



When Failure is an Option: Redundancy, reliability and regulation in complex technical systems

John Downer

When Failure *is* an Option

Redundancy, reliability and regulation in complex technical systems

John Downer

Contents

Introduction	2
Basic principles	4
Complexity	6
Independence	7
Propagation.....	11
Human elements	14
Conclusion.....	18
References	20

The support of the Economic and Social Research Council (ESRC) is gratefully acknowledged. The work was part of the program of the ESCRC Centre for Analysis of Risk and Regulation.

Published by the Centre for Analysis of Risk and Regulation at the
London School of Economics and Political Science
Houghton Street
London WC2A 2AE
UK

© London School of Economics and Political Science, 2009

ISBN 978-0-85328-395-9

All rights reserved

No part of this publication may be reproduced, stored in a retrieval system, or transmitted, in any form or by any means, without the prior permission in writing of the publisher, nor be otherwise circulated in any form of binding or cover other than that in which it is published and without a similar condition including this condition being imposed on the subsequent purchaser.

Printed and bound by Kube, May 2009

When Failure is an Option

Redundancy, reliability and regulation in complex technical systems

John Downer

Abstract

This paper argues that redundancy in engineering, should be understood as a ‘design paradigm’ that frames regulatory assessments and interpretations of all complex technical systems, profoundly shaping decisions and judgements about modern technologies. It will further argue that the ‘redundancy paradigm’ used by regulators contains epistemic ambiguities that lead to imperfect predictions about the effects of redundancy in practice. By deconstructing the logic of redundancy in relation to aviation regulation, this paper illuminates much wider issues about technology governance.

There is no safety in numbers

~ James Thurber

Introduction

On 24 June 1982, 37,000 feet above the Indian Ocean, a British Airways Boeing 747 began losing power from its four engines. Passengers and crew noticed a numinous blue glow around the engine housings and smoke creeping into the flight deck. Minutes later, one engine shut down completely. Before the crew could fully respond, a second failed, then a third and then, upsettingly, the fourth. The pilots found themselves gliding a powerless Jumbo-Jet, 80 nautical miles from land and 180 from a viable runway.

Unsurprisingly, this caused a degree of consternation among the crew and – despite a memorably quixotic request that they not let circumstances ruin their flight – the passengers, who began writing notes to loved ones while the aircraft earned its mention in *The Guinness Book of Records* for ‘longest unpowered flight by a non-purpose-built aircraft’. Finally, just as the crew were preparing for a desperate mid ocean ‘landing’, the aircraft spluttered back into life: passengers in life-jackets applauded the euphony of restarting engines and the crew landed them safely in Jakarta (Tootell 1985).

The cause of the close call was quickly identified by incident investigators. They determined that ash from a nearby volcano had clogged the engines causing them to fail and restart only when the aircraft lost altitude and left the cloud. This unnerved the aviation industry; before the ‘Jakarta incident’ nobody imagined a commercial aircraft could lose all its propulsion.

A Boeing 747 can fly relatively safely on a single engine so with four it enjoys a comfortable degree of redundancy. This is why Boeing puts them there. Under normal circumstances, engineers consider even a single engine failure highly unlikely, so, until British Airways inadvertently proved otherwise, they considered the likelihood of all four failing during the same flight to be negligible – as near to ‘impossible’ as to make no difference. In fact, the aviation industry considered a quadruple engine failure so unlikely that they taught pilots to treat indications of it as an instrumentation failure. Nobody had considered the possibility of volcanic ash.

The aviation industry had been misled by redundancy before. Peter Galison (2000) discusses another ‘impossible’ aviation failure, this time with an unhappier ending. In July 1989, the tail-mounted engine of a United Airlines DC-10 exploded. The two wing-mounted engines ensured that the aircraft, with its 285 passengers, still had ample thrust but shrapnel from the explosion tore into the fuselage, severing all three of the triple

redundant hydraulic systems and rendering the airplane uncontrollable.¹

MacDonald Douglas had designed the DC-10 to resist shrapnel. Each of the hydraulic systems had its own redundant pumps connected to redundant (and differently-designed) power sources with redundant reservoirs of hydraulic fluid. As with a quadruple engine failure, therefore, the aviation community had deemed a triple hydraulic failure impossible. They argued this was ‘so straightforward and readily obvious that [...] any knowledgeable, experienced person would unequivocally conclude that the failure mode would not occur’ (Haynes 1991). Again there was no training for such an eventuality.

These incidents, and others like them, are significant because all civil aircraft depend heavily on redundancy for their safety. If a critical system fails, another should always be there to do its work. This is true even of airframes, which are designed with redundant load-paths. Aeroplanes are far from unique in this respect, redundancy is the single most important engineering tool for designing, implementing, and – importantly – *proving* reliability in all complex, safety-critical technologies. It is the *sine qua non* of ultra-high reliability engineering: deeply implicated in everything from bridges and skyscrapers to computer networks and power plants. Rochlin et al. (1987) identify it as a key strategy of successful high reliability organizations and speak disparagingly of ‘mechanistic management models’ that seek to eliminate it in the name of efficiency.

Although redundancy may seem abstract and esoteric, it is a foundation stone of our knowledge about the technologies we depend on and is reflected in all the decisions we make about them. It lies at the heart of all modern reliability calculations and is central to the way we navigate what Beck (1992) refers to as the ‘risk society’.

Charles Perrow equates the engineering conception of redundancy with a Kuhnian scientific paradigm (1984: 196). Paradigms have almost become a cliché but the analogy is unusually apposite here. Petroski (1994) argues persuasively for the relevance of ‘design paradigms’ to engineering, describing them as engineering principles that are common across ostensibly disparate specialities and give engineering a theoretical foundation.

Thinking of redundancy as a design paradigm is useful because it highlights that, like a scientific paradigm, redundancy acts as a conceptual lens through which to ‘know’ the objects it frames. Redundancy allows engineers to build reliable systems, but the ‘lens’ of redundancy, this paper will argue, allows them to see levels of reliability far beyond those that would be visible in a laboratory. The significance of the latter is difficult to overemphasize. Risk quantification has become an essential element of modernity (Power 2007) and a panoply of oversight bodies and regulatory professionals exist to measure

¹ Amazingly, there was a credible attempt at a landing. The crew obtained some control by manipulating the power to different engines and were able to save some of the passengers.

and verify the reliability of dangerous technologies. The Federal Aviation Administration (FAA), for example, both demand and verify that every safety-critical element of a civil airliner has a reliability of no more than ‘one failure in a billion hours of use’.

Redundancy makes this plausible; it offers the surety that publics and policy-makers require to audit (and thereby *govern*) the technological world.

Like all paradigms, however, redundancy can distort as well as clarify. Its simple premise belies a deceptive complexity, and those who fail to recognize its nuances are prone to dangerously misconceive important decisions. This paper will explore the ambiguities of redundancy and their implications for its role as a design paradigm. It will look at how redundancy shapes our relationship to technologies, our knowledge of them and our choices regarding them.

