

NORMAL ACCIDENTS

Living with High-Risk
Technologies

*With a New Afterword and a
Postscript on the Y2K Problem*

CHARLES PERROW

Princeton University Press

CHAPTER 3

Complexity, Coupling, and Catastrophe

To make a systematic examination of the world of high-risk systems, and to address problems of reorganization of systems, risk analysis, and public participation, we need to carefully define our terms. Not everything untoward that happens should be called an accident; to exclude many minor failures, we need an exact definition of an accident. Our key term, *system accident* or *normal accident*, needs to be defined as precisely as possible, and distinguished from more commonplace accidents. We will define it with the aid of two concepts used loosely so far, which now require definition and illustration: *complexity* and *coupling*. We also need to explain the terms *catastrophe* and *victim* as they are used in our analysis. Then we will be in a position to chart the world of organized systems, predicting which will be prone to system accidents and which will not. This chapter, drawing upon numerous examples, will construct an apparatus that will carry us through the systems discussed in the rest of the book. It constitutes a theory of systems, of their potential for failure and recovery from failure. As such, it is, I believe, unique in the literature on accidents and the literature on organizations.

Complexity, Coupling, and Catastrophe

Perhaps the most original aspect of the analysis is that it focuses on the properties of systems themselves, rather than on the errors that owners, designers, and operators make in running them. Conventional explanations for accidents use notions such as operator error; faulty design or equipment; lack of attention to safety features; lack of operating experience; inadequately trained personnel; failure to use the most advanced technology; systems that are too big, underfinanced, or poorly run. We have already encountered ample evidence of these problems causing accidents. But something more basic and important contributes to the failure of systems. The conventional explanations only speak of problems that are more or less inevitable, widespread, and common to all systems, and thus do not account for variations in the failure rate of different kinds of systems.

What is needed is an explanation based upon system characteristics. In Chapter 6, on marine transport, you will see how technological fixes at best make no difference, and can even make the situation worse. Chapter 8, on the space program, reveals how the best talent and organizational resources, while they certainly help, cannot overcome system accident potentials. Several chapters show how system-related production pressures defeat safety improvements. The accidents covered all chapters challenge the ready explanation of "operator error." This is not to say that the concepts and definitions of system characteristics presented here will solve all analysis problems. They are preliminary; definitional problems remain. This is a first attempt at a "structural" analysis of risky systems, but I believe it goes substantially beyond conventional "item" analyses, and yields, as we shall see in the last chapter, long-run strategies for handling risks—rather than short-run tactics of posting safety notices or levying trivial fines.

Defining Accidents

What do we mean by an *accident*? At the minimum, an accident is an unintended and untoward event. If you are driving home and take a wrong turn and lengthen your trip, it is both unintended and unfortunate. But we generally mean something more serious than this by the term *accident*. You would not, upon arriving home late, announce, "I had an accident on the way home." If you did, you would be worriedly asked, "Was anyone hurt?" or, "Was the car damaged?"

An accident, then, involves some damage to people, objects, or to both. But uncertainties remain. Suppose you scratched the paint of the car on a post while leaving a parking lot. You probably would not call it an accident (even though it was accidental) if it didn't interrupt the trip or impair the functioning of either the car or the post. We might, then, say that the damage to objects or people must be sufficient to disrupt the ongoing "task" or future tasks that will be demanded of the objects or people.

This introduces a further complication: tasks analysis. What task is involved? That depends on what we want to call the system. If I had planned to take the car to a rally the next day and show it off to other automobile buffs, I might very well call the scratch in the parking lot an accident. The system, from my point of view, involves my going to a rally to meet people, impress people, and show off my car, and it is interrupted. If the post was bent over so much that cars trying to enter the lot were endangered by what I had done, that too might qualify the event as an accident. Drivers would notify the maintenance people, "There has been an accident, the post is knocked over so far that I can hardly get into the lot." It is part of the parking system to have a post separating lanes; it is part of my recreational system to show off a beautifully maintained and polished car. The event has disturbed these systems.

But there are also degrees of disturbance to systems. The rally will not be disturbed in any perceptible way if I show up with a scratch on my '57 De Soto, or if, ashamed of the state of my car, I do not go at all. But my system might be greatly disturbed if I stayed home rather than meeting or impressing people. The degree of disturbance, then, is related to what we define as the system. If the rally is the system under analysis, there is no accident. If that part of my life concerned with custom cars is the system, there is an accident. A steam generator tube failure in a nuclear plant can hardly be anything other than an accident for that plant and for that utility. Yet it may or may not have an appreciable effect upon the nuclear power "system" in the United States.

So far we have defined an accident as unintended damage to people or objects that affects the functioning of the system we choose to analyze. But we can also affect system functioning by "damaging" symbols, communication patterns, legitimacy, or a number of factors that are not, strictly speaking, people or objects. This observation will become important if we are to consider such organizations as universities.

An accident, then, involves damage to a defined system that disrupts the ongoing or future output of that system. But not all such disruptions

should be classified as accidents; the damage must be reasonably substantial. In a nuclear plant, the momentary failure of the off-site power supply (the electricity coming into the plant to run its machinery) usually means that the reactor must be shut down. Similarly, a valve failure may result in a shutdown. But we will not call these *accidents*, even though they were not desired and have unfortunate or untoward consequences—purchasing replacement power alone may cost a few hundred thousand dollars before the reactor is restarted and the plant placed back on-line. We will call them *incidents*. Though the reactor is shut down, it does not sustain any damage. We need, then, criteria, inevitably somewhat arbitrary and rough, to distinguish between "minor" events such as these and accidents. Furthermore, we need a scheme that can apply equally to a steam generator as a system, or to a plant, or to the nuclear power industry.

I propose that the system be divided into four levels. Disruption to the third or fourth levels will be called *accidents* and disruption to the first and second will be called *incidents*. It is a scheme that need not be applied carefully in all parts of this book; sometimes we will ignore it and speak quite loosely of accidents when the meaning is obvious. But it is important for selecting accidents for analysis, for understanding safety devices, and for risk analysis; at times it will be crucial.

Consider a nuclear plant as the system. A *part* will be the first level—say a valve. This is the smallest component of the system that is likely to be identified in analyzing an accident. A functionally related collection of parts, as, for example, those that make up the steam generator, will be called a *unit*, the second level. An array of units, such as the steam generator and the water return system that includes the condensate polishers and associated motors, pumps, and piping, will make up a *subsystem*, in this case the secondary cooling system. This is the third level. A nuclear plant has around two dozen subsystems under this rough scheme. They all come together in the fourth level, the nuclear plant or system. Beyond this is the environment.

With this scheme we reserve the term *accident* for serious matters, that is, those affecting the third or fourth levels; we use the term *incident* for disruptions at the first or second level. The transition between incidents and accidents is the nexus where most of the engineered safety features (ESFs) come into play—the redundant components that may be activated; the emergency shut-offs; the emergency suppressors, such as core spray; or emergency supplies, such as emergency feedwater pumps. The scheme has its ambiguities, since one could argue interminably over the dividing line between part, unit, and subsystem, but it is flexible and

adequate for our purposes. It must be flexible because at times we might want to consider, for example, the Apollo rocket and modules as a system; at other times we might want to consider all moon shots as a system. What is an accident in the first might be a mere incident in the second.

