
- Charge more *A*.
- Charge less *A*.
- Charge as well as *A*.
- Charge part of *A* (if *A* is a mixture).
- Charge other than *A*.
- Reverse charge *A* (that is, can flow occur from the reactor to the *A* container?).

For further details, see [1] and [2].

Batch-type operations carried out on a continuous plant — e.g., conditioning of equipment, or catalyst change — should be studied similarly, by listing the sequence of operations and applying the guide words to each step.

Who carries out a hazop?

A hazop is carried out by a team. For a new design, the usual team is:

Project engineer — Usually a mechanical engineer and, at

Plant engineer—Responsible for mechanical maintenance; knows many of the faults that occur.

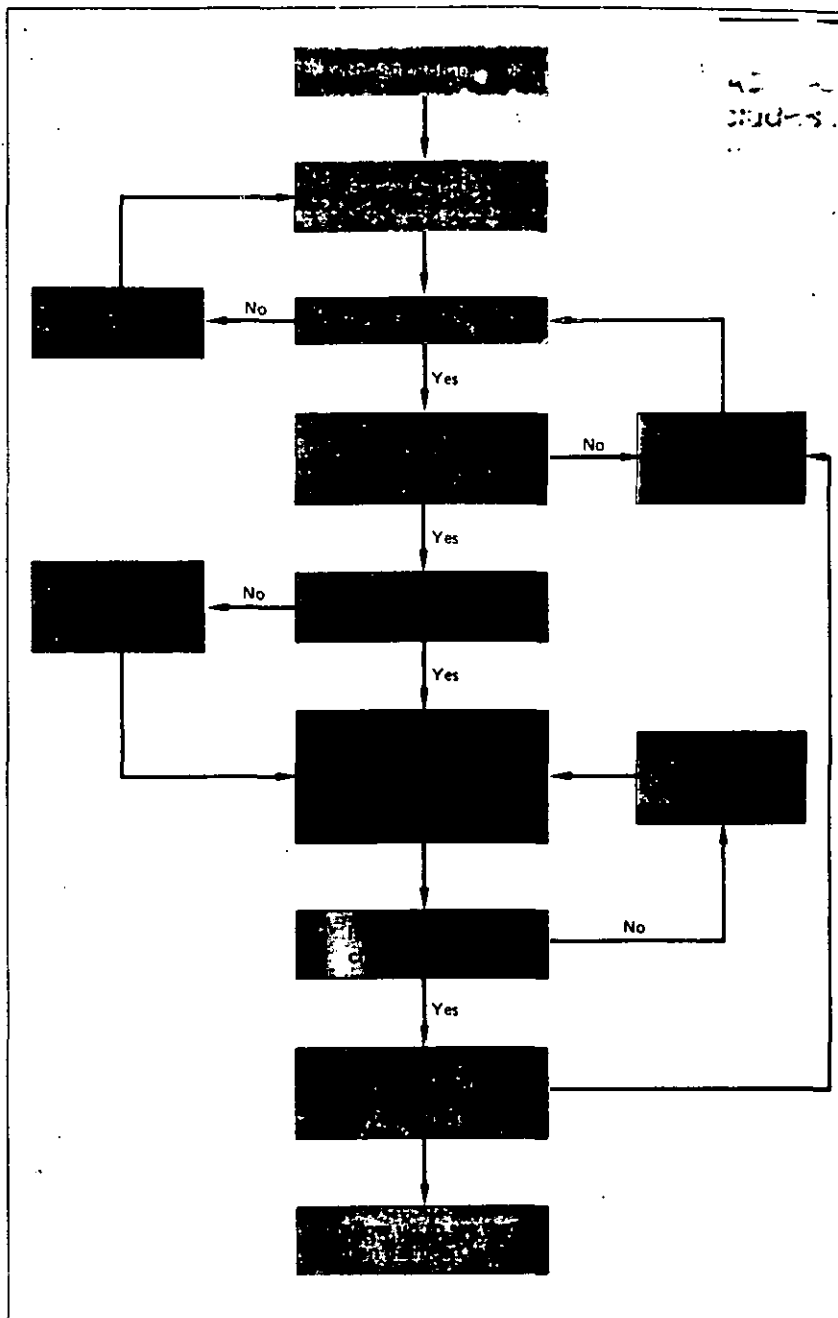
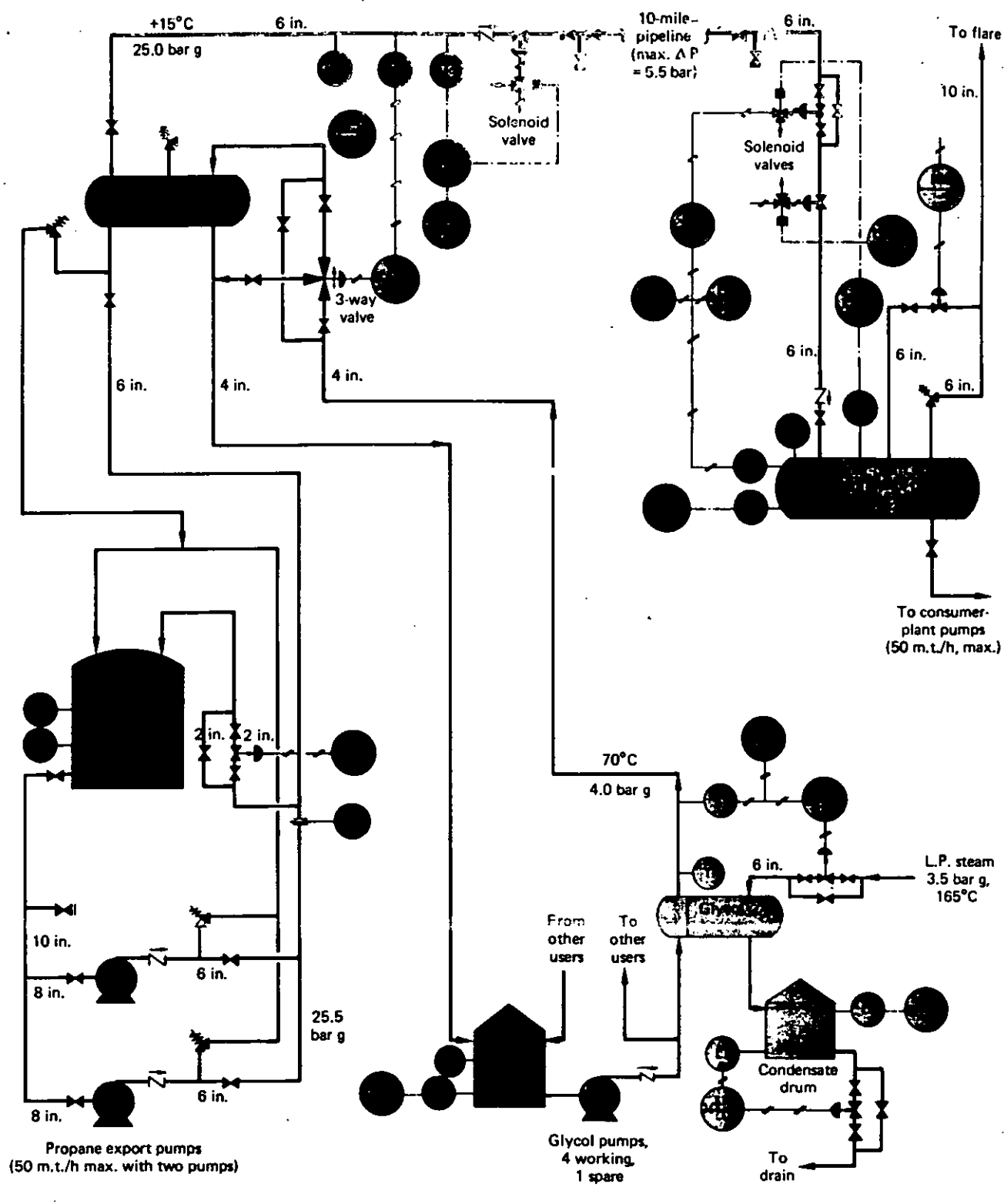


Figure 2 — Flowchart of the procedure for carrying out a hazop

Hazop teams, apart from the team leader, do not require much training. They can pick up the techniques as they go along. If anyone is present for the first time, the team leader



via a 10-mile cross-country pipeline to consumer plants (hazop example analyzed in Table II)

FEATURE REPORT

A

Table II — Hazop study of proposed liquid-propane transfer from main storage via pipeline to consumer's buffer tank

Deviation	Consequences
No flow	a. Online export pump(s) overheats, due to loss of suction, leading to pump seal leak and possible fire.
	b. Propane export to consumer plant lost, but limited storage is available in buffer tank for consumer units. Common low-flow kickback opens by normal action, and liquid upstream of export NRV (non-return valve) flashes back to tank. (But see also "reverse flow" for other consequences when pipeline backflow protection fails to operate on demand.)
	As for (b) above.
	c. As for (b), except major leakage in pump compound as well, leading to a flammable cloud that would almost certainly ignite.
	As for (a) and (b).
	d. As for (b), with full pump delivery-rate also passing back to tank via failed relief valve.
	e. Export-system-pressure rises to pump closed-head delivery pressure. Online pump(s) overheats, leading to pump-seal leak and possible fire.

carried out on the flowsheet before detailed design starts. This will take much less time than a full hazop.

Some points to watch during hazop

Don't get carried away — It is possible for a team to get carried away by enthusiasm and install expensive equipment to guard against unlikely hazards. The team leader can counter this by asking how often the hazards could occur and how serious the consequences might be. Sometimes, the leader may suggest a full hazard analysis, but more often can bring a problem into perspective by just quoting a few figures or asking a team member to do so. How often have similar pumps leaked in the past? How often do flanged

joints leak and how far do the leaks spread? How often do operators forget to close a valve when an alarm sounds? A later section on the use of quantitative methods describes a five-minute hazan carried out during a hazop meeting. (The most effective team leaders are trained in hazan as well as hazop.)