Basic principles

The word redundancy has a specific meaning in engineering: one that is simultaneously similar and yet strangely at odds with its common usage. Both common and engineering usages imply some manner of ‘repeating’, but, where redundancy usually has negative connotations in its common usage – as something superfluous and excessive – in modern engineering it is equated with safety and integrity (Landau 1969).

The engineering practice it denotes looks simple in principle. An element is redundant if it contains backups to do its work if it fails; a system is redundant if it contains redundant elements. This can mean having several elements that work simultaneously but are capable of carrying the ‘load’ by themselves if required – such as the engines on civil airliners – or it can mean having idle elements that ‘awake’ when the system needs them, such as backup generators. Redundancy can work at different levels: sometimes it is two capacitors on a circuit-board in case one fails and at others it is the duplication of a whole system, such as in the US missile programmes of the 1950s, where the Air Force entrusted mission success to the doctrine of ‘overkill’ and the redundancy of the entire missile (Swenson et al. 1998: 182). Often it means both.

Redundancy has served as a central tenet of high reliability engineering for over 50 years. It was first proposed as a design solution for reliability in complex tightly-coupled systems by Jon von Neumann in his classic 1956 treatise: ‘Probabilistic Logics and Synthesis of Reliable Organisms from Unreliable Components’, published in Princeton’s *Automata Studies*. His innovation lay not in the idea of duplicating elements itself but in envisaging a system that could make use of duplication by recognizing when it should engage backup elements. Grappling with the complexities of building a working computer from thousands of unreliable vacuum tubes, von Neumann realized – radically – that a redundant system could be *more* reliable than its constituent parts. ‘The basic idea in this procedure is very simple,’ he writes, ‘Instead of running the incoming data

into a single machine, the same information is simultaneously fed into a number of identical machines, and the result that comes out of a majority of these machines is assumed to be true. [...] this technique can be used to control error.’ (1956: 44). His insight represented a significant breakthrough at a time when the reliability of Western military paraphernalia was proving inadequate and engineers were pushing the limits of their ability to increase the dependability of a system by maximizing the reliability of each component. It was exemplary ‘system level’ engineering: enhancing the integrity of the system by allowing it to adapt and compensate for the shortcomings of its elements (Jones-Imhotep 2000: 153).

Although first established as a design tool, engineers soon realized that redundancy’s practical virtues were coupled with an epistemological advantage: redundancy not only allowed them to *design* for high reliability, it also allowed them to *quantitatively demonstrate* reliability – something that had previously proved very difficult. Redundancy could do this because it offered a powerful, straightforward, and convincing rubric with which engineers could mathematically establish reliability levels much higher than they could derive from lab testing.

Put simply, this was because engineers could multiply the reliability of redundant elements to ascertain the reliability of the system as a whole. The basic principle is again straightforward: where two redundant and independent elements operate in parallel to form a system, we can, in principle, calculate the chance of total system failure by establishing the probability that both systems will fail at the same time (Littlewood et al. 2002: 781). If each element only fails infrequently, then there is a vanishingly small probability of both failing simultaneously. For instance, if one element has a 0.0001 probability of failing over a given period, then the probability of two redundant elements failing over the same period is that number squared $(0.0001)^2$ or 0.00000001. By using redundancy, therefore, we have demonstrated a *ten-thousandfold* reduction in the probability of failure and, significantly, without recourse to lab tests (Shinners 1967: 56).

Whilst vital to modern technology regulation, however, these calculations are deceptively complex. Critics argue that they rely on unrealistic assumptions and ignore important variables (e.g. Littlewood & Strigini 1993; Littlewood et al. 1999; Littlewood 1996; Popov et al. 2000). Such detractors rarely object to the use of redundancy *per se*, but to the way engineers or regulators use it in calculations. In different ways, they articulate the common observation that the platonic niceties of mathematics fit imperfectly onto the ‘messiness’ of the real world.

Experience seems to support a disjuncture between the performance of redundancy in theory and its capability in practice: failures happen more often than engineering calculations predict (Perrow 1984). The following sections seek to explain why. They unpack some of the ‘messiness’ of redundancy in the real world, by visiting broad

epistemic difficulties, illustrating their relevance and explaining how they subvert the mathematical ideal of the ‘redundancy paradigm’.

The first looks at the dilemmas of complexity.

Complexity

Simplicity and reliability go hand in glove, or so many commentators suggest (e.g. Perrow 1984). Mary Kaldor, for instance, argues that Soviet technology was often more reliable than its Western equivalent because it was straightforward. Russian technology was ‘uncomplicated, unadorned, unburdened,’ she writes, ‘performing only of what [was] required and no more’ (1981: 111). She gives the absence of powered controls on large passenger jets as an example.

Redundancy upsets this relationship. It increases the number of parts in a system leading to more unexpected interactions and leaving the system harder to understand and to verify. Modern airliners, assembled from over a million parts, are extraordinarily complex machines. The electronics alone are Byzantine. Inside each aircraft lies a highly sophisticated computer network with hundreds of miles of wiring linking dozens of computer systems that regulate everything from in-flight entertainment to engine temperature.

Several students of complex systems argue that increasing redundancy can exacerbate complexity to the point where it becomes the primary source of unreliability (e.g. Rushby 1993; Hopkins 1999). Perrow makes this argument about the Space Shuttle, suggesting that its extensive redundancy makes its workings mysterious and unpredictable even to its designers, and that this is a source of risk (1984: 270).²

The extra elements needed to manage redundant systems deepen this problem. Redundant elements invariably require further ‘managerial’ systems to determine, indicate, and/or mediate failures (Shinners 1967: 57). Should a commercial aircraft engine fail during flight, for example, an electronic supervisor senses the failure, cuts the fuel, adjusts the rudder, and compensates for the missing thrust (Rozell 1996). It may sound simple to bolt extra engines on to an aeroplane, but this simplicity quickly dissolves if we consider the many extra management systems and sensors it entails, any one of which, we should remember, might fail and cause its own accident. Even if the system relies on a human mediator, that mediator relies on dials, sensors, and other indicators, all of which can fail. (And, of course, humans have their own reliability problems.)

² Sagan (1993) has found that false indications from safety devices and backup systems in the American missile-warning apparatus have nearly triggered nuclear wars!

More than simply exacerbating a system's complexity, redundancy management systems introduce a further complication. Since they are themselves critical elements of the system, *they too* must be 'ultra-reliable.' Or, in engineering parlance: 'a centralized redundancy management function will likely inherit the criticality level of the most critical task that it supports' (Rushby 1993: 69).