We are now ready for a formal definition. An *accident* is a failure in a subsystem, or the system as a whole, that damages more than one unit and in doing so disrupts the ongoing or future output of the system. An *incident* involves damage that is limited to parts or a unit, whether the failure disrupts the system or not. By disrupt we mean the output ceases or decreases to the extent that prompt repairs will be required. Since we have drawn a dividing line between the unit and the subsystem, and since many of the ESFs are clustered around that dividing line, it will often mean that an ESF will be one of the components that fails.

Victims

Note that, in contrast to common usage, we have not explicitly made damage to humans a part of our definition. This is because we are primarily interested in systems and the way they work. Humans are part of all the systems considered in this book, and a group of humans (the flight crew on an airliner) or a single human (the astronaut in a space capsule) may constitute a subsystem. Damage to them will mean that we have an accident.

But it is important for analysis to treat humans in most systems as mere parts. We kill about 5,000 people a year outright in U.S. industry. The vast majority of these "accidents," however, are only "incidents" in our scheme, for no subsystem or system damage is entailed. Only a "part" has been destroyed.

I am aware that this sounds like a quite heartless scheme, but while the ultimate concern of this book is reducing actual and potential damage to humans, I believe it is the character of systems that cause that damage, and thus we need a definition of accidents that focuses upon system characteristics. The failure of what we have defined as parts and units plays a vital role in subsystem and system failures, but if the analysis is limited to these we will lose our focus upon the kinds of systems that business and government leaders decide should be built. Our analysis,

then, must be of systems, and falling off a ladder becomes a mere incident.

More important, we are concerned about those systems that have catastrophic potential—can cause damage to a great many humans. Our concern here is not the blue-collar worker who backs into a grinding machine, or the scientist who accidentally creates the proper criticality conditions for the plutonium he is experimenting with and is fatally irradiated. More or less mundane precautions and training will reduce the frequency of such accidents; fairly minimal attention to mine safety greatly reduced the number of coal mining deaths in a few years. Our concerns are more cosmic. To bring them into focus, we will turn to a victim classification scheme that will be used throughout this book.

Most of the work concerned with safety and accidents deals, rightly enough, with what I call first-party victims, and to some extent second-party victims. But in this book we are concerned with third- and fourth-party victims. Briefly, first-party victims are the operators; second-party victims are nonoperating personnel or system users such as passengers on a ship; third-party victims are innocent bystanders; fourth-party victims are fetuses and future generations. Generally, as we move from operators to future generations, the number of persons involved rises geometrically, risky activities are less well compensated, and the risks taken are increasingly unknown ones. We will take a close look at each of the four classes.

First-party victims are the operators of the system. In this book operators will include not only those actually running the system (nuclear plant operators, pilots, ship officers) but others in attendance on regular shifts, such as first-level supervisors, maintenance personnel, low-level engineering personnel, and laborers and assisting personnel. Most industrial accidents involve operators; in fact, most involve only one operator. Most industrial accidents are attributed to "operator error" or "human error" by those who study and seek to prevent such accidents. There is a growing recognition, however, that this is a great oversimplification; worse, it involves blaming the victim. It also suggests an unwitting—or perhaps conscious—class bias; many jobs, for example, require that the operator ignore safety precautions if she is to produce enough to keep her job, but when she is killed or maimed, it is considered her fault. Some operators are compensated for the risks they run by receiving higher pay than the skill level would suggest, but in industry as a whole there is no clear relationship between risk and pay.¹ There is no impersonal, fair market that rewards those that risk their lives with higher wages. Indeed,

"jumpers" or the "glow boys" in the nuclear industry, temporary help who dash into a radioactive area to make repairs, will be hired for two or three weeks' work, at only six dollars an hour, even though they can receive high doses of radiation in the few minutes they are near the core.² There is no evidence of compensation for the long-term effects of toxic chemicals or contaminated atmospheres. Textile workers are not compensated for brown lung disease, nor are chemical plant workers compensated for cancer showing up ten or twenty years after exposure.

Second-party victims are those associated with the system as suppliers or users, but without influence over it. They are not innocent bystanders (third-party victims), because they are aware (or could be informed) of their exposure, even though such exposure may not be entirely voluntary. The largest class of second-party victims are passengers on ships, trains, airlines, cars, and busses. To some extent they "choose" to participate in the system, and thus elect to share at least some of the risk. If I accept a ride home from a party with an inebriated driver, I accept the risk. The difference between first- and second-party victims with regard to voluntary exposure may be slight when we compare an unemployed person who has no choice but to accept risky work, and an employee who must travel extensively to keep his job. Neither has realistic options to participating in the system. But there are less ambiguous examples. Comparing second-party victims to innocent bystanders illustrates the slightly voluntary nature of the former: we feel differently about airline passengers killed in a crash than we do about innocent bystanders on the ground killed by the airliner's crash. The former accepted the risk of flying, the latter did not.

There are other types of second-party victims. For example, consider the office personnel hurt by a refinery explosion or the truck driver delivering goods for his company who just happens to be on the site when the explosion occurs. These are the voluntary actions of people who elect to participate in a system but have no influence over its operation. Nevertheless, without these second party or system-associated participants there would be no system. The refinery could not run. Like passengers, they are participants in the system and aware of some risks.

Third-party victims, innocent bystanders, have no such involvement in the system. I have heard some nuclear proponents argue, in a conference on establishing safety goals for nuclear power plants, that people can choose not to live near a nuclear plant (or choose to avoid going to events at a stadium that is in the flight path of an O'Hare Airport runway). But I think these arguments can be dismissed. Nuclear plants exist near all densely populated areas in the United States; there is simply no

practical means to avoid being within 50 miles of one. Even if one could, a severe core meltdown with breach of containment under the proper weather conditions could contaminate areas as large as the whole Northeast. The plutonium that might have crashed on Madagascar during the *Apollo 13* emergency reentry of the atmosphere would have contaminated innocent bystanders in a distant part of the world that one would have thought safe from such high-technology disasters. People are aware of risks associated with flying; but I suspect that most people living downstream from large dams do not realize the dam is near enough to endanger their lives in the case of a failure.

Fourth-party victims are, for the most part, victims of radiation and toxic chemicals. They are fetuses being carried at the time of exposure; the would-be children that damaged parents will not be able to conceive; stillborn or deformed children conceived after exposure; and all those people who will be contaminated in the future by residual substances, including those substances that will become concentrated as they move up the food chain. Note that we do not include run-of-the-mill pollution here. Nuclear plants may regularly give off radiation in doses large enough to affect fetuses or future reproductive capacity, as a few scientists argue (though most deny); but these consequences are not the result of accidents. (Although one might argue that such radiation is a design failure that should, but doesn't, interrupt the output.) Nor are the pollutants from industrial plants of concern here, unless they are released in an accident. This routine and mostly conscious contamination of the planet probably has more serious long-term consequences than any accident we shall consider in this book, other than military accidents, but it is beyond the scope of our concern.

The importance of fourth-party victims in risk analysis is growing with the increasing concern and sophistication about the long-term consequences of some systems. Yet recognition is slow. A draft version of the NRC's "Safety Goals for Nuclear Power Plants" indicated that while the NRC was aware of "inter-generational" risks from a nuclear power plant accident through genetic effects or long-term contamination, "we cannot suggest a good way to handle the issue in a safety-goal context." So they ignored it.³ If the risks of an accident are kept low enough, they said, there will be no problem with ignoring inter-generational effects. This conclusion answers the question about consequences of accidents by saying they will be trivial because there will be so few accidents.