Hardware and software — The team consists mainly of engineers. They tend to like hardware solutions, but sometimes a hardware solution is impossible or too expensive, and we have to make a change in methods or improve the training of operators — that is, we change the software. We cannot spend our way out of every problem. Table II gives examples of software solutions as well as hardware ones.

Co
priat
Why
nor i
Mc
suits
hazo
How
ficat
ough
helpi
some
effect

Deviation		Consequences	
No flow (cont.)		f. As for (b), except that export system pressure and inventory is retained ready for restart. (Maximum kickback rate only about 10 m.t./h).	
		f1. Gross overpressuring of heater shell and glycol surge tank.	
		g. Large release of flashing liquid in public areas, with high chance of ignition and risk to life from radiation burns or deflagration. (The normal pipeline inventory is 160 m.t., and further quantities, due to continuing export and/or backflow from the buffer tank if the tank non-return valve fails to operate, could be discharged to atmosphere.)	
Reverse flow		h. As for (b), except that the liquid in the pipeline will also flash back to the main storage tank, leading to low pipeline temperatures and the chance of brittle fracture if the pressure falls below about 2 bar gage.	
More flow		i. Buffer tank overpressured to export-pump shutoff head, with high chance of rupture if no remedial action taken by operator on rising level, and if tank high-level trip and tank relief valve both fail to cope. Flare header at risk from brittle fracture when vessel overfills and relief valve operates.	
More pressure		j. Failure of pipeline, equipment or fittings due to surge pressure or hammer.	20. Facility operating

Contractors, in particular, should choose solutions appropriate to the sophistication and experience of their clients. Why install elaborate trips if the client has neither the skill nor the will to use them? Try less-sophisticated solutions.

Modifications—Many companies feel that hazop is unsuitable for small modifications. It is difficult to assemble a hazop team every time we wish to install an extra valve. However, many accidents have occurred because small modifications had unforeseen side-effects. They should be thoroughly probed before they are authorized; a guide sheet for helping to do so is shown in [3] and [4], which also describe some modifications that had unforeseen and hazardous side-effects.

Do we need a hazop?—Some may claim, "We don't need a hazop. We employ good people and rely on their knowledge and experience." Agreed, a hazop is no substitute for knowledge and experience. It is not a sausage machine that consumes line diagrams and produces lists of modifications. It merely harnesses the knowledge and experience of the team in a systematic and concerted way. Because designs are overwhelmingly complex, the team cannot apply its knowledge and experience without this crutch for their thinking. If the team lacks knowledge and experience, the hazop will produce nothing worthwhile.

"Good people" sometimes work in isolation. Hazop ensures that hazards and operating problems are considered

FEATURE REPORT

A

Table II (continued) — Hazop study of proposed liquid propane transfer from main storage via pipelining to consumer's buffer tank

Guideline	Deviation	Consequences
More pressure (cont.)		k. Failure of line, joints or valve glands.
		l. High metal temperature, leading to softening and loss of containment at normal operating pressure.
Less flow		m. Loss of vaporizer feed if no action taken, but several hours normally available.
		n. Propane escape to atmosphere and possible fire.
		o. High pressure, and possible subzero temperatures, developed in glycol circulation system.
Less pressure		p. Liquid propane from pipeline flashes into buffer tank as tank pressure falls, leading to low pipeline temperatures and to the chance of brittle fracture if the pressure falls below about 2 bar gage.
		q. Vaporizer feed system could be at risk from brittle fracture when buffer tank pressure is reduced.
		r. As for (p) for all causes listed. Consequence (q) also possible in the case of an emergency blowdown if vaporizer feed not isolated and shut down first.
Hand-control	Control room.	

systematically by people from different functions working together. Experience shows that startup, shutdown and other uncommon conditions are often overlooked by functional groups working in isolation.

"A hazop is done too late"—Some may charge that a hazop is carried out too late in design for basic changes to be made. True. By the time the line diagrams are available, it is usually too late to make major changes in design. Hazards are therefore *controlled* by adding on protective equipment rather than *removed* by changes in design. For this reason an increasing number of companies are carrying out preliminary hazops on flowsheets before detailed design starts.

At this stage, it is possible, for example, to replace a

flammable heat-transfer medium by a nonflammable one. At the line-diagram stage, all we can do is to reduce the risk by adding fire insulation, leak detectors, emergency isolation valves and so on. These preliminary hazops do not remove the need for a *full* hazop later.

An example of a hazop

Table II gives the results of a hazop on the plant shown in Fig. 3. Details of the operation are as follows:

Liquid propane is transferred via a 10-mile cross-country pipeline to a consumer plant, where it is vaporized for use as a feedstock for a group of cracking furnaces, and also as makeup for the site's fuel-gas system.

Deviation	Consequences
Less temperature	s. Chance of brittle fracture of pipeline if propane-export low-temperature trip fails to prevent liquid propane entering pipeline at below -15°C .
Change in composition (sulfur content)	t. Serious corrosion problems in combustion chambers of user units supplied from the site fuel-gas system, but other users not adversely affected by sulfur-rich propane gas.
Contamination (nitrogen)	u. Buffer-tank pressure higher than normal, and dissolved nitrogen present in feed to vaporizer.
Startup errors	v. Brittle fracture possible if liquid propane enters pipeline at below -15°C , or when pipeline pressure is below 2 bar gage.
Startup requirements	w. Pumps cannot be started.
Shutdown errors	x. Brittle fracture of pipeline possible if pressure falls below 2.0 bar gage, when significant amounts of liquid propane are present.
Routine testing	y. Automatic protection or alarms fail to operate on demand, thereby adversely affecting overall safety.

ble one. At the risk by isolation not remove

t shown in

ss-country l for use as and also as

Refrigerated liquid propane is pumped from an atmospheric-pressure storage tank (normal working conditions, -45°C and 20-in. water gage pressure) via the tubeside of a heater where the temperature of the hydrocarbon is raised to about 15°C before passing along a 10-mile overground cross-country pipeline to a buffer tank at the consumer plant that is operating on inlet level control and under autogenous pressure (6.5 bar gage; Fig. 3). This tank is required to provide limited storage for the consumer plant in the event of interruptions in the pipeline supply.

The two centrifugal export-pumps (closed-head delivery pressure, 30.5 bar gage) are supplied from different switchboards, and are designed for a combined maximum transfer

rate of 50 metric tons/h at the design delivery pressure of 25.5 bar gage. The pumps, which are each fitted with a single low-pressure trip for gearbox lube oil (not shown in Fig. 3), are on manual start, and are operated either one at a time or both together, depending on the consumer-plant requirements. Protection against possible damage at low or zero export rates is provided by a common low-flow kickback, which is designed to open on falling delivery flow, allowing propane to pass back to storage so as to maintain a minimum (8-m.t./h) delivery rate. The entire propane export system through to the buffer-tank inlet is designed to Class 300 specifications (49 bar gage) to allow for possible surge. The pump delivery relief-valves, each sized for full pump delivery

FEATURE REPORT

A

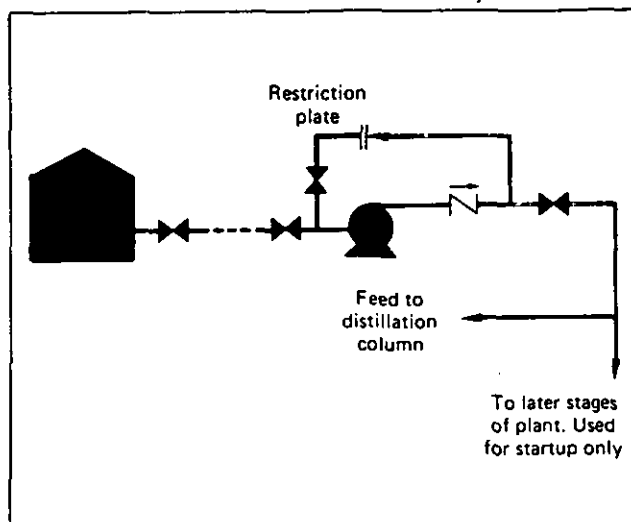


Figure 4 — 12 points came out of a hazop on this simple system

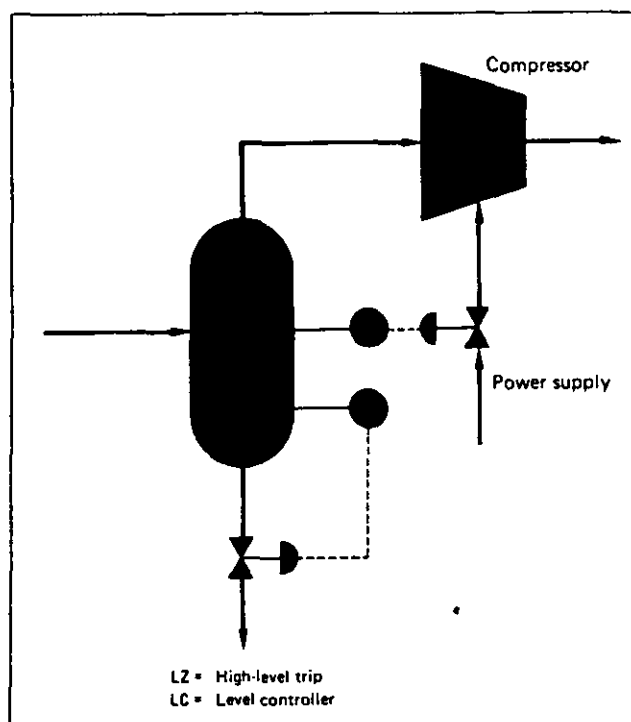


Figure 5 — Do we need to provide a second high-level trip?

rate, are thus set to lift at about 45 bar gage, as is the thermal relief valve on the tubeside of the heater.