Redundancy management systems have certainly caused their own disasters. In January 1989, for example, a British Midland 737-400 crashed at Kegworth, near Nottingham, killing 44 people. Accident investigators think that one of the aircraft's two – redundant – engines caught fire, and that somebody had miswired the – not redundant – warning system, leading the pilot to shut down the wrong engine (Krohn & Weyer 1994: 177).³ Unlike at Kegworth, management systems have sometimes even instigated disasters.⁴

Aircraft manufacturers must make management systems reliable but they cannot make them redundant as this would lead to an infinite regress. If mediating elements were themselves redundant then they, too, would require mediation. 'The daunting truth,' to quote a 1993 report to the FAA, 'is that some of the core [mediating] mechanisms in fault-tolerant systems are single points of failure: they just have to work correctly' (MacKenzie 2001: 229). The implicit hope is that engineers can make mediating systems simpler than the elements they mediate and therefore more reliable, but redundancy management systems have a way of becoming highly complex and, as such, especially prone to failure.⁵

Even if manufacturers *can* make mediating systems simple and reliable, the limits of empirical lab tests still preclude them from demonstrating or proving that reliability to an ultra-high level. Often they can partially bypass this problem by delegating mediating functions to a person – such as a pilot – rather than a machine (more on this below). On other occasions the problem dissipates because regulations do not classify mediating elements as 'safety critical,' and so manufacturers are excused from 'proving' these elements in the same way as the systems they mediate. Either way, the underlying epistemology is fragile.

Independence

On 4 June 1996, the European Space Agency (ESA) collectively winced as their heavy-lift rocket, *Ariane 5*, veered from its course 39 seconds into its maiden flight and then

³ Interestingly, the pilot disregarded vibration gauges indicating the error because these had gained a 'common knowledge' reputation for unreliability.

⁴ Rushby (1993) gives a list.

⁵ MacKenzie (2001: 228-29) highlights some of the problems inherent to redundancy management. He shows that it can be extremely complex for a management system to discern when one of the systems it is managing is malfunctioning, and which system is in error.

disintegrated with its cargo of four expensive and uninsured satellites. The ESA subsequently identified a glitch in the rocket's guidance computer as the cause. The software generated a number too big for the system to handle and so the computer shut down and passed control to its redundant twin, which, being identical to the first, came to the same conclusion and shut down a few milliseconds later. The rocket, now without guidance, changed direction to compensate for an imagined error and collapsed in its own turbulence (Gleick 1996).⁶

A significant critique of redundancy lies in the observation that many calculations assume that redundant systems behave completely independently of each other (Popov et al. 2003: 2). To call two elements 'independent' is to say the chances of one failing are not linked in any way to the chances of the other failing. This assumption underlies all the simple reliability calculations outlined above, yet, as Ariane 5 demonstrated, it is far from safe.

Indeed, there are many reasons to doubt the independence of redundant systems. 'Identical' elements will likely wear in similar ways and, consequently, fail at similar times when they both operate simultaneously. Most failures result from external pressures acting on a system and redundant elements in close proximity will likely face the same pressures at the same time. In this way, 'operating environment' can act as a source of interdependence.⁷ An external pressure such as a violent storm, cloud of ash, or flock of birds might stress all the elements in a system, as might an 'internal' event such as an unanticipated power drain.

Elements can also fail as they lie idle. Even as they wait in reserve, they may fall victim to external pressures such as vibrations or moisture. They may also have been faulty on installation. Such 'latent' or 'dormant' failures pose particular dangers because they can go undetected. Airlines regularly deal with this problem, as in May 1995, when the FAA criticized the Boeing 737's rudder control system. The system involved two 'slides': one, which did most of the work, and a second – redundant – slide that lay in reserve. The FAA argued that since the system rarely uses the second slide it could fail 'silently', leaving the aircraft 'a single failure away from disaster' for long periods of time (Acohido 1996).

The engineering problems arising from the question of 'independence', however significant, are more tractable than the epistemological problems. If there is good reason to assume that the failure-behaviour of different elements is not independent, then redundancy-reliability calculations become much more complicated because they now require a quantitative measure of independence and engineers face the daunting task of measuring the likelihood of two elements failing at the same time.

⁶ Probably the most costly software bug in history.

⁷ See for instance, Eckhard & Lee 1985; Hughes 1987; Littlewood 1996; Littlewood & Miller 1989.

One solution to the problem of measuring independence lies in testing two elements together as a single system and thereby getting an empirical figure for how often they fail simultaneously. Testing a system in this way, however, does not bring engineers closer to achieving the ultra-high reliability figures that engineering standards demand (and redundancy should deliver), because the tests are constrained by the same practical limits that redundancy supposedly allows engineers to transcend. At such high levels of reliability the probability of simultaneous failures are almost impossible to determine in a lab (Littlewood & Strigini 1993: 10).

Engineers must therefore deal with the problem of independence in other ways. To do this they usually rely on variations of what they call ‘design diversity’. Broadly speaking, design diversity is the practice of designing redundant elements differently whilst keeping their functions the same: producing interchangeable black boxes with different contents. The idea is that if elements differ from each other they will have different weaknesses and fail in different ways and at different times.

Manufacturers approach design diversity in varying ways. Sometimes they leave it to evolve spontaneously by giving responsibility for different elements to isolated design teams and hoping a lack of central authority will result in different designs with independent failure behaviour (Littlewood & Strigini 1993: 9). Sometimes they adopt an opposite approach where a single authority actively promotes diversity by explicitly requiring different teams to use divergent approaches, solutions, and testing regimens. An example from software engineering would be forcing teams to program in different languages (Popov et al. 2000: 2).

Both approaches are imperfect, however. The idea that different groups, left to their own ‘devices’, will design the same artifact differently finds some support in the social construction literature, which highlights the contingency – on some level – of technological designs by arguing that different groups have more ‘technological options’ than it might appear (Bijker et al. 1989). Yet the same literature also suggests that where designers come from similar professional cultures, and where they have problems specified for them in similar ways, their designs will likely converge. The ‘separate design team’ approach has its limitations, and we certainly cannot assume it produces mathematically perfect independence.

The same literature also suggests that ‘forced’ diversity has similar epistemic shortcomings. Engineers designing dissimilar artefacts require a well-defined notion of ‘dissimilarity’, yet ‘dissimilarity’ – like truth, beauty, and contact lenses – inevitably rests in the eye of the beholder. Harry Collins’ (1985) analysis of ‘similarity’ applies equally to ‘dissimilarity’. He argues that to replicate an artefact (or process) precisely, we need a precisely defined notion of ‘similarity’ (65). When building the ‘same’ TEA laser, for

example, engineers had to determine whether the gauge of the wires and their colour should be the same.

His point is that ‘similarity’ (and therefore ‘dissimilarity’) must be restricted to a finite number of variables. To force diversity between redundant elements, engineers must understand what counts as ‘different’. Should they use wires of different gauges, or should they allow only one element to use wires and force the other to use an entirely different method of conveying electricity (and should they both use electricity)? We can never make diversity ‘absolute’, so the term always has a bounded and socially negotiated meaning, yet the way we restrict it will always represent a way that we have restricted the ‘independence’ of different elements and undermined any ‘proof’ that rests on redundancy calculations.⁸

A further limitation, common to both forced and unforced diversity, lies in the fact that, even when designed in different ways, redundant elements are often expected to act in the same environment. This assumption alone imposes limitations on potential diversity. Even very different designs will not achieve complete independence as to how ‘difficult’ they find environmental demands if engineers build them around a similar concept of what constitutes a ‘normal’ or ‘routine’ environment, because all designs will then find the same environments unusual (Littlewood & Strigini 1993: 10). A 20-lb swan will stress engines built to ingest birds smaller than 4 lb whatever their design differences.