Fourth-party victims potentially constitute the most serious class of victims. Chemical or radioactive contamination of land areas could have far-reaching effects upon the health of future generations. Genetic defects

harm future generations in other ways, including adding the burden of lifetime care and treatment of victims. Future generations carry the burden; the present generation reaps whatever rewards there may be from the activity.

These issues are comparatively new—less than one generation old. Some influential scientists and academics, at the workshop that tried to formulate safety goals for the NRC, argued that the present generation is more important than future ones—for we need nuclear power to prevent economic and political crises—and who knows, there may be a technological fix that would mitigate the burden for the future generations of an accident in the present one. Thus at least some of the “experts” do not see the issue as a clear-cut one of our responsibility to future generations. Other experts, however, believe we have a responsibility to pass on to our offspring a world that is at least no more contaminated and degraded than the one we inherited.

Accident Definitions

We now have a definition of an accident, distinguishing it from incidents on the basis of levels in the system. The following list presents these and adds the definition of the two types of accidents: component failure accidents and system accidents, which we will now take up.

Systems are divided into four levels of increasing aggregation: units, parts, subsystems, and system.

Incidents involved damage to or failures of parts or a unit only, even though the failure may stop the output of the system or affect it to the extent that it must be stopped.

Accidents involve damage to subsystems or the system as a whole, stopping the intended output or affecting it to the extent that it must be halted promptly.

Component failure accidents involve one or more component failures (part, unit, or subsystem) that are linked in an anticipated sequence.

System accidents involve the unanticipated interaction of multiple failures.

Component failure accidents and system accidents are distinguished on the basis of whether any interaction of two or more failures is anticipated, expected, or comprehensible to the persons who designed the sys-

tem, and those who are adequately trained to operate it. A system accident, in our definition, must have multiple failures, and they are likely to be in reasonably independent units or subsystems. But system accidents, as with all accidents, start with a component failure, most commonly the failure of a part, say a valve or an operator error. It is not the *source* of the accident that distinguishes the two types, since both start with component failures; it is the presence or not of multiple failures that interact in unanticipated ways.

The vast majority of component failure accidents involve a series of failures. If a valve fails there is good chance that the pump will overheat and fail, and if this happens, the boiler is quite likely to overheat because the coolant is insufficient. Designers know this, and so do operators—though they may not be able to prevent it, or intervene in the series of failures. There will be some accidents, however, where the initial failure is so drastic that it is not worth tracing out the subsequent sequence, if there is one. If the wing comes off an airplane in flight, or an earthquake shatters a dam, we hardly need further analysis. These might be called “final accidents”; there is no possible intervention by operators, and no point in detailing sequences. We will not deal with any examples of these in this book; for one thing they are fairly rare, but more important, analysis of the system itself does little to contribute to understanding these accidents.

Incidents are overwhelmingly the most common untoward system events. Accidents are far less frequent. Among accidents, component failure accidents are far more frequent than system accidents. I have no reliable way to estimate these frequencies. For the systems analyzed in this book the richest body of data comes from the safety-related failures that nuclear plants in the United States are required to report. Roughly 3,000 Licensee Event Reports are filed each year by the 70 or so plants. Based upon the literature discussing these reports, I estimate 300 of the 3,000 events might be called accidents; 15 to 30 of these might be system accidents. So far, at least as far as we know, all the system and component failure accidents in nuclear plants have had only first-party victims, and very few of these.

Complex and Linear Interactions

What kinds of systems are most prone to system accidents? In the last chapter we kept mentioning two concepts: interactiveness, which could confuse operators, and tight coupling, which could prevent speedy recovery from an incident. With a more precise definition of these two terms, we can classify systems and be forewarned about those that are most prone to system accidents. We will take up interactions first.

The notion of baffling interactions is increasingly familiar to all of us. It characterizes our social and political world as well as our technological and industrial world. As systems grow in size and in the number of diverse functions they serve, and are built to function in ever more hostile environments, increasing their ties to other systems, they experience more and more incomprehensible or unexpected interactions. They become more vulnerable to unavoidable system accidents.

Interactiveness per se, though, is not a useful concept. Almost any organization of any size, whether public or private, will have many parts that interact once we look closely at them. The existence of many parts is no great trouble for either system designers or system operators if their interactions are expected and obvious. If a part or a unit fails in an assembly line, it is quite clear what will happen to the parts and units "downstream" of the failure, and we know that the products "upstream" will start piling up fast. We can shut down the line and fix it, or bypass that work station and assemble the missing parts later on, or temporarily shunt the product into a storage bin. These are "linear" interactions: production is carried out through a series or sequence of steps laid out in a line. It doesn't matter much whether there are 1,000 or 1,000,000 parts in the line. It is easy to spot the failure and we know what its effect will be on the adjacent stations. There will be product accumulating upstream and incomplete product going out downstream of the failure point. Most of our planned life is organized that way.

But what if parts, or units, or subsystems (that is, components) serve multiple functions? For example, a heater might both heat the gas in tank A and also be used as a heat exchanger to absorb excess heat from a chemical reactor. If the heater fails, tank A will be too cool for the recombination of gas molecules expected, and at the same time, the chemical reactor will overheat as the excess heat fails to be absorbed. This is a good design for a heater, because it saves energy. But the interactions are no longer linear. The heater has what engineers call a "common-mode"

function—it services two other components, and if it fails, both of those "modes" (heating the tank, cooling the reactor) fail. This begins to get more complex.

This source of complexity was slow to gain recognition in the nuclear power field. The first analytical model to consider common mode failures appeared only in 1967, according to a perceptive and disturbing article by one of the editors of *Nuclear Safety*, E. W. Hagen.⁴

The monumental reactor safety study issued in 1975, WASH-1400 or the Rasmussen Report, was criticized by a subsequent NRC study for giving insufficient and overly simplified attention to this problem. Hagen concludes that potential common-mode failures are "the result of adding complexity to system designs." Ironically, in many cases, the complexity is added to reduce common-mode failures. The addition of redundant components has been the main line of defense, but, as Hagen illustrates, also the main source of the failures. "To date, all proposed 'fixes' are for more of the same—more components and more complexity in system design."⁵ The Rasmussen safety study relied upon a "PRA" (probabilistic risk analysis), finding that core melts and the like were virtually impossible. Hagen notes (p. 191) that PRAs, "using established techniques of reliability and statistical analysis," constitute the main source of public reassurance about such potential dangers. But, he notes, this involves a "narrow definition" of common-mode failures, and the analysts "are busily working in an area which has not been shown to be the main problem." The main problem is complexity itself, Hagen argues. With that we can agree.

Common-mode failures are just one indication of interactiveness in systems. Proximity and indirect information sources are two others. For a graphic illustration of complexity in the form of unanticipated interactions from these sources, let us go to a different system, marine transport. A tanker, the *Dauntless Colocotronis*, traveling up the Mississippi River near New Orleans, grazed the top of a submerged wreck. The wreck had been improperly located on some of the charts. Furthermore, it was closer to the surface than the charts indicated because its depth had been determined when the channel was deep, and the listing for mariners had not been corrected for those seasons when the river would be lower. The wreck, only a foot or so too high for the tanker, sliced a wound in the bottom of the tanker, and the oil began to seep out. Unfortunately, the gash occurred just at the point where the tank adjoined the pump room. Some of the oil seeped into the pump room. At first it was probably slow-moving, but the heat in the room made it less viscous,

allowing it to flow more rapidly, drawing more oil into the pump room. When enough accumulated it reached a packing "gland" around a shaft near the floor that penetrated into the engine room next to it. The oil now leaked into the engine room. In the hot engine room, evaporation was rapid, creating an explosive gas. There is always a spark in an engine room, from motors and even engine parts striking one another. (In fact, nylon rope can create sparks sufficient to cause explosions in tanker holds.) When enough gas was produced through evaporation, it ignited, causing an explosion and fire.