Heat is supplied to the shellside of the propane heater by a continuous flow of warm glycol (70°C, 4.0 bar gage), which, being common to other users, is circulated by a group of pumps (four working, one spare, all from different boards) also on manual start. The glycol passing to the propane heater is heated by low-pressure steam (165°C, 3.5 bar gage) via a heater situated just upstream and operating on temperature control from the steam control valve.

The temperature of the propane export is controlled by means of a 3-way valve in the glycol circulation system, which permits bypassing of the propane heater when re-

quired. A manual bypass is also provided around the 3-way valve for online maintenance purposes. The steam condensate from the glycol heater is collected in a condensate drum that operates on level control via a 1-in. valve to drain. The glycol heating system is designed to cope adequately with the maximum propane export rate when the two pumps are operating under design conditions.

All equipment and pipework in the propane export area upstream of the 10-mile pipeline, and also in the consumer-plant area from pipeline battery limit through to and including the buffer tank plus associated fittings, is insulated and is designed for low-temperature duty down to -50°C. On the other hand, the 10-mile pipeline itself (Schedule 40; wall thickness 7.1 mm) is only to API 5L Grade B specifications (uninsulated) and thus would be at some degree of risk from brittle fracture at metal temperatures below about -15°C.

Accordingly, an independent low-temperature trip is provided on the propane export line just upstream of the change in piping specifications, to safeguard against this eventuality (proposed trip setting -5°C). Low-temperature alarms are also used on the circulating-glycol and propane-export temperature control loops to give pre-warning of low-temperature condition. In addition, there is an independent high-level alarm on the steam condensate drum to inform the operator of potential loss of steam inflow to the glycol heater due to overflowing of the drum, and an independent low-level alarm on the glycol surge tank to give pre-warning of loss of glycol pump suction due to loss of tank inventory.

The alarm instrumentation at the consumer-plant end of the pipeline comprises high- and low-level pre-alarms on the buffer-tank level-control loop, together with an independent, extra-low-level alarm and a low-pressure alarm on the tank. As an additional precaution against tank overflowing, there is also an independent, extra-high-level trip (with trip alarm) designed to close the inlet level-control valve, as well as an adjacent trip valve. The tank is fitted with conventional relief protection (set to lift at 13 bar gage), and a remotely-operable blowdown valve (HIC) to the flare header is provided for use at startup, shutdown, or in an emergency.

The normal pipeline startup procedure is to pressure up the equipment and lines, from export pump delivery through to and including the consumer buffer tank, to about 6.5 bar gage with nitrogen, using the site 10-bar-gage supply, and then to establish normal conditions on the glycol heating system before propane export is begun. This is to avoid low temperatures due to flashing of the hydrocarbon (if system is not first pressured up) and/or inadequate heating (if full heating is not available at the outset). Liquid propane is then pumped through the pipeline until normal level is established in the buffer tank, the displaced nitrogen being purged to flare as necessary via the tank blowdown valve.

Deburdening of the cross-country pipeline may be necessary from time to time — e.g., for repairs in the event of minor leaks. Because of the presence of the nonreturn valves, this will normally involve displacing the pipeline inventory through to the consumer buffer tank, using nitrogen at a liquid displacement rate matching the vapor requirements on the consumer units (i.e., with the buffer tank operating throughout on automatic level-control), care being necessary during this operation to avoid low pipeline pressures and temperatures, and hence possible brittle fracture.

The shutdown and deburdening procedure proposed is

first to isolate the pipeline at the battery limit and stop the export pump(s), thereby allowing liquid propane upstream of the export nonreturn valve to flash off back into the main storage tank. The bleed valve at the battery limit is then connected up to the nitrogen supply (10 bar gage) and, with the propane vaporizer still online, deburdening of the pipeline is begun via the buffer tank operating on inlet level control.

As soon as nitrogen breaks through to the buffer tank, the pipeline is isolated at the tank inlet, the nitrogen supply turned off, and the vaporizer kept on line until most of the feed has been worked off. The vaporizer unit is then isolated from the tank and shut down. Residual liquid in the tank is blown down via the tank HIC to the flare, before gradual depressuring of the pipeline via the same route.

A very low rate of depressuring is important in the final stages to avoid localized chilling of the pipeline due to pockets of residual liquid. The pipeline is finally swept out with nitrogen via the buffer tank HIC route to the flare. If purging of the export equipment is also required, this is achieved by using nitrogen from the upstream branch at the export pump suction, with the kickback flow control valve set isolated, for the final sweepout.

When only the buffer tank has to be deburdened, the procedure is similar, except that the pipeline is isolated at both ends, and nitrogen is connected to the bleed branch at the consumer end for the final purge.

The outcome of the studies recommended in items 11A and 26 of Table II was that emergency valves were installed on the propane and glycol lines inlet and exit of the heater, operated by a high-pressure switch on the shellside and by a gas-in-glycol detector. This action was necessary as it would be difficult to protect the low-pressure surge tank from gross overpressure by using a relief valve. The hazard analysis requested in Items 14, 28, 30, 32, 35 and 37 of the table is written up in detail in Ref. 5.

"Why hazop this plant?"

"Why should we hazop this plant? It is only a simple project (or it is similar to the last one)."

So many of the things that go wrong occur on small, simple or repeat units where people feel that the full treatment is unnecessary. "It is only a storage project and we have done many of these before!" "It is only a pipeline and a couple of pumps."

If designers talk like this, suggest they try a hazop and see what comes out of it. After the first meeting or two, they usually want to continue.

Fig. 4 shows part of a line diagram on which the design team were persuaded, somewhat reluctantly, to carry out a hazop. Twelve points that had been overlooked came out of the study. Here are four:

- If the pump stops, the backflow will occur through the kickback line. The nonreturn valve should be downstream on this line
- If the pump stops, backflow may occur through the startup line. Should there be a nonreturn valve in this line?
- The restriction plate in the kickback line might be replaced by a flow controller to save power.
- Since no provision has been made for slip rings or spectacle (figure-8) plates, the pump cannot be isolated by slip plates (blinds) for maintenance.

The design team then readily agreed to study the rest of the plant.

Use of quantitative methods during hazop

The following example shows how a quick calculation can resolve a difference of opinion between members of a hazop team. It acts as a link to the next section of this article, in which numerical methods are considered in more detail.

On a design, a compressor suction catchpot was fitted with a level controller, and a high-level trip that shut down the machine (Fig. 5). The commissioning manager asked for a second, independent, trip, as failure of the trip could result in damage to the machine, which would be expensive to repair. The project engineer, responsible for controlling the cost, was opposed. This, he said, would be "gold-plating." A simple calculation helps to resolve the conflict.

The trip will have a fail-danger rate of once in two years. With monthly testing, the fractional deadtime will be 0.02.

The demand rate results from the failure of the level controller. Experience shows that a typical figure is once every two years, or 0.5/yr.

A hazard will therefore occur once in 100 years; or, more precisely, there is a 1 in 100 chance that it will occur in any one year, or a 1 in 10 chance that it will occur during the 10-year life of the plant.

Everyone agreed that this probability was too high to risk. They also saw that there was more than one way of reducing the hazard rate. They could improve the control system and reduce the demand rate, or they could improve the trip system and reduce the fractional deadtime. It may not be necessary to duplicate the entire trip system — duplicating only the trip initiator may suffice.

HAZARD ANALYSIS (HAZAN)

Hazard analysis is the term used to describe the application of numerical methods to safety problems.

It is not an esoteric technique that can be practiced only by an elite band of the initiated. It can be employed by any competent technologist, provided he or she discusses the first attempts with someone more experienced.

Why apply numerical methods?

The horizontal axis of Fig. 6 shows expenditure on safety over and above that necessary for a workable plant, and the vertical axis shows the money we get back in return.

In the left-hand area, safety is good business — by spending money on safety, apart from preventing injuries, our plants blow up or burn down less often and we make more profit.

In the next area, safety is poor business — we get some money back for our safety expenditure but not as much as we would get by investing our money in other ways.

If we go on spending money on safety, we move into the third area where safety is bad business but good humanity. Money is spent so that people do not get hurt, and we do not expect to get any material profit back in return.

Finally, in the fourth area, we are spending so much on safety that we go out of business. Our products become so expensive that nobody will buy them; our company is bank-

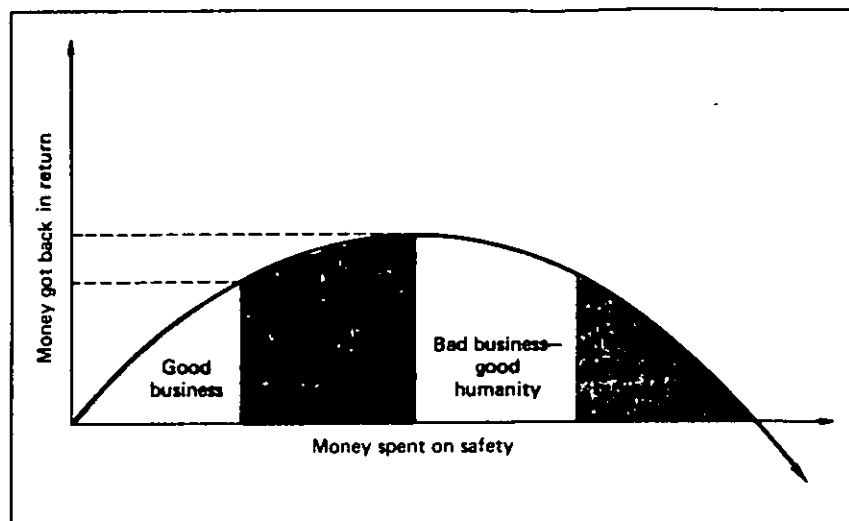


Figure 6 — Effects of increasing expenditures on safety

rupt and we are out of a job. The public is deprived of the benefits it could get from our products.