A design strategy does exist for reducing the extent to which the environment challenges different elements at the same time. It is an extension of design diversity known as ‘functional diversity’. Systems designed in a ‘functionally diverse’ manner use different inputs. A functionally diverse instrument, for instance, might comprise two redundant elements: one using temperature and the other pressure in that hope that environmental conditions that challenge one element will not challenge the other (Littlewood et al. 1999: 2). Hutchins (1995), for example, explains how US Navy navigators establish a ship’s position using both satellite data and measures of the ship’s bearing relative to known landmarks. The navigators compare results from one source with those from another, and because the information comes from independent sources it offers a high degree of confidence when the sources correspond. In a similar fashion, aircraft sometimes use pressure, radar, and GPS to measure altitude.

Again, however, even functionally diverse systems cannot be expected to offer perfect independence. Seemingly different and separate inputs are often interrelated; extremes of temperature, for instance, will correlate, at least loosely, with extremes in pressure and moisture. Functional diversity also neglects the fact that more environmental factors work

⁸ There are also many practical limitations to this approach, there are only so many technological options, and elements often must co-ordinate with each other, it would be a challenge to combine a jet engine and a propeller on the same aircraft.

on any given element than simply those it uses directly. Strong buffeting, for example, could threaten both a system reading pressure and one reading temperature, as could moisture damage or any number of other factors.

Clearly, therefore, design diversity is far from being a logically incontrovertible or epistemologically simple solution to the problem of independence. This is not to say it is not a useful engineering practice or that it does not lead to increased reliability. Indeed, experiments do appear to show that design-diversity brings increased safety over ‘identical’ redundancy. The same experiments also show, however, that it brings less safety than we would expect of completely independent redundant elements (Littlewood & Strigini 1993: 10). Design diversity, therefore, fits awkwardly into quantitative reliability calculations. It does not guarantee perfect independence in the failure behaviour of redundant elements, and so it does not remove the need for a way of quantifying independence (which, in turn, would require a way of quantifying ‘diversity’).

Propagation

In 2005, a Boeing 777 departing from Perth, Australia, spontaneously pitched upwards, activating stall warnings and startling the crew. On its immediate return to Perth, investigators determined that a faulty accelerometer had caused the incident. The accelerometer had a redundant backup in case it failed but its designers had misconceived *how* it would fail: they assumed a failure would always result in an output of zero volts, but this failure produced a high voltage output, confusing the computer.⁹

It is often difficult to predict ‘how’ machines will fail¹⁰ (as one engineer put it: ‘you can’t always be sure your toilet paper is going to tear along the perforated line’) yet redundancy calculations lean heavily on such predictions. In principle, the failure of a redundant element should not compromise a system, but this principle often depends on the element failing in a predictable way, and this does not always happen.

For instance, on 17 July 1996, just outside Long Island, New York, a Boeing 747 – TWA flight 800 – suddenly and tragically exploded. The cockpit and most of first class broke away from the fuselage, gracelessly plummeting 13,000 feet into the dark Atlantic. The nose-less fuselage stayed aloft under its own momentum for 30 terrifying seconds before beginning its own terminal dive towards the water. The FBI initially suspected sabotage or a missile strike but the subsequent investigation concluded that a spark – probably caused by (poorly understood) corrosion of the aircraft’s ageing wiring – had ignited

⁹ The Australian Transport Safety Bureau's investigation report is at http://www.atsb.gov.au/publications/investigation_reports/2005/AAIR/aaair200503722.aspx

¹⁰ Especially, perhaps, those designed not to fail at all.

volatile fuel vapours in the central fuel tank (Negroni 2000).¹¹ A stark, if fairly obvious reminder that not all failures can be mitigated by redundancy.

TWA 800 illustrates an important dimension of malfunction that redundancy calculations often ignore: failures do not always keep to themselves and have a tendency to *propagate*. This is a straightforward observation, but one with important implications. It means that a complete measure of reliability-from-redundancy needs to account not only for the independence *between* redundant element, but also for the independence of these elements from other – functionally unrelated – elements. In Perrow's terms, this amounts to a measure of the 'coupling' of a system.

'Two engines are better than one,' writes Perrow, 'four better than two' (1984: 128). This seems simple enough and Perrow is probably correct, but his axiom is less simple than it appears. If, when engines fail, they do so catastrophically, fatally damaging the aircraft – as happened to United Flight 232 in July 1989 – then it is not at all clear that four engines are safer than two. In fact, it is hypothetically possible that four engines could be much *less* safe than two.

Imagine, for instance, that an aircraft can either have a configuration of two or four. It can function adequately with only one engine, and the engines enjoy (for the sake of argument) perfect independence in their failure behaviour. The chances of any given engine failing during a flight are one-in-ten. (It is a very unreliable design!) Also, however, one out of every ten engines that fail will explode and destroy the aeroplane. (Of course, the aeroplane also suffers a catastrophic failure if all the engines fail in the same flight.) Now, it follows that an aircraft has a higher chance of an 'all-engine' failure with the two-engine configuration than with four, but the aeroplane enjoys a lower chance of experiencing a single engine failure, and so less chance of an explosion. The maths works out such that the combined risk of any catastrophic event during flight (an all-engine failure or an explosion) is higher with a four-engine configuration than with two. This is to say that two engines would be safer than four!

Boeing has, in fact, come to this very conclusion (albeit using very different probabilities), arguing that its 777 is safer with two engines because of the lower risk of one failing catastrophically (Taylor 1990, cited in Sagan 2004: 938). Even if four engines were unquestionably safer, it would still be true that accurate reliability calculations demand an assessment of the independence of a system's functionally unrelated elements (its *coupledness*), and that, in some instances, redundancy may even *detract* from a system's reliability.

¹¹ The investigation was marked by controversy, with the NTSB vying for authority with the FBI and the latter memorably enrolling a psychic to help with their analysis, much to the professional chagrin of the NTSB and incredulous ire of a congressional panel of inquiry.

Even where failures are not immediately catastrophic in themselves, they often cause new and unexpected links between even functionally unrelated elements. It is not enough to calculate the probability of all four engines failing, for example, because if one engine failed at the same time as the rudder hydraulics then the pilot could not compensate for the uneven thrust and the aircraft would be imperilled, so that too must be calculated, together with an indefinite and incalculable number of other combinations. Even functionally unrelated elements can interact; a fire caused by one component might damage elements it would never usually interact with.