In this example, an unanticipated connection between two independent, unrelated subsystems that happened to be in close proximity caused an interaction that was certainly not a planned, expected, linear one. The operators of the system had no way of knowing that the very slight jar to the ship made a gash that would supply flammable or explosive substances to the pump and engine rooms; nor of knowing, until several minutes had passed, that there was a fire in the engine room; or of the extensive amount of oil involved. They tried to put it out with water hoses, perhaps not realizing it was an oil fire, but the water merely spread the oil and broke it up into finer, more flammable particles.

This accident went on for several hours, because the crew made various mistakes (e.g., a key fire door had been tied open and was not closed by an escaping crew member, allowing the fire to spread), and they did not have a clear notion of the location of the many rooms and closets and passageways in that part of the ship. Late in the accident a trained fire crew boarded the ship equipped with protective devices and proper extinguishing equipment. But when they opened one door there was a series of explosions. They immediately closed the door and made no further attempt to put out the fire in that part of the ship, believing that the oil was producing explosive gases. Subsequently, however, it was determined that there were no explosive gases involved by this time; instead, three small empty gas tanks stored just inside the door, before being exchanged for full ones, had exploded when the heat expanded the residual amounts of freon, oxygen, and acetylene in them.⁶

Thus even in the recovery stage of the accident a nonlinear interaction intervened to mislead the recovery attempt. It was a sensible place to store the tanks; who could imagine that there might be a fire, that firefighters would not be certain that there were not explosive mixtures behind the closed doors, and that these tanks would go off just as the firemen were entering the passageway? A requirement that empty tanks be stored elsewhere would hardly make the ship safer; who could know where the next fire might be?

Recall the illustration that opened this book, linking three "subsystems": breakfast, getting to the appointment, and the job interview. In the world we plan out and think through, this seems like a very linear, straightforward problem—get some food and coffee, get the car, drive to the interview. One would expect the car keys to be linked to using the car, but one would not expect the failure of the coffeepot to be linked to using the car. One would also not expect that even if the car failed, the alternative of a taxi would be linked to a contract dispute, and the neighbor's car would be unavailable just that day. These represent interactions that were not in our original design of our world, and interactions that we as "operators" could not anticipate or reasonably guard against. What distinguishes these interactions is that they were not designed into the system by anybody; no one intended them to be linked. They baffle us because we acted in terms of our own designs of a world that we expected to exist—but the world was different.

I will refer to these kinds of interactions as *complex interactions*, suggesting that there are branching paths, feedback loops, jumps from one linear sequence to another because of proximity and certain other features we will explore shortly. The connections are not only adjacent, serial ones, but can multiply as other parts or units or subsystems are reached.

The much more common interactions, the kind we intuitively try to construct because of their simplicity and comprehensibility, I will call *linear interactions*. Linear interactions overwhelmingly predominate in all systems. But even the most linear of systems will have at least one source of complex interactions, the environment, since it impinges upon many parts or units in the system. The environment alone can constitute a source of failure that is common for many components—a common mode accident. But even the most complex systems of any size will be primarily made up of linear, planned, visible interactions.

Based upon a general knowledge of various systems, considerable experience with reading accident reports, and some site visits, I suggest that only 1 percent of all possible parts or units in a linear system are capable of producing "complex" interactions, while about 10 percent of those in a complex system will be capable of doing so. But that 10 percent represents more than a tenfold increase in the potential for system accidents. The potential interactions produced by, say, four parts or units that are interrelated in a complex, rather than linear way, is twelve. (There are twelve possible paths between the four parts.) Suppose this exists in a system where there are 400 parts or units. If 10 percent of the units had such characteristics, rather than only 1 percent, there would be forty such

parts or units. The potential complex interactions of each one of these with the remaining 399 would be in the millions, since the possible paths increase exponentially. Of course, in a large system some parts are so remote from each other that the chances of their interacting in unexpected ways can be disregarded; but many will not be so remote. Thus, increasing the proportion of possible complex interactions from 1 percent to 10 percent will have an enormous impact upon the potential for system accidents—given the fact that, since nothing is perfect, component failures are inevitable.

The possibility of many unintended interactions is recognized by designers, so they introduce buffers and other safety devices to prevent some kinds of interactions. Imagine a chemical plant where gas is expected to flow from tank A to tank B, because the pressure is kept higher in tank A than in tank B. Various things go into tank A besides the basic gas—reagents, purifiers, “inerting” gases, and so on. Then the gas flows to B, where additional substances are piped in to alter that mixture even further. This is quite straightforward and linear.

But there is a danger that it might become nonlinear. A feedback loop could occur if pressure declined in tank A because of a failure there or elsewhere, and the now-modified gas in tank B flowed back into A. This could create problems; it might even be dangerous. The engineers are aware of this, so they design a butterfly valve to be placed between the two tanks to prevent reverse flow. An engineered safety device (ESD) is installed to keep the system as linear as possible (in this case, flowing in one direction). But valves can fail, especially rarely used ones. If the pressure in tank A were to fall and the butterfly valve to fail, the feedback loop would open again, and the operators might not expect the interaction. Actually, this example is so simple that most operators would take the unplanned interaction into account as a possibility.

For a more complicated example involving quite remote and independent units, recall the case in the last chapter of the demineralized water needed to clean up the containment room in the nuclear plant. A faucet (valve) was opened to see if the water had been sent to the men, and a complex series of backflows and pressure adjustments occurred that allowed radioactive materials to enter containment. There was nothing very linear about that, to the operators in containment, the operators in the control room, or the engineers who designed the system. Nor did anyone make an error or even a mistake in design. Subsequently they could figure out what had happened and introduce measures to prevent a recurrence, but that meant altering an unexpected and unplanned inter-

action between two normally independent subsystems, core maintenance and steam generation. As I suggested, there are probably a number of other complex interactions ready to be revealed by a casual event in that very complicated piping system.

Linear interactions are overwhelmingly those of an “expected production sequence”—this is the way the system is designed to run, and any one operating it will know that. Complex interactions will generally be those not intended in the design. No designer of the *Dauntless Colocotronis* said, “Let’s put the number 5 tank next to the pump room so they can interact.” There is just no way to isolate all tanks in a tanker from a room that generates sparks. Or the nonlinear interactions may be intended but rarely activated, and thus operators or designers forget about them. It could have been foreseen by a designer that demineralized water might sometimes be needed in containment, so she could have made it possible to line up various valves to provide the water. But if it is rarely used or is usually lined up before a crew enters containment, the faucet creates no problems, it is not an expected production sequence but an infrequently used system possibility (in this case, for maintenance, not production). Thus, complex interactions may be unintended ones, or intended but unfamiliar ones.

While linear interactions occur overwhelmingly in an anticipated production sequence, there is another kind of interaction that does not occur in a production sequence but is nevertheless obvious, and thus can be defended against. This is a visible interaction, even though it is outside of the normal sequence. If the operator of a crane sees that the part of the crane that holds up the load (cable, ratchet, motor, hook) has failed and the load is about to drop on a boiler on the floor, he knows what this interaction will produce. There is nothing mysterious about the connection between these events, though they are certainly not in any expected production sequence. Knowing there is a remote possibility of a large load falling, the operator usually tries to avoid passing it over the boiler, but he cannot always do so. Similarly, the designer may have considered covering the boiler or moving it, but estimated that the problems involved were too great, given the remote possibility of a large load falling at just that point.