We have to decide where to draw the line between the last two areas. Usually this is a qualitative judgment, but it is often possible to make it quantitative. The methods for doing so are known as hazard analysis or hazan.

The stages of hazard analysis

Every hazard analysis consists of three steps:

1. Estimating how often the incident will occur.
2. Estimating the consequences to:
 - Employees.
 - Members of the public.
 - Plants and profits.

In both (1) and (2), estimates should be based, whenever possible, on past experience. However, sometimes there is no past experience, either because the design is new or the incident has never happened, and in these cases we must use synthetic methods.

3. Comparing the results of (1) and (2) with a target or criterion, in order to decide whether or not action to reduce the probability of occurrence or minimize the consequence is desirable, or whether the hazard can be ignored (at least for the time being).

Many writers on plant safety are reluctant to discuss Step 3, but it is of little use knowing that a plant will blow up once in 1,000 years, with a 50% chance that someone will be killed, unless we can employ this information to help decide whether we should reduce the probability of the incident occurring (or protect people from the consequences), or whether the risk is so small, compared with all the other risks around us, that we should ignore it and devote our attention to bigger risks.

In the following pages, we first discuss Step 3, then Step 1. Discussion of Step 2 is not attempted; the methods used differ for each type of hazard: fires, explosions and releases of toxic gas, and reference should be made to specialist textbooks or to Ref. 4.

In brief, the stages in hazard analysis are:

- How often? • How big? • So what?

If you can remember these six words, you will know what

to do (though not how to do it) if you are ever asked to carry out a hazard analysis. Also, you will know what to look for in hazard analyses carried out by others.

Some of the targets or criteria

When injury is unlikely, we can compare the annual cost of preventing an accident with the average annual cost of the accident:

Suppose an accident will cause \$1 million worth of damage (but with no injuries) and is estimated to occur once in 1,000 years — an average cost of \$1,000/yr. Then it is worth spending up to \$1,000/yr to prevent it, but not more. Capital costs can be converted to maintenance, depreciation and interest. Actually, future costs should be discounted, but the data are often not accurate

enough to make this worthwhile.

This method *could* be used for all accidents if we could put a value on injuries and life, but since there is no generally agreed figure, we set a target instead.

For example, in fixing the height of handrails around a place of work, the law does not ask us to compare the cost of fitting them with the value of the lives of the people who would otherwise fall off. It fixes a height for the handrails (36 to 45 in. in the U.K.). A sort of intuitive hazan has been done to show that with handrails of this height the chance of falling over them is so small (though not zero) that we are justified in ignoring it.

Similarly, we fix a "height" or level for the risk to life. In setting this level, we should remember that we are all at risk all the time, whatever we do, even staying home. We accept the risks when we consider that, by doing so, something worthwhile is achieved.

We go rock-climbing or sailing or we smoke cigarettes because we consider the pleasure worth the risk. We take jobs as airline pilots or soldiers or we become missionaries among cannibals because we consider that the pay, or the interest of the job, or the benefit it brings to others, makes the risk worthwhile.

At work, there is likely to be a slight risk, whatever we do to remove known perils. By accepting this risk, we earn our living and we make goods that enable ourselves and others to lead a fuller life.

A widely used target for the risk to life of employees (discussed in the next section) is the Fatal Accident Rate (FAR).

Risks to the public will be discussed in a later section.

However, it is not always necessary to estimate the risk to life. When we are making a change, it is often sufficient to say that the new design must be as safe as, preferably safer than, one that has been generally accepted without complaint. For example:

- If trips are used instead of relief valves, they should have a probability of failure 1/10 as high [6,7].
- If equipment that might cause ignition is introduced into a Zone 2 (Div. 2) area, it should be no more likely to spark than the electrical equipment already there.

• A new form of transport should be no more hazardous, preferably less hazardous, than the old form.

Risks to employees — the FAR

The fatal accident rate (FAR) is the number of fatal accidents per 1,000 workers in a working lifetime (10^5 h).

For weekly-paid employees in the chemical industry, the FAR = 4 (the same as the average for all activities covered by the U.K. Factories Act).*

This is made up from:

• Ordinary industrial risks (e.g., falling downstairs or getting run over). FAR = 2

• Chemical risks (e.g., fire, toxic release or spillage of corrosive chemicals). FAR = 2

If we are sure that we have identified all the chemical risks attached to a particular job, we say that the person doing the job should not be exposed to an FAR greater than 2 for these chemical risks. We will eliminate or reduce, as a matter of priority, any such risks on new or existing plants.

It would be wrong to spend our resources on reducing the risk to people who are already exposed to below-average risk. Instead, we should give priority to above-average ones.

Often, we are not sure that we have identified all the chemical risks, and so we say that any single one, considered in isolation, should not expose an employee to an FAR greater than 0.4. We will eliminate or reduce, as a matter of priority, any hazard on a new or existing plant that exceeds this figure. We are thus assuming that there are about five significant chemical risks in a typical plant.

Experience has shown that the costs of doing this, though often substantial, are not unbearable. They may involve the company in an expenditure that some of its competitors do not incur. Some of the extra expenditure can be recouped in lower insurance premiums. Some can be recouped by the greater plant reliability that safety measures often produce. The rest is a self-imposed "tax" that has to be balanced by greater efficiency.

Note that the FAR is estimated for the person or group at the greatest risk, not for all employees in a plant as a whole.

Note also that we are not suggesting that our target FAR's are "acceptable." No risk to life is ever really acceptable. We should never knowingly fail to act when life is at risk. But we cannot do everything at once; we have to set priorities. We have to decide which are the biggest risks and deal with them first.

To convert FAR to hazard rate

The hazard rate is the rate at which dangerous incidents occur. Suppose the person at greatest risk is killed every time the dangerous incident occurs (this is an example, not a typical situation). Then it must not occur more often than:

0.4 occasion in 10^5 working hours, or once in 2.5×10^5 working hours (= 30,000 years)

or 3×10^{-5} occasion/year — i.e., the probability of occurrence should not exceed 3×10^{-5} /year (for shift jobs).

For a job manned only during day hours, the corresponding figures are once in 120,000 years, or 8×10^{-6} occasion/year.

*If you spend your working lifetime in a typical factory of 1,000 workers, then during your time there 4 of your fellows will be killed in industrial accidents — but about 20 will be killed in other accidents (mostly on the roads and in the home) and about 370 will die from disease, including about 40 from the results of smoking (if present rates continue).

If the worker at greatest risk is killed every 10th time the incident occurs, then the target hazard rate is:

once in 3,000 years, or 3×10^{-4} occasion/yr, and so on.

Multiple casualties

What is the target hazard rate if more than one person is killed?

Consider two cases:

Case 1. One person killed every year for 100 years.

Case 2. 100 people killed once in 100 years.

Should the prevention of Case 1 have higher priority than the prevention of Case 2, or vice versa?

The arguments in favor of giving priority to the prevention of Case 2 are:

The press, public and legislatures make more fuss about Case 2, while they usually ignore Case 1. The public perceives Case 2 as worse; as servants of the public, we must therefore give priority to the prevention of Case 2.

Case 2 disrupts the organization and the local community, and the wounds take longer to heal. It may cause production to be halted for a long time (perhaps even forever) and new legislative requirements may be introduced.

Various writers have proposed that the acceptable hazard rate for Case 2 should be the acceptable rate for Case 1, divided by either $\log N$, or N or N^2 , where N is the number of people killed per incident. However, these formulas are quite arbitrary, and if we divide the hazard rate by N^2 , or even N , we may get such low rates that they are impossible to achieve.

Gibson [8] has suggested that we can allow for the wider effects by estimating the financial costs of disruption of production, etc., and comparing them with the costs of prevention. This may be a more effective and defensible method than introducing arbitrary factors.

It is true that as servants of the public we should do what they want. But a good servant does not obey unthinkingly; he points out the consequences of his instructions. If we think the public's perception of risks is wrong, we should say so, and say why we think so. Perhaps the public think that preventing events like Case 2 will reduce the number of people killed accidentally; it actually would have very little effect on the total number killed.

The argument in favor of giving priority to the prevention of Case 1 is that Case 2 will probably never happen (if the plant lasts 10 years the odds are 10 to 1 against) but that Case 1 almost certainly will happen — one person will almost certainly be killed every year — so why not give priority to preventing the deaths of those who will almost certainly be killed, rather than to preventing events that will probably never happen? This argument becomes stronger if we consider Case 3:

Case 3. 1,000 people are killed once in 1,000 years.

In this case, it is 100 to 1 that nobody will be killed during the life of the plant.

The simplest and fairest view seems to be to give priority to the prevention of both Case 1 and 2 — the victims are just as dead in Case 1 as in Case 2.

If we give priority to the prevention of Case 2 we are taking resources away from the prevention of Case 1 and, in effect, saying to the people who will be killed one at a time that we consider their deaths as less important than others.

FEATURE REPORT

We should treat all persons the same. There may, however, be an *economic* argument for preventing Case 2, as stated by Gibson, even though the risk is so small that we would not normally spend resources on reducing it further.

Consider now two more cases:

Case 4. A plant blows up once in 1,000 years, killing the single operator.