The implications of this ‘coupling’ are often overlooked, and where they are not, they are difficult to quantify with any precision. This is not to say that the potential for unexpected interactions has escaped the FAA and the manufacturers, of course. They use techniques such as ‘failure modes and effect analysis’ in an attempt to map out each possible interaction between elements. Sophisticated though these techniques are, few engineers deny that using them requires as much art as science. Whilst engineers might find these techniques useful as design tools, they fit awkwardly into reliability predictions because their conclusions are difficult to quantify – it being impossible to foresee every possible interaction between parts or to calculate the degree to which one has.¹²

If such analysis is difficult, it is also vital. The majority of fatal accidents involve unanticipated chains of failures, where the failure of one element propagates to others in what the NTSB call a ‘cascade’ (NAS 1980: 41). All systems have elements with ‘catastrophic’ potential (loosely defined as a capacity to fail in a way that might ‘propagate’ damage to other elements in the same system). Good examples of these, Perrow writes, are ‘transformation’ devices, involving ‘chemical reactions, high temperature and pressure, or air, vapour, or water turbulence’, all of which, he says, make a system particularly vulnerable to small failures that propagate unexpectedly and with disastrous results (1984: 9). The accident record supports this view. As well as TWA flight 800, described above, several other commercial aircraft have crashed because their fuel tanks unexpectedly exploded, such as on 8 December 1963, when a lightning strike to a fuel tank left a Boeing 707 – Pan Am flight 214 – buried in a Maryland cornfield. Perrow is probably too limited in his scope here as it makes sense to expand his set of potentially catastrophic elements to include anything that contains large amounts of energy of any kind be it electrical, kinetic, potential or chemical. An element may operate at high pressure or speed, it may be explosive, corrosive, or simply have a potentially destabilizing mass. An element may also propagate failure simply by drawing on a resource required by other systems such as electricity, fuel, or oil. A faulty engine that leaks fuel will eventually threaten the other engines, as in August 2001 when a faulty

¹² In the words of one report: ‘The failure of a neighboring system or structure is not considered within a system’s design environment and so is not taken into account when analyzing possible ways a system can fail.’ (National Academy of Sciences. 1980: 41)

‘crossfeed’ valve near the right engine of an Airbus A330-200 leaked fuel until none remained to power the plane and both engines failed.¹³

Engineers attempt to mitigate the effects of coupling by ‘isolating’ or ‘partitioning’ different elements, physically separating them and shielding them from each other. Manufacturers of military-grade microelectronics, for example, frequently encase them in ceramic.¹⁴ In a similar vein, the FAA require that aircraft designs, (since Concorde) separate engines on the wing, and house them in casings capable of containing the shrapnel from broken fan-blades. Aircraft designers also isolate resources; for instance, they use different computers to manage the in-flight entertainment systems than they use to manage the avionics.¹⁵

While isolating elements undoubtedly helps improve system safety, it again makes calculating reliability more complicated. Like independence, ‘isolation’ is a subjective virtue. During the certification of the Boeing 747-400, for example, the FAA and the JAA¹⁶ differently interpreted an identically worded regulation governing the segregation of redundant wiring. The JAA interpreted the word ‘segregation’ more conservatively than their American counterparts, forcing Boeing to redesign the wiring of the aircraft late in the certification process (GAO 1992: 16).¹⁷ The confusion arose because the rules were open to what Pinch and Bijker (1984) call ‘interpretative flexibility’: the two regulators could not defer to a common and unambiguous yardstick of ‘isolation’.

Hollnagel (2006: 15) neatly summarizes the wider point: ‘Most major accidents are due to complex concurrences of multiple factors, some of which have no *a priori* relations.’ He writes, ‘Event and fault trees are therefore unable fully to describe them – although this does not prevent event trees from being the favourite tool for Probabilistic Safety Assessment methods in general.’

Human elements

On 29 December 1972, just outside Miami, the crew of Eastern Airlines Flight 401 became so fixated with a faulty landing gear light that they failed to notice the autopilot

¹³ The aircraft – *en route* to Lisbon from Toronto with 291 passengers – glided for 115 miles before making a high-speed touchdown in the Azores that wrecked the undercarriage and blew eight of the ten tires. [Significant Safety Events for Air Transit < <http://www.airsafe.com/events/airlines/transat.htm>>]

¹⁴ Mil hdbk 217, Appendix d. Washington DC: Department of Defense.

¹⁵ NASA spacecraft, similarly, observe a rigorous separation between the components and resources that are critical to the craft itself, and those that control the on-board scientific experiments. The consequences of the absence of such strict fault containment are well-illustrated by the failures of the Russian Phobos spacecraft (Rushby, 1993)

¹⁶ The JAA (Joint Aviation Authorities) – now the European Aviation Safety Agency (EASA) – is the FAA’s European counterpart.

¹⁷ Because of this, two different designs of the 747-400 now exist: one for FAA standards and one for JAA standards.

was disengaged. They continued in their distraction until the aircraft smashed into the Everglades, killing 101 of the 176 passengers.

The American National Transportation Safety Board (NTSB) estimates that 43 percent of fatal accidents involving commercial jetliners are initiated by pilot error (Lewis, 1990: 196). A surprising number happen when pilots misread navigational instruments – usually under stress – and fly into the ground. (Euphemistically known as Controlled Flight Into Terrain or CFIT). There were at least 43 CFIT incidents involving large commercial jets in the decade between 1992-2002 (Flight Safety Foundation 2004). Although many such ‘errors’ can indirectly be blamed on design (misleading displays for instance), it is undeniable that people are unavoidably fallible, even with the very best design (e.g. Reason 1990). They get ill. They get impatient, stressed, scared, distracted and bored. They make mistakes.

An often-hidden aspect of redundant systems, therefore, is that they require people to build and work them. Seemingly redundant and isolated elements frequently link to each other at the level of the people who operate and maintain them. Failures are not passive events: they frequently instigate *actions*. A failure of an aircraft system may require the pilot to reduce demand on that system – by cutting the power on a failing engine, for instance – or by altering the flight plan, altitude, or speed (Lloyd & Tye 1982: 14). The pilot may find such actions unusual and s/he may have to perform them under stress whilst making complex inferences from instruments and training.¹⁸ As such, technological failures open a window for human error. A relatively common mistake in twin-engine aircraft, for instance, is for the pilot to respond to an engine failure by shutting down the wrong engine, as happened in the 1989 British Midland crash, outlined earlier (AAIB 1990). Here we see a failure ‘cascade’ with one engine indirectly precipitating the failure of another because of their common link at the level of the pilot.

People other than ‘operators’ can link redundant elements. Investigators of a multiple engine failure on a Lockheed L-1011, for instance, determined the engines were united by their maintenance. The same personnel had checked all three engines, and on each they had refitted the oil-lines without the O-rings necessary to prevent in-flight leakage.¹⁹

As well as simply being a common link between redundant systems, people can sometimes act as redundant elements in and of themselves. One NASA engineer, for example, wrote of the Apollo space capsule that whilst ...