To summarize our work so far: *Linear interactions* are those interactions of one component in the DEPOSE system (Design, Equipment, Procedures, Operators, Supplies and materials, and Environment) with one or more components that precede or follow it immediately in the sequence of production. *Complex interactions* are those in which one

component can interact with one or more other components outside of the normal production sequence, either by design or not by design.

To elaborate on the properties of these interactions as they affect the operators:

Linear interactions are those in expected and familiar production or maintenance sequence, and those that are quite visible even if unplanned.

Complex interactions are those of unfamiliar sequences, or unplanned and unexpected sequences, and either not visible or not immediately comprehensible.

A note on the terms "complex" and "linear" is in order. It is difficult to find precise terms that are also brief; I have opted for brevity. "Complex" should read "interactions in an unexpected sequence"; "linear" should read "interactions in an expected sequence." One problem with the terms *complex* and *linear* is that the opposite of complex is "simple," and the opposite of linear is "nonlinear." Linear interactions are "simple" in the sense that they are easily comprehensible, but "simple" implies unsophisticated processes and technologies, or systems with few parts and units or routine and uncomplicated operation and maintenance. But the production of, say, pharmaceuticals or F-16 fighter planes is anything but simple, though it is linear. On the other hand, "nonlinear" does not readily convey the notion of possible incomprehensibility, while the term *complex* does. I will occasionally slip in such phrases as "complexly interactive" and "linear sequences" to remind the reader that neither of the chosen terms is really satisfactory.

Another warning is in order. Since linear interactions predominate in all systems, and even the most linear systems can occasionally have complex interactions, systems must be characterized in terms of the degree of either quality. It is not a matter of dichotomies. Furthermore, systems are not linear or complex, strictly speaking, only their interactions are. Even here we must recall that linear systems have very few complex interactions, while complex ones have more than linear ones, but complex interactions are still few in number.

Finally, the reader should not equate the notion of a linear system with the physical layout of the plant or production process. It does not necessarily imply an assembly line, although these production systems tend to be linear. Nor does the designation "complex system" necessarily imply highly sophisticated technology, numerous components, or many stages of production. We will characterize universities as complex systems, but

as most large-scale organizations go, they have few of the above characteristics.

In order to understand systems, we need to go beyond the distinction between the two types of interactions. By examining, in the next section, how systems cope with hidden interactions, we will explore the basic attributes of complex systems, and be able to characterize them more systematically.

Coping with Hidden Interactions

Linear systems also have interactions that are not visible, of course, but they occur within well-defined and segregated segments of the production or maintenance sequence. Controls, such as dials, warning lights, audible alarms, and switches read the presence of these interactions and inform the operators and allow them to intervene. In systems with a fair degree of complex interactions, however, well-defined and segregated segments do not necessarily exist. Instead, jiggling unit D may well affect not only the next unit, E, but A and H also. This increases the number of controls that must be installed and monitored. The control panel of a nuclear plant is considerably denser and larger than that of an oil or coal plant, in part because so many components are linked in branching paths and feedback loops.

Attempts are continually made to reduce the number of controls by automating the subsidiary interactions and leaving only the main parameters for the operators to worry about. But this decreases the system's flexibility; the operator loses the ability to correct a minor failure in a part rather than shutting down a whole unit or subsystem. The operator cannot exit from the high-level, summary controls to the low-level specific one required to deal with a single part. Larry Hirschhorn provides an instructive example in a perceptive article on the limits of "cybernetic" (self-correcting) systems, which I shall elaborate on considerably.⁷

Consider a piston engine and a fuel pump. If particles get into a cylinder, perhaps during a maintenance procedure, or from wear from some other part, they can cause the piston to move more sluggishly. The result is that there is less output from this part of the engine. A control device monitors the output (but not the many causes of reduced output, such as foreign particles, piston wear, low octane fuel, over- or under-heating).

Noting the reduced power, the automatic control calls for more fuel from the fuel pump. This does compensate for the new source of friction, and the machine functions adequately, though less efficiently. But suppose that a surge in power is called for from the engine. The fuel pump is asked to increase its output. Soon its limit is reached. A typical cybernetic monitor will now decide that the fuel pump has failed; it will not know that the requirements expected of a (now degraded) piston/cylinder part cannot be met. An engineered safety device comes on, sending fuel from another source to the cylinder.

Any number of things might happen at this point. If the first pump is left running (since it has not, in fact, failed at all) there will be too much fuel supplied, which might burst some fuel carriers, stall the engine with fuel flooding in, overpressurize the fuel supply so that it tries to flow backwards (an eventuality that was not conceived by designers, so no protection is provided), or cause a runaway engine that rapidly overheats and explodes.

If, instead of the above, power to the fuel pump is automatically shut off because it is believed to have failed, the pump may be shut off with the valves in the open position, rather than the closed position they would be in if the pump had indeed failed. This could create back-flow or other problems since the other systems have not been shut off. Designers might have anticipated that if an undamaged pump were shut off, it would be because of a power failure that would affect other units of the system, and thus no back-flow problems would occur. In this case, however, an undamaged pump loses its power but keeps its valve setting, while other related parts or units do not lose their power and thus keep operating.

After the failure, the pump would probably be replaced, since the monitor indicated that this is what had failed. But the problem would remain, since it resided in the piston; only the automatic cybernetic control system thought the problem was the pump. This elaborated example is fabricated, but it illustrates the difficulty control systems have in reading the nature and source of failures. In a system without this added complexity, an operator presented with this sequence would probably cut back on the power demands, reasoning that either the engine or the fuel supply was malfunctioning, check the pump by varying engine speed, and conclude that there was sluggishness in the engine, which in fact is the correct conclusion.

One limitation of this example is that there is no pressing reason to have automatic controls in such a simple system (though if the motor exists in connection with many other units or subsystems, it may be

advisable). But in some systems automatic controls are necessary because there is simply not sufficient time for operator reaction. The nuclear reactors designed by Babcock and Wilcox are of this type. The reactor is extremely responsive, and this has important economic benefits. But some of its subsystems may be too responsive. With a loss of coolant, the operator may have only thirty to sixty seconds before the steam generator boils dry. Consider this accident at the Crystal River, Florida, plant in February 1980.