Case 5. A similar plant, less automated, blows up once in 1,000 years but kills all 10 operators.

The FAR is the same in both cases, the risk to all operators is the same, but some way of drawing attention to the high *exposure* involved in Case 4 is desirable. Lees [20] suggests that the number killed, the *accident fatality number*, should be quoted, as well as the FAR.

Risks to the public

We accept voluntarily a number of risks such as driving, flying and smoking that expose us to a risk of death of 10^{-5} or more, sometimes a lot more, per person per year (FAR = 0.1). We also accept, with little or no complaint, a number of involuntary risks (for example, from lightning or falling aircraft) that expose us to risk of death of about 10^{-7} per person per year (FAR = 0.001).

We thus have a possible basis for considering risks to the public at large from an industrial activity. If the average risk to those exposed is more than 10^{-7} per person per year, we should eliminate or reduce the risk as a matter of priority. If it is already less, it would not be right to spend scarce resources on reducing the risk further. It would be like spending money on protecting people from lightning. There are more important hazards to be dealt with first.

As well as considering the average risk, we should consider the person at greatest risk. A man aged 30 years has a probability of death from all causes of 1 in 1,000 per year. (The figure for a younger man is not much less.) An increase of 1% from industrial risks is hardly likely to cause him much more concern and an increase of 0.1% should certainly not do so. This gives a range of 10^{-5} to 10^{-6} per year.

Many other criteria have been proposed for risks to the public. They are reviewed in [21]. The criteria vary but it is generally agreed that the public should be exposed to much lower risks than employees are. People choose to work for a particular company or industry but members of the public have risks imposed on them against their will. However, members of the public are farther away from the source of the hazard, so in practice the risk to employees may be more important. For example, the pressure developed by an explosion decreases with distance; the risk to the public is so much less than the risk to employees that reducing the latter is usually the more important task.

Why consider only fatal accidents?

As pointed out many years ago, there is a relationship between fatal, lost-time, minor and no-injury accidents (in which only material damage is caused). If we halve fatal accidents from a particular cause, we halve lost-time accidents, minor accidents, and no-injury accidents from that cause. If we halve the number of deaths from explosions, for example, in a particular plant, we probably also halve the number of lost-time accidents and minor accidents caused by explosions, as well as the material damage they cause.

Note that halving the *total* number of fatal accidents will

not necessarily halve the *total* number of lost-time (or minor) accidents, because the ratio of lost-time to fatal accidents differs for various types of accidents. For example, it is about 1:250 for transport accidents, but about 1:20,000 for accidents involving the use of tools.

ASSESSING HOW OFTEN AN INCIDENT WILL OCCUR

The methods to be described in this section are used when we have no past experience to go by.

Some definitions

First we will define some terms we will need:

Hazard rate — This is the rate (occasions/year) at which hazards occur — for example, the rate at which the pressure in a vessel exceeds the design pressure, or the rate at which the level in a vessel gets too high and the vessel overflows.

Protective system — This is a device installed to prevent the hazard from occurring — for example, a relief valve or a high-level trip.

Test interval (T) — Protective systems should be tested from time to time to see if they are inactive or "dead". The time between successive tests is the test interval *T*.

Demand rate (D) — This is the rate (occasions/year) at which a protective system is called on to act — for example, the rate at which the pressure rises to the relief-valve set pressure, or the rate at which a level rises to the setpoint of the high-level trip. The word "demand" is used in the French sense (demander = to ask).

Failure rate (f) — The rate (occasions/year) at which a protective system develops faults. The faults of most interest to us are fail-danger ones that prevent the protective system from operating. But fail-safe faults also can occur; these result in the protective system operating when it should not. For example, a relief valve lifts below its set pressure, or a high-level trip operates when the level is normal.

Most failures are random, and this is assumed as we develop this section. However, failures can be high when equipment is new and when it is worn out (that is, just after birth and during old age).

Fractional deadtime (Fdt) — This is the fraction of the time that a protective system is inactive — i.e., the probability that it will fail to operate when required.

If the protective system never failed to operate when required, then the hazard rate would be 0.

If there were no protective system, then the hazard rate would be equal to the demand rate.

Usually, the protective system is inoperative or dead for a (small) fraction of the time.

A hazard results when a demand occurs during a dead period, hence:

Hazard rate = Demand rate $\times$ Fractional deadtime (but see the later section headed "More-accurate formula").

Some of the figures used in the following examples are typical, while the others are merely assumed.

Example 1: Relief valves

Tests on relief valves show that fail-danger faults that will prevent them lifting within 20% of the set pressure occur at a rate (*f*) of 0.01/yr, or once in 100 yr (a typical figure).

Let test interval *T* = 1 year (a typical figure).

Failure occurs, on average, halfway between tests. There-

fore, the relief valve is dead for 6 months ($\frac{1}{2}T$) every 100 ($1/f$) years, or for $1/200$ or 0.005 of the time ($\frac{1}{2}fT$).

This is the fractional deadtime.

Suppose the demand rate D is 1/year (an example).

A hazard results when a demand occurs during the time that the relief valve is dead. The valve is dead for $1/200$ of the time, and there is one demand per year, so there is a hazard once in 200 years.

Expressed more precisely:

$$\begin{aligned}\text{Hazard rate} &= \text{Demand rate} \times \text{Fractional deadtime} \\ &= D \times \frac{1}{2}fT \\ &= 1 \times 0.005 \\ &= 0.005/\text{year}\end{aligned}$$

or once in 200 years. (The "more-accurate formula" given later yields once in 250 years.)

We could not determine this figure by counting the number of occasions on which vessels have been overpressured, because this occurs so rarely; but we have been able to estimate it from the results of tests on relief valves.

Note that in this example a hazard has been defined as taking a vessel more than 20% above its design pressure. Not all such "hazards" will result in vessel rupture, or even a leak.

Example 2: Simple trips

Assume that: Fail-danger faults develop at a rate f of once every two years, or $0.5/\text{year}$ (a typical figure), much more frequently than with relief valves.

The test interval, T , is 1 week.

The demand rate, D , is 1/year.

Calculate the fractional deadtime and the hazard rate.

Answer:

The trip is dead for $3\frac{1}{2}$ days every 2 years.

Therefore, fractional deadtime $= 3.5/(2 \times 365) = 0.005$

Hazard rate $= 1 \times 0.005 = 0.005/\text{yr}$ or 1 in 200 years

With monthly testing, fractional deadtime $= 0.021$

Hazard rate $= 1$ in 48 years

With annual testing, fractional deadtime $= 0.25$

Hazard rate $= 1$ in 4 years

(The more-accurate formula gives 1 in 5 years)

If we take into account the time the trip is dead while it is being tested, then weekly testing may not give the lowest hazard rate, and monthly testing may be better.

If the trip is tested yearly, then the hazard rate is only reduced from $1/\text{yr}$ with no trip, to once in 5 years. If the trip is so unimportant that annual testing is sufficient, then the trip is probably not necessary.

Because trips fail more often than do relief valves, they have to be tested more often.

Example 3: Frequent demands on a trip

Let demand rate $D = 100/\text{yr}$ (an example).

Test interval $T = 0.1$ yr (a typical figure).

Failure rate $f = 0.5/\text{yr}$ (a typical figure).

Calculate the fractional deadtime and the hazard rate.

Answer:

Using the formula:

$$\begin{aligned}\text{Hazard rate} &= D \times \frac{1}{2}fT \\ \text{Hazard rate} &= 100 \times \frac{1}{2} \times 0.5 \times 0.1 \\ &= 2.5/\text{yr}\end{aligned}$$

In fact, the hazard will be almost the same as the failure rate ($0.5/\text{year}$), because:

• There will always be a demand in the dead period.

• The fault will then be disclosed and repaired.

($2.5/\text{yr}$ would be the right answer if, when a hazard occurred, we did not repair the trip but left it in a failed state until the next test was due.)

Testing in this situation is a waste of time, as almost all failures are followed by a demand before the test is due.

If you find the last example hard to follow, consider the brakes on a car:

Frequent demands on a trip — Auto brakes

Let failure rate $f = 0.1/\text{yr}$ (a typical figure).

Test interval $T = 1$ yr (as required by law in some places).

Demand rate $D = 10^4/\text{yr}$ (a guess).

Using the formula:

$$\begin{aligned}\text{Hazard rate} &= D \times \frac{1}{2}fT \\ &= 10^4 \times \frac{1}{2} \times 0.1 \times 1 \\ &= 500/\text{year!}\end{aligned}$$

The true answer is $0.1/\text{yr}$.

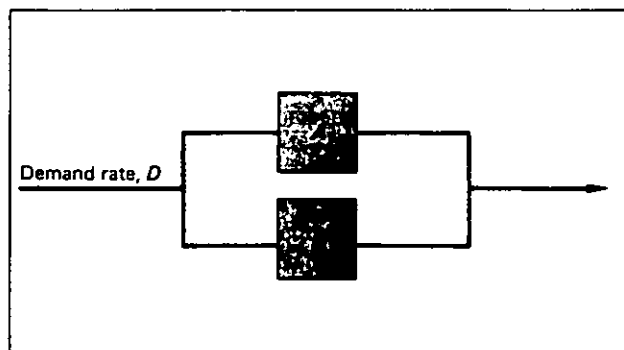


Figure 7 — Two protective systems in parallel

Why? The simple intuitive formula that we had derived earlier:

Hazard rate = Demand rate $\times$ Fractional deadtime must be incorrect.

More-accurate formula

$$\text{Hazard rate} = f(1 - e^{-DT/2})$$

Where: f = failure rate; T = test interval; and D = demand rate.