... primary control is automatic, for vehicle operation, man has been added to the

¹⁸ Military pilots evocatively refer to the confusion borne of stress and information overload as ‘helmet fires’.

¹⁹ Similarly, in 1983, all the engines failed on a Boeing 767 outside Lisbon. The aircraft had run out of fuel because the engineer in charge of refuelling confused kilograms with pounds and loaded the aircraft less than half the required amount. For further examples of maintenance-induced common mode failures see (Ladkin 1983).

system as a redundant component who can assume a number of functions at his discretion dependent upon his diagnosis of the state of the system. Thus, manual control is secondary (Swenson et al. 1998: 194).

This is to say that – as far as the capsule engineers were concerned – Neil Armstrong was a redundant element in the Apollo 11 mission! People are sometimes uniquely useful in this role because they can respond creatively to unanticipated errors and interactions (Hollnagel et al. 2006: 4). NASA, for instance, could never have ‘MacGyverd’ the crippled *Apollo 13* capsule 400,000 km back to Earth without the astronauts on board to duct-tape things together and otherwise reconfigure the system on the fly.²⁰

There is also evidence, however, that people have unique limitations as redundant elements. The job of monitoring stable systems for very rare failures often requires an unrealistically high capacity for boredom, without offering much practice for whatever interventions any failure might require. Failures, meanwhile, can be overwhelming for people with less of the ‘right stuff’ than NASA’s early cadre of test-pilot astronauts (see e.g. Reason 1990: 180-82).

Interestingly, engineers will sometimes mitigate *human* reliability problems through redundancy. Commercial aircraft have two pilots, for instance, and the Navy make navigation work onboard their ships ‘robust’ by redundantly distributing knowledge among the navigation team and making workloads light enough to permit mutual monitoring and assistance (Hutchins 1995: 223). Indeed, many complex social organizations lean heavily on redundancy, including aircraft carriers, missile command and control centres, and air traffic control networks (see e.g. La Porte 1982). Airlines even make some maintenance practices redundant. As a prerequisite for being allowed to fly twin-engined (as opposed to three- or four-engined) aircraft over wide oceans (so-called ETOPS flights) airlines must ensure that different people perform the same maintenance job on different engines. (An analogue to design diversity, perhaps.)

Human errors, like mechanical failures, can also be ‘latent’. A nuclear engineer who misunderstands the procedures for shutting-down a reactor may never discover his or her error until it is too late. Reason (1990: 173-83) suggests that most significant human errors are latent, and that this is especially true of redundant systems where – ‘unlike driving in Paris’ – errors do not automatically and immediately reveal themselves. Vaughan (1996) makes an analogous claim, arguing that latent human errors or ‘deviances’ become ‘normalized’ over time.

The analogy between human reliability and its relationship to redundancy only goes so far, however. Combining people to make reliable systems comes with unique social and psychological problems. Extensive work has gone into studying the interactions between

²⁰ Although, of course, without the astronauts there would have been no need to bring the capsule home anyway.

pilots and co-pilots. This field, known as Cockpit Resource Management or CRM, was born of the realization that when two officers are charged with operating an aircraft certain authority relationships and communication protocols are more conducive to safe practice than others (see e.g. Wiener et al. 1993).

Scott Sagan, in a fascinating (2004: 939) article, outlines a series of caveats against the use of redundancy in social systems. As well as suggesting that redundant social systems suffer many of the same drawbacks as redundant mechanical systems (especially complexity), he argues they are prone to unique socio-psychological shortcomings. He argues, for instance, that where more than one person has responsibility for a single task there exists the risk of what he calls ‘social shirking’: the mutual belief of each person that the other will ‘take up any slack’ – a phenomenon well known to social psychologists (Sagan 2004: 939).

A closely linked dimension of redundancy, and one even harder to quantify, lies in what Sagan refers to as ‘overcompensation’ (2004: 941). The extra security that redundancy offers, for example, can lead people to act less cautiously. Perrow echoes this claim. He suggests that people may have a ‘risk homeostasis’ in which they become accustomed to a specific level of risk and compensate for lower risks in one area by taking greater risks in another (e.g. Peltzman 1975; Perrow 1984: 171; Wilde 1994).

Closely related to this idea of overcompensation is a tendency towards overconfidence in redundant systems by their designers, where the use of redundancy may reduce the perceived need for other kinds of safety, such as over-specification and/or rigorous testing.²¹ We saw above that airlines deemed quadruple engine failures to be so unlikely that they neglected to train the pilots for such an eventuality. As Diane Vaughan (1996) recounts in her important book about the *Challenger* disaster, NASA had long known that the Shuttle’s O-rings were problematic, yet they were confident in the integrity of the system because they knew that behind every O-ring was a redundant twin, a backup-up for the first. No failure had ever breached both O-rings, and NASA considered the chances of this happening to be marginal. In an important respect, therefore, the fateful ‘Challenger launch decision’ might be attributed to redundancy-induced overconfidence.

Engineers certainly realize that human reliability is a safety issue and a design problem, and many manufacturers invest much effort in making human-machine interfaces as intuitive, and error-tolerant (‘foolproof’ or even ‘drool-proof’ in more irreverent industries) as possible. Such efforts may come at their own cost in terms of

²¹ There is good evidence, for example, that pilots have come to rely on the automated flight control systems (in their role as a redundant pilot) to correct their errors and allow them to fly more recklessly. At Habsheim air show, in June 1988, for instance, it is alleged that the pilot of an A320 Airbus, relying on his computer to monitor the aircraft, attempted to take off so slowly that he was unable to clear a line of trees at the end of the runway. ‘He was so used to the system keeping the aircraft safe,’ writes one commentator ‘that he felt it could wave a magic wand and get him out of any problem.’ (Race 1990: 13-15) The pilot, it should be noted, contests this explanation, attesting instead to an altimeter malfunction.

overconfidence and morale: Soviet nuclear power plant operators reportedly considered the highly redundant – foolproof – design of American nuclear power plants all but an insult (Schmid 2004).

The real significance of human reliability, however, lies in its relationship to the processes of validation and accountability. Regulators struggle to quantify human reliability; it fits awkwardly into formal reliability measures, and the people who calculate the reliability of redundant systems usually separate human factors such as pilot error or prosaic externalities such as maintenance.²² Regulators do try to assess such factors, of course, but their findings in these areas usually black boxed have only an indirect bearing on their assessment of the technology itself. Most reliability measures, therefore, come with implicit caveats, such as ‘given proper maintenance’ or ‘if handled properly’. When constructing proofs, therefore, engineers can ‘dump’ epistemic ambiguities – points where the logic becomes intractable – by exploiting the interstices between different regulatory regimes. The FAA assesses the reliability of the flight crew separately from the reliability of the airframe, so if a pilot, rather than another machine, mediates between two elements, the proof of the ‘technology’ can escape an infinite regress by passing the epistemological ‘buck’ to the flight crew. This need not make the aeroplane safer, but – in a formal sense – it can help ‘prove’ its reliability. This has significant policy ramifications: courts and newspapers rarely hold airframe manufacturers accountable for pilot errors.²³

Conclusion

Redundancy is indispensable to a world where technological risks must be closely regulated and where ‘reliability’ is construed as a variable that can be defined, calculated and designed. The paradigm built around it provides an epistemic frame that makes technology-related regulatory regimes possible: underlying everything we know about the reliability of complex technological systems, invisibly implicated in the calculations that regulators and manufacturers pass down to the public and policy makers.