For an unknown reason a short circuit occurred in some of the controls in the control room. The utility said it could have been due to a bent connecting pin in the control panel, so sensitive are these devices, or the malfunction may have been caused by some maintenance work being done on an adjacent panel. The short circuit distorted some of the readings in the control system, in particular the important and very sensitive one of coolant temperature. The computer "thought" the coolant was growing too cold, so it speeded up the reaction in the core. (The Babcock and Wilcox reactor operates within a very efficient, but quite narrow temperature band.) The reactor overheated, the pressure in the core went up to the danger level, and then the reactor automatically shut down. The computer was apparently now puzzled, and correctly ordered the pressure relief valve (PORV) to open, but incorrectly ordered it to remain open until things settled down. This was an "error" on its part, because pressure dropped so quickly that it automatically caused high pressure injection to come on, and stay on, flooding the primary coolant loop—including the core, steam tubes, and pressurizer. A valve stuck and 43,000 gallons of radioactive water were dumped on the floor of the reactor building. Fortunately it was not worse; after several minutes an operator noticed the computer's error in keeping the relief valve open and closed the valves manually. Had the frequent injunction been followed that the computer knows best, and that the (dumb) operators should keep their hands off the system until its routines are carried out, the sump would have overflowed and we would indeed have had a wet reactor.⁸

Operators tend to resist the introduction of more general, high-level controls such as cathode ray tubes (CRTs), which produce a television screen display of the state of a number of units or subsystems, because they feel they cannot make selective interventions as easily. Only these high-level controls can be manipulated; the selective ones are more difficult to get to because it is assumed they will not be needed. They may be out of the way, or accessible only by a complicated series of steps that inactivate the more general controls. On the other hand, operators also

complain about the arrangement of less automated systems that allow selective intervention, because they are confronted with 15-foot banks of identically designed switches with tiny numbers above them. These are not even grouped in any way that reflects operation, but only in the way that reflects ease of installation. One of the most common "operator errors" in nuclear plants is, understandably, operating the wrong switch. Surely the choice need not be reduced discretion versus endless, identical displays.⁹

Complicating the dilemma of too few vs. too many controls is the problem of uncertain "default" status. Default status is the normal status of a control; for example, you must choose to change a switch to "on"; by default it is "off." The position of a relief valve is normally closed. The temperature reading on a drain pipe is normally cool. The position of most block valves is normally open. One might decide to have the default status represented by a green light on the control panel—go ahead, there is no danger. But one also might decide to have green represent open valves or circuits that are open to current flow—they are "on," and on is the default status. One can readily see the problem: if something is normally closed, this status cannot be represented by the color green, since green also means it is on, or open. Nor can red be reserved for "danger"; if the reactor is supposed to be on, you may want a red signal to warn you that it *is* on, even though it is not dangerous, but normal under these circumstances. But then what would green mean? (I learned the hard way that on my personal computer closing a "dip switch" means to open it!)

Some switches and valves have to be on at certain times and off at others, so there is no default position. Some systems use colored lights to handle this. If operating mode A is in effect, then switch 1 should be on, and an amber light next to the switch indicates this. If operating mode B is in effect, the switch should be off. If the switch is in the wrong position for the operating mode, a red or blinking light appears. If that part of the system is not in operation, a green light may be on, regardless of the position of the switch, or no light may be on. But modes and switches are not always clearly linked; there may well be a mode A and a mode AI, and the correct switch positions for the two would be different. Operators sometimes disconnect these sophisticated indicators (or much more commonly, ignore them) because of such complications.

These problems exist in all industrial and transportation systems, but they are greatly magnified in systems with many complex interactions. This is because interactions, caused by proximity, common mode connections, or unfamiliar or unintended feedback loops, require many

more probes of system conditions, and many more alterations of the conditions. Much more is simply invisible to the controller. The events go on inside vessels, or inside airplane wings, or in the space craft's service module, or inside computers. Complex systems tend to have elaborate control centers not because they make life easier for the operators, saving steps or time, nor because there is necessarily more machinery to control, but because components must interact in more than linear, sequential ways, and therefore may interact in unexpected ways.

In addition to there being many interactions to control, the information about the state of components or processes is more indirect and inferential in complex systems. The reactor core at Three Mile Island had no direct measure of coolant level. Given the high pressures and the flow of coolant in the core, such a measure would be very hard to provide, and would introduce more penetrations of the vessel, something to be avoided. So, operators had to estimate coolant level from indirect indicators. The presence of steam voids and hydrogen bubbles made even these unreliable.

To cite other examples, pilots in ships or airplanes can fix on the wrong star. Ships may find navigation beacons obscured by shore lights, distorted by refractions, or simply out. Routine fluctuations in pressure and temperature in chemical plants can mislead operators. Operators discount an abnormal reading, thinking it is part of the routine fluctuation. Generally they are correct; occasionally they are wrong. At TMI the operators knew the temperature on the hot leg of a drain pipe was abnormal because of a leak. During the accident, when the temperature exceeded even this abnormal reading, they treated it as a particularly wide fluctuation, whereas it really indicated an open relief valve.

In 1977 New York City experienced a massive and very costly blackout. One key contribution to the accident was an operator's expectation about the default reading for current flowing over a particular line. Normally that line carried little or no current. The operator did not know there had been two relay failures—one that would automatically lead to a high flow of current over that line; and a second that blocked the flow over the line. The operator treated the zero current reading as normal. In fact, it was abnormal, but only in this particular set of circumstances. This ambiguity led to a systematic, by-the-book sequence of actions to handle the problems that were showing up in other parts of the system, ending in the system being brought to a halt.¹⁰ The only evidence the operator could really "see," in terms of sensory confirmation, was the lights going out in his control room.

In a typical (linear) manufacturing plant both errors and interactions

are more visible. If an operator misunderstands an instruction the supervisor will probably know it soon enough because she will see the operator set up for the wrong task. I am told that one of the advantages of old-style steam valves, where the stem of the valve will rise when opened, is that a quick glance by an operator or supervisor over a huge room will show which valves are open, and which are shut—they just stick up when open. If an employee heads for the wrong valve after misunderstanding an order, that will be quite visible too. In complex systems, where not even a tip of an iceberg is visible, the communication must be exact, the dial correct, the switch position obvious, the reading direct and "on-line."

The problem of indirect or inferential information sources is compounded by the lack of redundancy available to complex systems. If we stopped to notice, we would observe that our daily life is full of missed or misunderstood signals and faulty information. A great deal of our speech is devoted to redundancy—saying the same thing over and over, or repeating it in a slightly different way. We know from experience that the person we are talking to may be in a different cognitive framework, framing our remarks to "hear" that which he expects to hear, not what he is being told. The listener suppresses such words as "not" or "no" because he doesn't expect to hear them. Indeed, he does not "hear" them, in the literal sense of processing in his brain the sounds that enter the ear. All sorts of trivial misunderstandings, and some quite serious ones, occur in normal conversation. We should not be surprised, then, if ambiguous or indirect information sources in complex systems are subject to misinterpretation.

Transformation Processes

As we gain more experience with systems, and design them more effectively, the high degree of interactiveness may be reduced. In Chapter 5 we shall examine the case of air traffic control, where exactly this appears to have happened. It is also true that a poorly trained or inexperienced operator may see a system as replete with unsuspected interactions or "traps," but after gaining experience may find it to be more linear. Finally, new technological breakthroughs may bring linearity into a once complexly interactive system, as when the jet engine replaced the piston engine, or the transistor replaced the vacuum tube. Though there is a

tendency to then demand more of the redesigned system, and thus build in new interactions, as has been dramatically recounted in the case of a jet fighter,¹¹ we should emphasize that new or improved designs and more operator experience may reduce the possibilities of unexpected interactions. Some we shall examine in this book are aspects of marine transport, the air traffic control system, dams, and mining. But our major concern is with systems that appear to be irretrievably complex, or at least will remain so for the next few decades. Generally these are systems that transform their raw materials, rather than fabricate or assemble them.

Transformation processes exist in recombinant DNA technology, chemical plants, nuclear power production, nuclear weapons, and some aspects of space missions. Most of these are quite new, but it is significant that chemical processing is not. While experience has helped reduce accidents, accidents continue to plague transformation processes that are fifty years old. These are processes that can be described, but not really understood. They were often discovered through trial and error, and what passes for understanding is really only a description of something that works. Some of industrial chemical production is of this nature; much of iron and steel production was of this nature, although it has been greatly altered through scientific research.