If $DT/2$ is small, this becomes: Hazard rate $= \frac{1}{2}fDT$

If $DT/2$ is large, this becomes: Hazard rate $= f$

As a rough rule of thumb, if $\frac{1}{2}fDT$ less than f , hazard rate $= \frac{1}{2}fDT$. But, if $\frac{1}{2}fDT$ is greater than f , hazard rate $= f$.

The exponential formula above is correct only when fT is small, and applies only to a single protective system.

Two protective systems in parallel

Examples are two 100% relief valves in parallel or two high-level trips (Fig. 7).

Let F_A and F_B be the fractional deadtimes of systems A and B .

The setpoints of the two systems are, by accident or design, never exactly the same.

Assume A responds first — that is, if A and B are two relief valves, A is set at a slightly lower pressure; if A and B are two high-level trips, A is set at a lower level.

The demand rate on $A = D$.

FEATURE REPORT

The frequency of demands to which *A* does not respond is $F_A D$.

This is the demand rate on *B*.

Therefore, the fractional deadtime of the combined system should be $F_A F_B$, and the hazard rate should be $D F_A F_B$.

Actually, the fractional deadtime is $4/3 F_A F_B$ and the hazard rate is $4/3 D F_A F_B$.

(Because the demands on *B* tend to occur toward the end of a proof-test interval when there is a more than average likelihood that *B* will have failed. If *A* and *B* are tested at different times, the hazard rate can be shown to be $0.83 D F_A F_B$.)

Like the example in the section "Frequent demands on a trip," this shows the perils of intuitive mathematics.

Two protective systems in series

An example is a relief valve and a bursting disc in series (Fig. 8). (Failure of a bursting disc in this context means failure to burst when the required bursting pressure is reached.)

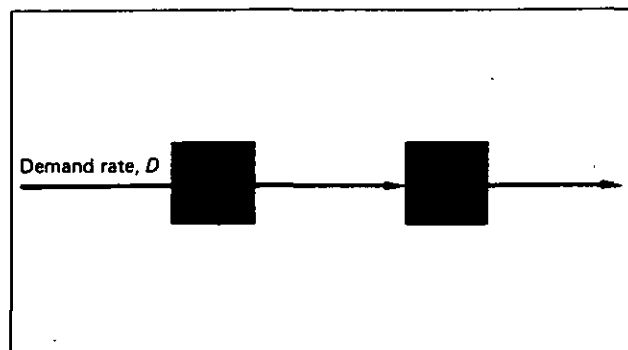


Figure 8 — Two protective systems in series

If *A* or *B* fails, the system is dead.

Fdt of the whole system = $F_A + F_B - F_A F_B$
or, $Fdt = F_A + F_B$ (if F_A, F_B are small).

Fault trees

Some examples of fault trees are shown in Figs. 9 and 10. They are widely used in hazard analysis to set down, in a logical way, the events leading to a hazardous occurrence. They allow us to see the various combinations of events that are needed and the various ways in which the chain of events can be broken. They allow us to calculate the probability of the hazardous event from other probabilities that are known.

In drawing a fault tree, we start on the left with the hazardous event — for example, that common industrial hazard — a free meal* (the logic is the same if you regard a free meal as desirable rather than a hazard), or the overpressuring of a vessel. Some people start at the top instead of the left, so the hazardous event is often called the top event. We then work from left to right (or top to bottom), drawing in the various events that lead up to the top event. We then work back, inserting numbers and then estimating the frequency of the top event.

The points at which two branches of a fault tree join are known as gates; they can be "AND" or "OR" gates.

Fig. 9 shows two example of "AND" gates. Both a meeting with lunch AND an invitation are required for a free

*Not a problem in universities

meal. Note that a *frequency* is multiplied by a *probability*. A common beginner's mistake is to multiply two frequencies. (Two probabilities *can* be multiplied together.)

In Fig. 10, the logic trees have been extended and OR gates are shown. We need visitors OR a training course but not both to get a free meal. Note that at an "OR" gate the two rates are added together. In practice, we stop drawing when we have data for the frequency or probability of the events on the right of the tree.

Suppose we are asked to revise Fig. 10 (a). We examine records for 10 years, carry out a regression analysis, allow for the effect of the changing economic situation, and conclude that the visitor rate is more likely to be 12/yr or 20/yr instead of 15/yr. The effect on the frequency of the top event is negligible. Similarly, detailed study may show that instead of 5 training courses per year we should expect 3, or perhaps 8. Again, the effect on the final answer is small. The number of free meals is between 1.5 and 2.8/yr and is unlikely to be near these limits.

A more serious source of error is that we have not considered that some visitors may stay to dinner. If half of them do and the probability of an invitation is the same, the free meal rate rises to 2.75/yr.

More serious still, suppose a new boss decides that all the staff should meet together over a free lunch once per week for an informal discussion. The free meal rate rises to 48/yr (assuming 4 weeks holiday) + 2/yr from other causes = 50/yr. Our original result is out by a factor of 25!

The simple example shows that most errors in hazard analysis are not due to errors in the data but to errors in drawing the fault tree — that is, to a failure to foresee all the ways in which the hazard could arise. Time is usually better spent looking for *all* the sources of hazard than in quantifying with ever greater precision those we have already found.

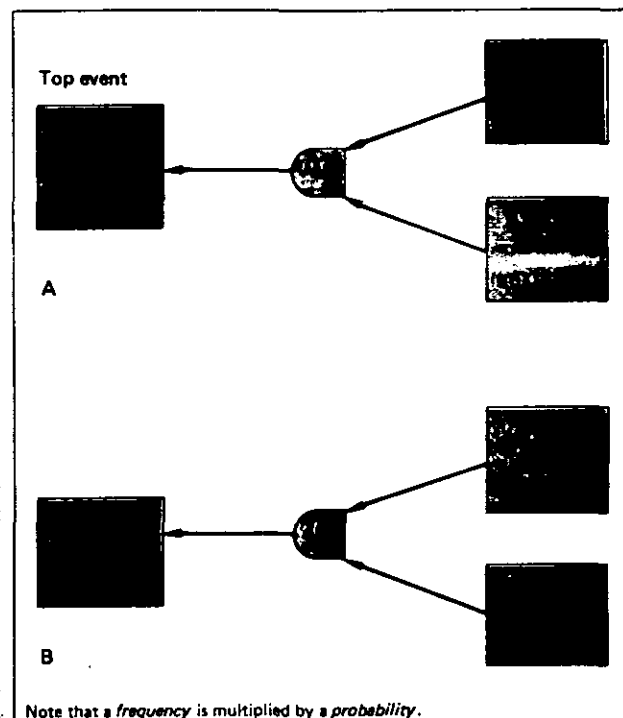


Figure 9 — Fault trees with "AND" gates

In Figs. 9 and 10, we assumed that the probability of being invited to lunch is the same for the two sorts of lunch. This may not be so. In Fig. 11, Fig. 10 (a) has been redrawn to allow for the fact that the probability of being invited to lunch with visitors may be different from the probability of being invited to lunch with a training course.

An industrial equivalent might be that the probability that an operator will take the correct action when an alarm sounds is not fixed, but differs for different alarms. Some alarms might be more noticeable or the operator might be trained to pay more attention to them.

Let us summarize what has been said about "AND" and "OR" gates. At school, we were taught that AND means add. Remember that in drawing logic trees: OR means add, and AND means multiply (as in probability calculations).

As already stated, estimating hazard rates is not the only use of fault trees. They help us think out all the ways in which the hazard can arise and they show us which branches of the tree contribute the most toward the hazard rate. They show us how we can reduce the hazard rate and which methods will be most effective.

For example, in the case of the free meal, we can reduce the hazard rate — the number of free meals per year — by reducing the numbers of visitors or the number of training courses, or by reducing the probability that we shall be invited. We also see that reducing the number of visitors will be more effective than reducing the number of training courses.

Pitfalls in hazard analysis

So far, the methods of hazard analysis appear straightforward. However, a number of pitfalls await the unwary. Two have already been discussed in "Example 3: Frequent demands on a trip," and under "Fault trees." Others are

discussed below. We start with data. Although errors in data, as shown in the section on "Fault trees," are not the most important errors, they nevertheless do occur and we should be on the lookout for them.

Data may be inapplicable

For example, published data on pumps may apply to different types, liquids, pressures, temperatures, corrosivities, etc. If we use the data without checking that conditions are similar, we may introduce serious errors. Leakage rates from flanged joints in a factory handling a corrosive chemical have been found to be many times higher than in a factory handling clean petroleum liquids.

Instruments are similar wherever they are installed and their failure rates in different industries are unlikely to differ by a factor of more than 3 or 4 [9]. This is not true of mechanical equipment.

Data apply to the past

Designs may have changed, not necessarily for the better. For example, a component in an instrument might be made of aluminum alloy or plastic instead of steel. The manufacturer regards the change as trivial and does not tell his customers. But the new component fails more frequently or sooner than the old one.

Here is another example: A plant contained equipment to restart it automatically if power failed and was restored within 0.1 s. The manufacturer of the equipment, without telling anyone, changed the delay time to 1 s. This led to an explosion.