It can only perform this function if regulators disregard its epistemic ambiguities, however. Genuinely accurate measures of the reliability that redundancy offers require the measure of many immeasurable things, such as the degree of ‘independence’, ‘isolation’ and ‘similarity’ that separate elements enjoy (which, again, would be subject

²² Simon Cole observes how courts would routinely fail to grasp that the reliability of the people who collect and match fingerprints was likely to be much lower than the theoretical and ‘abstract’ reliability of the technology itself. ‘Knowing how often the examiners made errors’, he writes, ‘was just as important as the philosophical proposition that no two people had identical fingerprint patterns’ (Cole 2001: 210).

²³ The idea of ‘epistemic dumping’ is slightly tangential to the central thesis of this paper, which is concerned with redundancy as a design paradigm. Nevertheless, it suggests a mechanism through which the epistemology of audit practices shapes systems and institutions, and is, I believe, worth pursuing further. It is intended as the subject of a future publication.

to a similar regress, *ad infinitum*). Aviation manufacturers could never demonstrate the levels of reliability required of their products unless regulators were willing to overlook the uncertainties inherent in the frame it provides (and regulators would be unable to demand the kinds of proof that the public and policy makers request).²⁴

The ambiguities intrinsic to redundancy do not make it an ineffective engineering tool, and the dilemmas of quantifying its effects are no different from those inherent in any attempt to quantify a complex and ‘unruly’ world. We should be mindful of its limitations, however, and remember to ‘think twice about redundancy,’ as Sagan (2004) adroitly puts it. It is important we recognize that redundancy comes with costs as well as benefits and avoid the misconception that necessarily conflates it with reliability.

Understanding redundancy as a deeply entrenched design paradigm helps illuminate the complex epistemic relationship between risk and regulation, and adds weight to Porter’s (1995) argument that an over-emphasis on ‘hard’ numbers may distort the information we use to make important choices about modernity. The fact that aeroplane manufacturers operate under a regulatory regime that demands quantitative proof of aeroplane reliability, for instance, means that redundancy’s utility for quantitatively demonstrating such proof almost forces designers to use it, even where it might be *more* rather than *less* risky to do so.

More significantly, the constrictive lens of the redundancy paradigm fosters a misleading understanding of the technologies we depend on. Aircraft in operation generally prove to be as safe as regulators predict, but, despite appearances to the contrary, regulatory surety in new aircraft designs owes relatively little to quantitative reliability calculations. In practice, FAA reliability assessments lean much more heavily on ‘legacy’ data than on the *a priori* calculations that invoke redundancy. In other words, the regulators have extensive data about the performance of previous aircraft and are careful to ensure the safety of new designs by requiring that they closely resemble their predecessors (see Downer 2007). In this instance, therefore, reliability calculations are largely *performative*: they satisfy a popular misconception that technologies should be objectively and quantitatively ‘knowable’ (see Wynne 1988). In doing so, however, they perpetuate and reify that misconception. Few complex technologies have the long and painful legacy of civil aircraft, and so their regulators are forced to rely on redundancy calculations much more directly.

In technological domains where innovation is less incremental and experience is shallower, therefore, the paradigms become commensurately more significant. Where regulators must assess the reliability of complex technologies in these circumstances, the numbers they generate inevitably reflect the distortions of the paradigms through which

²⁴ This is a problem shared between manufacturers and regulators. Regulators cannot simply disqualify every aircraft because of epistemic uncertainties.

they frame their calculations. This is unavoidable; but we, as users, citizens and policy makers should learn to treat such calculations with an appropriate degree of circumspection. As one engineer put it:

The key to a reliable design is understanding. Not materials, not fancy manufacturing processes, not computer controls, not nifty feedback loops. Just simple human understanding. [...] really, it does all get down to who is doing what to whom and how they are doing it.

References

- Acohido, B. (1996) 'Pittsburgh disaster adds to 737 doubts', *Seattle Times*, 29 October.
- Air Accidents Investigation Branch (AAIB) (1990) 'Report on the accident to Boeing 737-400 - G-OBME near Kegworth, Leicestershire on 8 January 1989.' *Aircraft Accident Report* No: 4/90 (EW/C1095). London: Department of Transport.
- Air Line Pilots Association International (ALPA) (1999) 'Comments on rules docket (AGC-200)'. Docket no. FAA-1998-4815 (52539).
- Beck, U. (1992) *Risk society: towards a new modernity*. London: Sage.
- Bijker, W. Hughes, T. and Pinch, T., (eds.) (1989). *The social construction of technological systems: new directions in the sociology and history of technology*. Cambridge, MA: MIT Press.
- Cole, S. (2001) *Suspect identities: a history of fingerprinting and criminal identification*. Cambridge, MA: Harvard University Press.
- Collins, H. (1985) *Changing order*. London: Sage.
- Downer, J. (2007) 'When the chick hits the fan: representativeness and reproducibility in technological tests', *Social Studies of Science* 31 (1): 7-26.
- Eckhard, D. and Lee, L. (1985) 'A theoretical basis of multiversion software subject to coincident errors', *IEEE Transactions on Software Engineering* 11: 1511-17.
- Flight Safety Foundation (FSF) (2004) *Controlled flight into terrain and approach-and-landing accident reduction*. < <http://www.flightsafety.org/cfit1.html> > 23 September.
- Galison, P. (2000) 'An accident of history' in P. Galison and A. Roland (eds.),

Atmospheric flight in the twentieth century. Boston: Kluwer, pp. 3-43.

Government Accountability Office (GAO) (1992) 'Aircraft certification: limited progress on developing international design standards.' Report to the Chairman, Subcommittee on Aviation, Committee on Public Works and Transportation, House of Representatives. Report no. 147597 August. Washington DC: GAO.

Gleick, J. (1996) 'A bug and a crash: sometimes a bug is more than a nuisance', *New York Times Magazine*, 1 December.

Haynes, A. (1991) 'The crash of United Flight 232'. Paper presented at NASA Ames Research Center, Dryden Flight Research Facility, Edwards, California, 24 May.

Hollnagel, E. (2006) 'Resilience – the challenge of the unstable' in E. Hollnagel, D. Woods, and N. Leveson, *Resilience engineering: concepts and precepts*. Aldershot: Ashgate, pp. 9-17.

Hopkins, A (1999) 'The limits of normal accident theory', *Safety Science* 32: 93-102.