The existence of transformation processes without full understanding certainly characterizes nuclear power; recall that the nuclear scientist who advised Governor Thornburg of Pennsylvania during the accident had asserted three years before in a scientific journal that there could be no problem with a zirconium-water reaction—the process was well understood. Yet precisely this problem produced the hydrogen bubble. Each space mission introduces new systems that may be unreliable because of lack of knowledge. Experience will increase understanding of some of these problems (our rockets do not blow up on the launch pad as readily as they once did), but the performance of space vehicles is still subject to much uncertainty. Recombinant DNA research is fraught with gaps in knowledge. Limited knowledge, then, allows unsuspected interactions, and requires many control parameters and indirect sources of information.

This completes our discussion of the attributes of systems with complex interactions. To summarize, complex systems are characterized by:

- Proximity of parts or units that are not in a production sequence;
- many common mode connections between components (parts, units, or subsystems) not in a production sequence;

- unfamiliar or unintended feedback loops;
- many control parameters with potential interactions;
- indirect or inferential information sources; and
- limited understanding of some processes.

Complex systems are not necessarily high-risk systems with catastrophic potential; universities, research and development firms, and some government bureaucracies are complex systems, as I will argue shortly. We can now briefly contrast complex and linear systems.

Linear Systems

The parts or units of linear systems that are not in a direct production sequence tend to be spatially spread out. Fabrication or assembly allow this, while transformation processes often have to be compact. Linear systems lack the common-mode connections that require proximity. It is also a design criterion to separate various stages of production for sheer ease of maintenance access or replacement of equipment. Linear systems not only have spatial segregation of separate phases of production, but within production sequences the links are few and sequential, allowing damaged components to be pulled out with minimal disturbance to the rest of the system. In complex systems, removing a component or shutting it down means temporarily severing numerous ties with consequent readjustments, capping, product storage, removal to get access, and reconfigurations because parts and units tend to be multiply linked. Linear systems also favor serial production—a series of linked but semi-independent production steps—rather than what organizational theorist James Thompson calls “pooled interdependence,” where all components (including operators) must coordinate their input if the system is to function at all.¹²

In contrast to complex systems, there is minimal specialization of labor, materials, and pools of supply in linear systems. Equipment is specialized (“dedicated,” engineers call it), but the people who run it tend to be generalists. Operators are trained on several tasks because they tend to rotate, bid on various jobs, or fill in for people. Maintenance people can operate equipment in an emergency, operators perform emergency maintenance. There are, of course, limits to such substitutions, including the demands set by unions. Later we shall see that substitutability is important for recovery from an accident—people can fill in and know some-

thing about the other person's job. Here, in discussing complexity and linearity, the emphasis is upon awareness of interdependencies. In complex systems, not only are unanticipated interdependencies more likely to emerge because of a failure of a part or unit, but those operating the system (or managing it) are less likely, because of specialized roles and knowledge, to predict, note, or be able to diagnose the interdependency before the incident escalates to an accident.

Consider some key jobs in complex systems and note how little substitutability there is between them: fighter pilots and maintenance; bombardiers and pilots and navigators; nuclear plant engineers and operators; operators and welders and other maintenance people rated to do specialized jobs; lab technicians and biochemists in a recombinant DNA firm; chemists with advanced degrees and lab technicians, in a chemical plant, or chemical engineers and control room operators; astronauts and ground control managers; navigators, radio operators, captain, and helmsmen on a ship.

Though I don't want to claim a vast difference between employees in complex and linear systems, the latter appear to have fewer specialized and esoteric skills, allowing more awareness of interdependencies if they appear. The welder in a nuclear plant is more specialized (and specially rated), and presumably more isolated from other personnel, than the welder in a fabrication plant. Specialized personnel tend not to bridge the wide range of possible interactions; generalists, rather than specialists, are perhaps more likely to see unexpected connections and cope with them.

What is true of labor is also true of materials and supplies. If materials and supplies are substitutable, more latitude of response is available, limiting failures to incidents and preventing failures in the first place. But complex systems appear to have more exacting requirements for materials and supplies; the fuel cannot be off-standard, nor one fuel substituted for another, whether we are dealing with nuclear plants, aircraft, spacecraft, or chemical production. Substitutions are more likely in linear systems.

Finally, linear systems have minimal feedback loops, and thus less opportunity to baffle designers or operators. There are fewer interactions of control parameters, because controls are more decentralized, and attached to special-purpose equipment. And the information used to run the system is more likely to be directly received, and to reflect actual operations.

The two systems are summarized in Table 3.1, with some summary terms listed that will be used in the rest of the book.

TABLE 3.1
Complex vs. Linear Systems

Complex Systems		Linear Systems	
Tight spacing of equipment		Equipment spread out	
Proximate production steps		Segregated production steps	
Many common-mode connections of components not in production sequence		Common-mode connections limited to power supply and environment	
Limited isolation of failed components		Easy isolation of failed components	
Personnel specialization limits awareness of interdependencies		Less personnel specialization	
Limited substitution of supplies and materials		Extensive substitution of supplies and materials	
Unfamiliar or unintended feedback loops		Few unfamiliar or unintended feedback loops	
Many control parameters with potential interactions		Control parameters few, direct, and segregated	
Indirect or inferential information sources		Direct, on-line information sources	
Limited understanding of some processes (associated with transformation processes)		Extensive understanding of all processes (typically fabrication or assembly processes)	

Summary Terms		Linear Systems	
Complex Systems		Linear Systems	
Proximity		Spatial segregation	
Common-mode connections		Dedicated connections	
Interconnected subsystems		Segregated subsystems	
Limited substitutions		Easy substitutions	
Feedback loops		Few feedback loops	
Multiple and interacting controls		Single purpose, segregated controls	
Indirect information		Direct information	
Limited understanding		Extensive understanding	

Which Is Best?

This litany of problems with complex systems and the advantages of linear systems might suggest the latter are much preferable and complex systems should be made linear. Unfortunately, this is not the case. Complex systems are more efficient (in the narrow terms of production efficiency, which neglects accident hazards) than linear systems. There is less slack, less underutilized space, less tolerance of low-quality performance, and more multifunction components. From this point of view, for design and hardware efficiency, complexity is desirable.

It also appears that we have comparatively little choice in the design of some of our systems. Some complex systems can be redesigned to be

more linear, such as in air traffic control or with the substitution of jet engines for the very interactive piston engines. Nuclear plants could be made marginally less complex if the spent storage pool were removed from the premises. (Should a plant with a full load of fresh rods in its pool have to be evacuated and the pool left unattended or without power, the water could boil off in a few days. Without the water, they would commence to fission like sparklers on the Fourth of July.) This would eliminate some proximity and common-mode problems. Using one control room to run two different reactors also seems like an unnecessary source of common-mode problems. The steam generation and turbine systems could be separated from the nuclear system to reduce unintended feedback loops, though at considerable expense. And perhaps, as we noted in the last chapter, more "forgiving" designs exist. But by and large no extensive reductions in complexity seem possible in the nuclear power industry. The transformation system simply requires many nonlinear interactions. The same is true of chemical plants such as refineries; there is probably no efficient way to crack crude oil except in a highly interactive system.