Data affected by maintenance policy

On beverage vending machines, for every 100 demands:

The *right drink* was obtained 94 times, whereas the

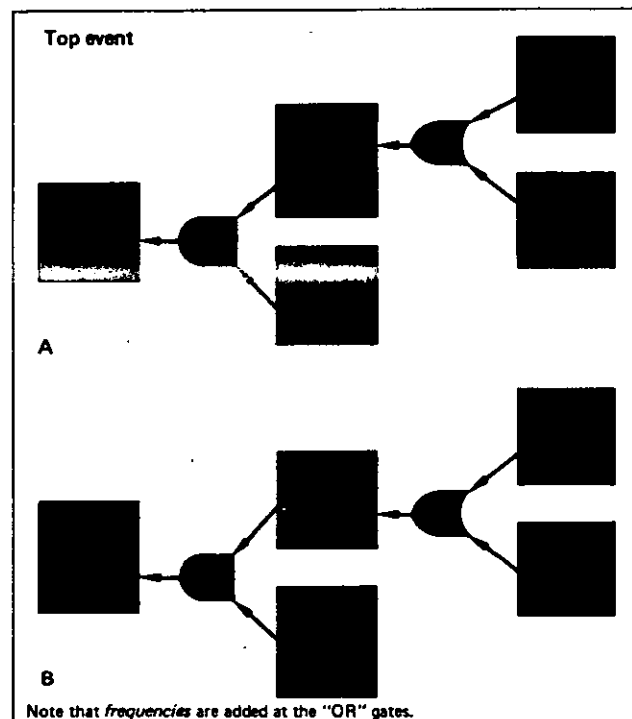


Figure 10 — Fault trees with "AND" and "OR" gates, and showing probabilities

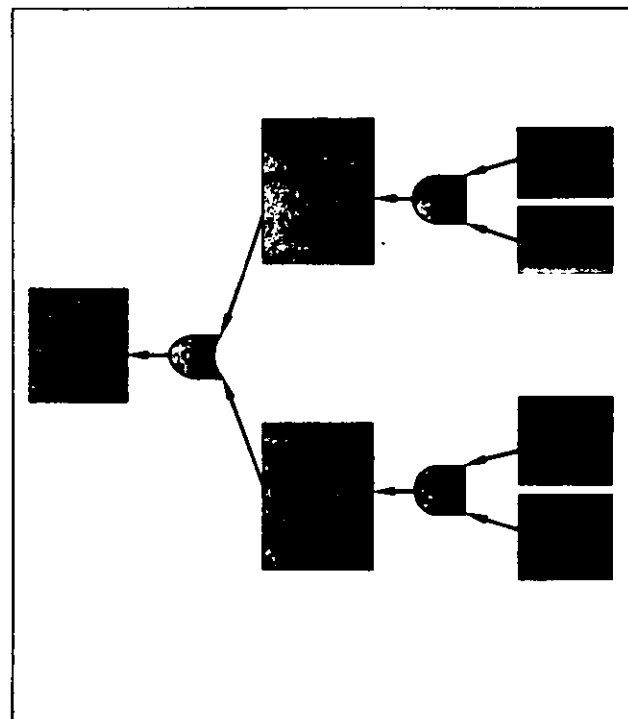


Figure 11 — Fig. 10A redrawn to show probabilities on different branches

FEATURE REPORT

wrong drink was obtained 6 times. Consequently, the *failure rate* = 6%

Before we assume that better machines are needed, let us see how the failure rate is made up.

Wrong drink includes warm drinks, no drinks, short measures, etc. (We must always define what is meant by a failure.)

1. Two of the failures in every 100 were due to the operator pressing the wrong button.

Therefore:

Operator failure rate = 2%

Machine failure rate = 4%

Better mechanical reliability will therefore remove, at the most, two-thirds of the faults. To eliminate the others, we would have to look at the factors that affect operator error (such as better layout of the panel, locating the machine where distraction is less, and so on).

2. 98 demands in every 100 were made on machines in the office and there were 2 failures. The remaining 2 demands

that is more liable to break down or is the management — the system for reporting and repairing faults — different? Perhaps the users treat the machines differently.

Here is a more technical example of the way in which data can be affected by maintenance policy.

Bellows were found to fail at a rate of 1 in 50 per year. Most of the failures did not result in large leaks but they caused shutdowns and loss of production. The failure rate seems high. Do we need a better product?

Analysis of the failures showed that some were due to specifying the wrong material of construction but most were due to poor installation.

The failure rate does not give us information about bellows, but rather about the engineers who specify and install them.

If we wish to reduce the failure rate, we should:

- Specify the material of construction correctly.
- Take more care over installation.

The first should not be difficult but the second one is.

In practice, bellows should be avoided when possible (by building expansion bends into the pipework) and more care taken over the installation of those we have to use.

A man had three Ford cars and crashed each of them, so he decided to try another make. Does this tell us something about Ford cars or rather about the man?

Too-low fractional deadtime

Consider a 1-out-of-3 system.

Assume that the fractional deadtime of each system = 10^{-2} .

Then the fractional deadtime of the total system = $2 \times (10^{-2})^2 = 2 \times 10^{-6}$ (or, in other words, a deadtime of only 1 minute per year).

(It would be 10^{-6} if testing were staggered.)

Do we really believe that our instrument engineers can provide us with a

protective system that is dead for only 1 minute per year?

This calculation is wrong because it ignores two factors:

1. The time the trips are out of action for testing.
2. Common-mode failures — for example, all three instruments are from the same manufacturer's batch and have a common manufacturing fault, or all three instruments are affected by contaminants in the instrument air or process stream, or all three impulse lines are affected by mechanical damage or flooding of a duct, or all three instruments are maintained by the same technician who makes the same error in working on each one. Two or three protective systems are never completely independent.

Therefore, we assume that the fractional deadtime of a redundant system is never less than 10^{-4} (that is, 1 hour per year) and is often only 10^{-3} (that is, 10 hours per year).

As we get 10^{-4} with two trips, a third trip is not worth installing (except as part of a voting system).

For example, wearing a second pair of braces (suspenders) attached to the same buttons may reduce the chance of our trousers falling down. Failure of the buttons (the common

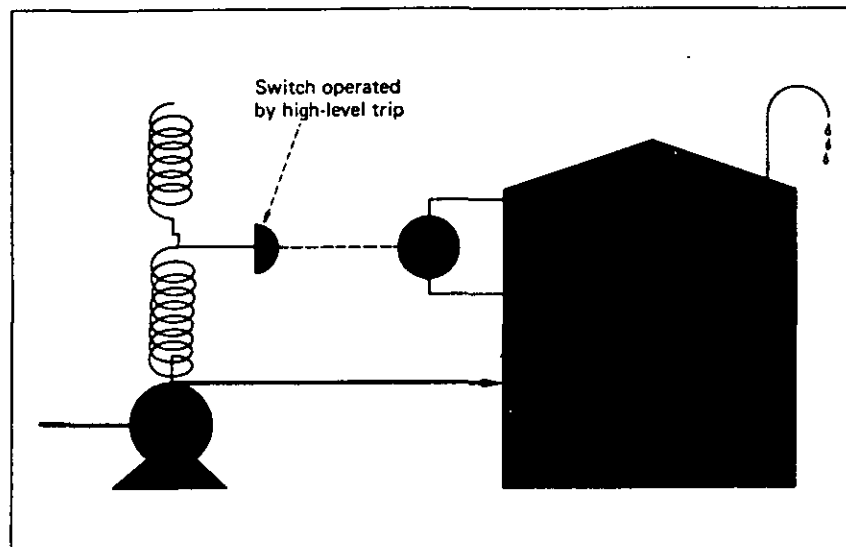


Figure 12 — Storage tank that has been fitted with a high-level trip

were made on machines in a local entertainment center and every demand (2% of the total) resulted in a failure.

Therefore:

Operator failure rate = 2%

Machine failure rate — office = 2%

Machine failure rate — entertainment center = 100%

This shows that misleading results can be obtained if we average widely differing data. (You can drown in a lake of average depth 6 in.)

3. One failure was due to a broken cup.

Hence:

Operator failure rate = 2%

Failure rate due to raw-material quality = 1%

Machine failure rate — office = 1%

Machine failure rate — entertainment center = 100%

We now see that a more reliable machine would reduce the failure rate by only 1%. We could do as well by buying better cups or perhaps by redesigning the panel to reduce operator error.

Are the machines at the entertainment center of a type

mode) is now the biggest cause of failure and adding a third pair of braces, attached to the same buttons, will make no further improvement.

With a diverse system (that is, one in which the approach to a hazardous condition is measured in different ways — say by a change in an analysis, a change in pressure and a change in temperature), 10^{-5} (6 minutes per year) may be possible. For example, belt and braces are better than two pairs of braces.

This example illustrates the perils of using thorough mathematics and ignoring practicalities.

Designer's Intentions not followed

The tank shown in Fig. 12 was filled once/day. Originally, the operator switched off the pump when the tank was full. After five years, the inevitable happened. One day, the operator allowed his attention to wander and the tank was overfilled. A high-level trip was then installed.

To everyone's surprise, the tank was overfilled again after only one year.

Why? What happened was that the trip had been used as a process controller to switch off the pump when the level rose to the setpoint. The operator no longer watched the level. The manager knew this but thought that better use was being made of the operator's time.

When the trip failed, as it was bound to do after a year or so, another spillage occurred.

It is almost inevitable that the operator will use the trip in this way. We should either remove the trip and accept an occasional spillage or install two trips — one to function as a process controller and one to act when the controller fails. The single trip *increased* the probability of a spillage.

In this example, we saw that no trip was a reasonable solution, as was a good trip. The compromise solution was a waste of money. On occasion, either of two extremes makes sense but a compromise does not.

(Note: Because this is true of instrumentation, do not assume it is true elsewhere.)

Non-random failures

A new plant had two 100% compressors (1 working, 1 spare). The failure and the time required for repair were known. Calculation showed that if failures were random, the offline time would be 0.04% (3 h/yr). The actual offline time was 1.8% (144 h/yr).

Why? The failure rates and repair times were as expected, but the failures were not random; most occurred soon after a compressor had been put online.

This may have been due to wrong diagnosis of the fault, installation of wrong parts, or incorrect reassembly.