Hughes, R. (1987) 'A new approach to common cause failure', *Reliability Engineering* 17: 2111-36.

Hutchins, E. (1995) *Cognition in the wild*. Cambridge, MA: MIT Press.

Jasanoff, S. (1995) *Science at the Bar: law, science, and technology in America*. Cambridge, MA: Harvard University Press.

Jones-Imhotep, E. (2000) 'Disciplining technology: electronic reliability, Cold-War military culture and the topside ionogram', *History and Technology* 17: 125-75.

Kaldor, M. (1981) *The baroque arsenal*. New York: Hill and Wang.

Krohn, W. and Weyer, J. (1994) 'Society as a laboratory: the social risks of experimental research', *Science & Public Policy* 3 (21): 173-83.

Ladkin, P. (1983) *The Eastern Airlines L1011 common mode engine failure*. Online: <http://www.nts.gov/ntsb/brief.asp?ev_id=20001212X19912&key=1> Accessed 5 May.

La Porte, T. (1982) 'On the design and management of nearly error-free organizational control systems' in D. Sills, V. Shelanski and C. Wolf (eds.), *Accident at Three Mile Island*. Boulder, CO: Westview.

Landau, M. (1969) 'Redundancy, rationality, and the problem of duplication and

overlap', *Public Administration Review* 38: 346-57.

Lewis, H. (1990) *Technological risk*. New York: Norton & Co.

Littlewood, B. Popov, P. and Strigini, L. (2002) 'Assessing the reliability of diverse fault-tolerant systems', *Safety Science* 40: 781-96.

Littlewood, B. and Strigini, L. (1993) 'Validation of ultra-high dependability for software-based systems', *Communications of the ACM* 36 (11): 69-80.

Littlewood, B.; Popov, P. & Strigini, L. (1999) 'A note on the reliability estimation of functionally diverse systems', *Reliability Engineering and System Safety* 66: 93-95.

Littlewood, B. (1996) 'The impact of diversity upon common cause failure', *Reliability Engineering and System Safety* 51: 101-13.

Littlewood, B. and Miller, D. (1989) 'Conceptual modelling of coincident failures in multi-version software', *IEEE Transactions on Software Engineering* 15 (12): 1596-1614.

Littlewood, B. & Wright, D. (1997) 'Some conservative stopping rules for the operational testing of safety-critical software', *IEEE Transactions on Software Engineering* 23 (11): 673-83.

Lloyd, E. and Tye, W. (1982) *Systematic safety: safety assessment of aircraft systems*. London: Civil Aviation Authority.

MacKenzie, D. (2001) *Mechanizing proof: computing, risk, and trust*. Cambridge, MA: MIT Press, pp. 228-29.

MacKenzie D. (1996) *Knowing machines: essays on technical change*. Cambridge, MA: MIT Press.

National Academy of Sciences (NAS) (1980) Committee on FAA Airworthiness Certification Procedures 'Improving Aircraft Safety', *FAA Certification of Commercial Passenger Aircraft Committee on FAA Airworthiness Certification Procedures Assembly of Engineering National Research Council*. Washington, DC: NAS.

Negroni, C. (2000) *Deadly departure: why the experts failed to prevent the TWA Flight 800 disaster and how it could happen again*. New York: Cliff Street Books.

Neumann, J. von (1956) 'Probabilistic logics and synthesis of reliable organisms from unreliable components', *Annals of Mathematics Studies* 34: 43-98.

- Peltzman, S. (1975) 'The effects of Automobile Safety Regulation', *Journal of Political Economy* 83 (4): 677-725.
- Perrow, C. (1984) *Normal accidents: living with high-risk technologies*. New York: Basic Books Inc.
- Petroski, H. (1994) *Design paradigms: case histories of error and judgment in engineering*. Cambridge: Cambridge University Press.
- Pinch, T. (1993) 'CTesting – one, two, three ... testing!': toward a sociology of testing', *Science, Technology, & Human Values*, 18 (1): 25-41.
- Pinch, T. and Bijker, W. (1984) 'The social construction of facts and artefacts: or how the sociology of science and the sociology of technology might benefit each other', *Social Studies of Science* 14: 339-441.
- Popov, P., Strigini, L., May, J., and Kuball, S. (2003) 'Estimating bounds on the reliability of diverse systems', *IEEE Transactions on Software Engineering* 29 (4): 345-59.
- Popov, P.; Strigini, L. and Littlewood, B. (2000) 'Choosing between fault-tolerance and increased V&V for improving reliability', *DISPO Technical Report*.
- Porter, T. (1995) *Trust in numbers: the pursuit of objectivity in scientific and public life*. Princeton NJ: Princeton University Press.
- Power, M. (2007) *Organized uncertainty: designing a world of risk management*. Oxford: University Press.
- Race, J. (1990) 'Computer-encouraged pilot error', *Computer Bulletin* (August): 13-15.
- Reason, J. (1990) *Human error*. Cambridge: Cambridge University Press.
- Rochlin, G.I., La Porte, T.R. and Roberts K.H. (1987) 'The self-designing high-reliability organization: aircraft carrier flight operations at sea', *Naval War College Review* 40 (4): 76-90.
- Rozell, N. (1996) 'The Boeing 777 does more with less', *Alaska Science Forum*, 23 May.
- Rushby, J. (1993) 'Formal methods and the certification of critical systems', *SRI Technical Report CSL-93-7* December. <<http://www.csl.sri.com/papers/csl-93-7/>>

Sagan, S. (2004) 'The problem of redundancy problem: why more nuclear security forces may produce less nuclear security, *Risk Analysis* 24 (4): 935-46.

Sagan, S. (1993) *The limits of safety*. Princeton NJ: Princeton University Press.

Schmid, S. (2004) 'Reliable cogs in the nuclear wheel: nuclear power plant operators in the Soviet Union'. Paper presented at the Society of the History of Technology (SHOT) conference, Amsterdam, 7-10 October.

Shinners, S. (1967) *Techniques of system engineering*. New York: McGraw-Hill.

Swenson, L., Grimwood, J. and Alexander, C. (1998) *This new ocean: a history of Project Mercury*. Washington DC: NASA History Office.

Taylor (1990) *Twin-engine transports: a look at the future*. Seattle, WA: Boeing Corporation Report.

Tootell, B. (1985) '*All four engines have failed*': the true and triumphant story of BA 009 and the "Jakarta incident" London: Andre Deutsch.

Vaughan, D. (1996) *The Challenger launch decision*. Chicago: University of Chicago Press.

Wiener, E., Kanki, B. and Helmreich R., eds. (1993) *Cockpit resource management*. San Diego, CA: Academic Press.

Wilde, G. (1994) *Target risk: dealing with the danger of death, disease and damage in everyday decisions*. London: PDE Publications.

Wynne, Brian (1988) 'Unruly technology: practical rules, impractical discourses and public understanding', *Social Studies of Science* 18: 147-67.