Mathematical techniques (Weibull analysis) for handling non-random failure are available, but one must recognize that they need to be used.

Motor cars provide another example of non-random fail-

ure — they are more likely to require attention during the week after servicing than at any other time.

If you had two cars (one working, one spare) and one had just been serviced, would you leave it unused until the other broke down or required servicing?

The man in the middle

Fig. 13 illustrates a common plant situation. When the alarm sounds, the operator has to go outside and close a valve within, say, ten minutes. He is well-trained and knows what to do, is physically and mentally capable of doing it, and wants to do it.

The reliability of the alarm is known. If it is too low, it is easy to improve it by adding redundancy or diversity — that is, by adding identical components in parallel, or different components capable of performing the same function. The reliability of the valve is known roughly, and if we do not think it is high enough, we can use a better-quality valve or two valves in series. But what about the reliability of the

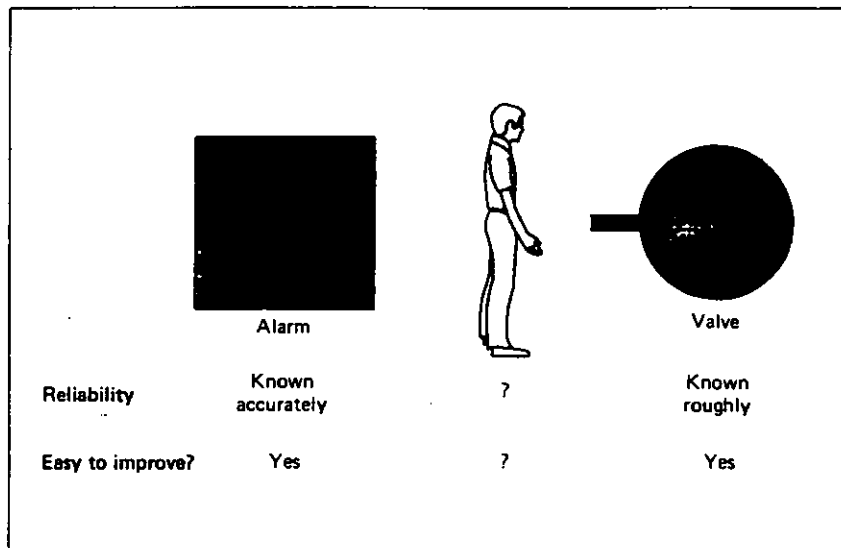


Figure 13 — Comparison of the reliabilities in a man/machine system

operator? Will he always close the right valve in the required time?

At one time, people assumed he would — or should. If he did not, he should be told to pay more attention. Other people have gone to the other extreme and said that sooner or later all operators make mistakes. "We need fully automatic equipment."

Both these extremes are unscientific. We should not say "The operator always should" or "The operator never will" but ask *how often* he will close the right valve in the required time. The answer will depend on the degree of stress and distraction. The following suggestions may help us make a judgement:

Suggested human failure rates

Rate = 1 in 1 (i.e., one failure per demand): This may come about when complex and rapid action is needed to avoid a serious incident. The operator will not really be as unreliable as this but will still be very unreliable, so we should assume this figure and install fully automatic systems.

Rate = 1 in 10: In a busy control room where other alarms are sounding, the telephone is ringing, people are asking for permits-to-work and so on.

Rate = 1 in 100: In a quiet control room, (for example, a storage-area control room) — if the operator is present.

A figure between these last two may be estimated.

Rate = 1 in 1,000: If the valve to be closed is immediately below the alarm.

In carrying out a familiar routine, such as starting up a batch reactor, a typical failure rate is 1 in 1,000 for each operation (for example, "close valve"). Some of these failures will be immediately apparent but others will not [10].

Note that the above figures assume that the operators are well-trained, capable and willing. It is impossible to give a figure for the probability that this assumption is correct, as it can vary from 0 to 1, depending on the policy of the company. We can, however, make a rough estimate of the probability that an operator will have a moment's aberration — as we all do in everyday life — and forget to carry out a prescribed task.

Remember that not all tasks can be prescribed. Sometimes the operator has to *diagnose* the correct action from the alarm and other instrument signals and may not do so correctly, particularly if the instruments are not reading correctly. This happened at Three Mile Island [11].

Finally, remember that installing a fully automatic system does not remove our dependence on personnel. Instead of relying on the operator, we are now dependent on the engineers who design, install, test and maintain the fully automatic equipment. They also make mistakes. They work

under conditions of less stress, so we may improve the overall reliability by installing fully-automatic systems but we should not delude ourselves that we have removed our dependence on people.

A final note

To many people, the calculations of this article, and others on the subject, may seem coldblooded or even callous.

Safety, like everything else, must be bought at a price. The more we spend on safety, the less we have with which to fight poverty and disease, or to spend on those goods and services that make life worth living, for ourselves and others. Hence, whatever money we make available for safety we should spend in such a way that it produces the maximum benefit. There is nothing humanitarian in spending lavishly to reduce a particular hazard that has been brought to our attention and ignoring the others.

Those who make the sort of calculations described in this paper, far from being coldblooded or callous, are the most effective humanitarians, as they allocate the resources available in a way that will produce the maximum benefit to their fellow humans.

Further reading

For detailed accounts of other hazard and operability studies, see Refs. [1,2,12-17]. The first five papers describe the same study. Ref. [12] gives the most detailed account.

For further detailed accounts of other hazard analyses, see Refs. [5,12,18,19].

Roy V. Hughson, Editor

References

1. "Hazard and Operability Studies," Chemical Industries Assn., London, 1977.
2. Knowlton, R. E., "An Introduction to Hazard and Operability Studies," Chemetics International, Vancouver, 1981.
3. Kletz, T. A., *Chem. Eng. Prog.*, Vol. 72, No. 11, p. 48, Nov. 1976.
4. Lees, F. P., "Loss Prevention in the Process Industries," Chapter 21, Butterworths, London, 1980.
5. Lawley, H. G., *Reliability Eng.*, Vol. 1, No. 2, p. 89, Oct.-Dec., 1980.
6. Kletz, T. A., *Chem. Process. London*, p. 77, Sept. 1974.
7. Kletz, T. A., and Lawley, H. G., *Chem. Eng.*, p. 81, May 12, 1975.
8. Gibson, S. B., *Chem. Eng. Prog.*, Vol. 72, No. 2, p. 59, Feb. 1976.
9. Lees, F. P., Institution of Chemical Engineers Symposium, Series No. 47, Process Industry Hazards — Accidental Release, Assessment, Containment and Control, p. 73, 1976.
10. Kletz, T. A., Proceedings of the Third International Symposium on Loss Prevention and Safety Promotion in the Process Industries, Swiss Society of Chemical Industries, p. 2/205, 1980.
11. Kletz, T. A., *Hydrocarbon Process.*, Vol. 61, No. 6, p. 187, June 1982.
12. Lawley, H. G., *Chem. Eng. Prog.*, Vol. 70, No. 4, p. 45, Apr. 1974.
13. As Ref. 4, Chapter 8.
14. Kletz, T. A., "Hazop and Hazan — Notes on the Identification and Assessment of Hazards," Institution of Chemical Engineers, London, 1983.
15. Lawley, H. G., *Hydrocarbon Process.*, Vol. 55, No. 4, p. 247, April 1976. Reprinted in Vervain, C. H., ed., "Fire Protection Manual for Hydrocarbon Processing Plants," Vol. 2, p. 94, Gulf Publishing Co., Houston, Tex., 1981.
16. Rushford, R., North-East Coast Institution of Engineers and Shipbuilders: Transactions, Vol. 93, p. 117, Mar. 21, 1977.
17. Austin, D. G., and Jeffreys, G. V., "The Manufacture of Methyl Ethyl Ketone from 2-Butanol," Chapter 12, Institution of Chemical Engineers, London, 1979.
18. Kletz, T. A., and Lawley, H. G., in "High Risk Safety Technology," ed. by Green, A. E., Chapter 2.1, Wiley, London, 1982.
19. As Ref. 4, Chapter 9.
20. Lees, F. P., Proceedings of the Third International Symposium on Loss Prevention and Safety Promotion in the Process Industries, Swiss Society of Chemical Industries, p. 6/426, 1980.
21. Kletz, T. A., *Reliability Eng.*, Vol. 3, No. 4, p. 325, July 1982.
22. Troyan, J. E., and Le Vine, L. Y., *Loss Prev.*, Vol. 2, p. 125, 1968.

Acknowledgements

I would like to thank the many colleagues, past and present, who contributed ideas for this paper, and in particular Dr. H. G. Lawley who kindly provided the detailed example of a hazop, based on one he has used in teaching the subject. I would also like to thank the Institution of Chemical Engineers for permission to quote extensively from Ref. [14], which should be consulted for other examples and for a more detailed treatment of some of the ideas, and also the U.K. Science and Engineering Council for financial support.

Reprints

Reprints of this 22-page report on safety will be available shortly. To order, check No. 121 on the Reprint-Order Form in the back of this or any subsequent issue. For a free catalog of reprints, circle No. 305 on the Reader Service Card.



The author

Trevor A. Kletz is a senior research fellow and professor at the University of Technology, Loughborough, Leicestershire LE11 3TU, U.K. He studied chemistry at the University of Liverpool, and joined Imperial Chemical Industries in 1944. He later carried out various production jobs, ranging from plant manager to assistant works manager. In 1968, he was appointed Safety Advisor to the Heavy Organics Chemicals Div. (later the Petrochemicals Div.), with responsibility for process safety. He has published over 70 articles on safety topics. He retired from ICI in 1982. He is a Fellow of the Fellowship of Engineering, and a member of the Royal Soc. of Chemistry, the Institution of Chemical Engineers and AIChE.