On the whole, we have complex systems because we don't know how to produce the output through linear systems. If these complex systems also have catastrophic potential, then we had better consider alternative ways of getting the product, or abandoning the product entirely. Short of that there does not seem to be a choice possible between complexity and linearity for high-risk systems. Complexity is inherent in some forms of production. It is not intrinsically undesirable; we welcome complexity in some bureaucracies and resist the "rationalization" of our disorderly life because the unexpected interactions lead to innovations, amuse or interest us, or provide variety. If the system has catastrophic potential, however, and we cannot prevent the propagation of incidents and intervene before an accident has occurred, the issue is much graver. Let us see what the idea of tight and loose coupling has to say about this. It is our second major dimension in analyzing systems.

Tight and Loose Coupling

Engineers speak of tight and loose coupling in the normal course of their work, and in that context the terms are quite clear. It's not as simple as coupling a garden hose to the faucet, and doing it tightly if we don't want to get sprayed, but that is a fairly close analogy; *tight coupling* is a me-

chanical term meaning there is no slack or buffer or give between two items. What happens in one directly affects what happens in the other.

Sociologists and social psychologists took up the term in the mid-1970s to conceptualize a particular phenomenon.¹³ Some public service organizations, public schools in particular, seemed to be characterized by an unusually large gap between official programs and actual behavior. In public schools researchers investigated the gap between new techniques and actual changes in the behavior of teachers, or between new programs and what the students actually learned. One might expect hierarchically organized bureaucracies to be capable of altering the behavior of subordinate personnel such as teachers and capable of producing outcomes in students that bear some relationship to the official goals. The explanation for the school's failure to achieve its goals was that while the programs and goals were real enough, they were only loosely connected to other matters with which the organization had to be preoccupied, such as political demands from the environment, the autonomy of professional teachers, and the ineffective mobilization of parental demands.

For example, there could be a tight connection (responsiveness) between a remedial program required by the school district and the school in that it is funded, staffed, and given space and a place in the curriculum. But the program might only be loosely coupled to other parts of the school. For example, two new teachers, presumably hired for the program, are actually the ones the school has long wanted to hire—a respected art teacher and a person who can teach a course in computers. They are used for these positions, and the two teachers who are the least respected by parents and students are assigned to the new program. The decisions might almost seem to have been made by themselves, so loosely coupled is the demand for the new program with the existing problems of the school.

The program budget is a concrete, discrete item, but its appearance does not result in purchasing the proper supplies and doing the expected remodeling. Instead, the supply orders are delayed, and the money used for gym equipment and refinishing the gym floor. There is no intent of fraud; the latter have long been pressing items, subject to an informal agreement with the district that the expenditures can be made just as soon as there is some slack in the system. The school principal reasons that the remedial program will be taken care of when the new budget passes. Space for the new program might turn out to be the most inconvenient space, on the incontrovertible grounds that all the convenient space is utilized, and the program should not disrupt successful existing programs. Existing programs are buffered, in this way, from the impact

of the new program. However, the district's intent was to alter the school's existing priorities, which would require that the remedial program should have an impact on—be tightly coupled to, and change—the existing priorities.

Finally, the school could find it easy—almost inadvertent—to assign students to the program not on the basis of their reading difficulties, but because they are disruptive, impolite, or racially stereotyped. They all look as if they need “extra help.” And the classroom materials, carefully designed by educators in elite universities, may bear little relationship to the needs of the stigmatized students.

In all this a number of connections are being made—student social status and program assignment; supply shortages and the power of certain departments; low performing teachers and periphery assignments. The system is not lacking in connections, but they are “loose” ones. Some are invoked for this program but perhaps not for the next. Some were invoked by accident—a chance discovery that the way the budget was drawn up it could cover costs of the gym floor. Others might have depended upon fleeting events, perhaps a recent controversy regarding the effectiveness of two teachers. Certainly the connection between the official mandate for a remedial program and the actual practice is quite loose. The characteristics of the system—its loosely coupled nature—make it possible for it to respond to the outside demand in such a loose fashion.

Loosely coupled systems tend to have ambiguous or perhaps flexible performance standards, and they may, as in the case of the school, have little consumer monitoring, so the absence of the intended connection can remain unobserved.

You might be tempted to call it a very inefficient system, or even a rip-off of state or federal funds, and call for “tight coupling.” This would make the remedial program work as everyone outside the system thought it should work. But it would be a mistake to call it inefficient. The system is quite efficient for accomplishing many things that many participants desire; they are just not what the school district or the federal government, which supplies the money, had in mind. A bit of authoritarian leadership designed to put the official program in effect might well dis-able parts of the system that others value. If the poor teachers remained in the regular curriculum, parents would continue to object to them; if the college preparatory subjects were put in the temporary building (temporary since the end of World War II) they might also object; everyone uses the gym and it needs fixing; and what is one to do with disruptive students who reportedly interfere with the learning of the docile ones?

Loose coupling, then, allows certain parts of the system to express themselves according to their own logic or interests. Tight coupling restricts this. Loose coupling, however, is not the same as disorganization, unless we mean lack of centralized control by that term. Informally, the school is well organized, with a variety of coherent (though occasional shifting or episodic) interests, mechanisms for accommodative arrangements, slack resources to meet unexpected challenges, and stable interaction patterns. The degree of organization is independent of the degree of coupling.

We have not strayed as far from the engineer's usage of tight and loose coupling as it may seem. Elaborating the concept as used by organizational theorists will allow us to examine the responsiveness of systems to failures, or to shocks. Loosely coupled systems, whether for good or ill, can incorporate shocks and failures and pressures for change without destabilization. Tightly coupled systems will respond more quickly to these perturbations, but the response may be disastrous. Both types of systems have their virtues and vices.

For example, a continuous processing plant requires tight coupling. In some continuous processing plants where the technology is well understood and only linear interactions take place, such as pharmaceutical plants, bakeries, confectionaries, or ball-bearing plants, the characteristics of the product are altered frequently in response to market demands. Processes are also altered frequently to reflect changes in raw materials or operating conditions. Decisions to alter the process or shift to a slightly different product must proceed quickly through the organization, producing nearly unreflective changes in operator behavior. Resources are strictly allocated, schedules strictly followed, and reporting systems must be exact. Surveillance, both of processes and people, is continuous, though it is usually through remote and unobtrusive indicators—at least when higher-level personnel are monitored. Deviations from standards are quickly noted or reported, since a long product stream is affected. Responses to deviations must be standardized and immediate. The result is high operating efficiency. To loosely couple such a production process would be to invite disasters, as well as inefficiencies. If the system is linear, rather than complexly interactive, tight coupling appears to be the optimum mode of organization.

In contrast, consider a loosely coupled linear aircraft manufacturing and assembly plant. Fabrication of the tail section will be separate from the fabrication of the fuselage section it will be attached to. Different metals are involved, different tolerances, and different heat-treating processes. The constraints on the two sections are different, and their

interdependence is minimal; the interdependence is taken care of in the design stage. Personnel practices may be different too, since particularly skilled personnel are needed for the tail construction, whereas fewer skills are needed for the fuselage, which involves more routine tasks than tail construction. This can lead to personnel decisions that resemble the allocation practices we discussed in the school, reflecting unit power and politics. Quality control, instead of being "built into" the system as in continuous processing, may be a "stand-alone" function, performing its own tests and inspections, physically isolated in the plant, and deliberately not integrated into the other functions because its independence must be maintained. But this buffering of quality control also means loose coupling, with the possibility of parochial interests developing, or unauthorized bargains with units.¹⁴

Characteristics of Coupling In Systems

We are now ready to more systematically lay out what is meant by tight and loose coupling in systems. These characteristics are reasonably independent of our other major dimensions—complex interactions and linearity.

1. Tightly coupled systems have more time-dependent processes: they cannot wait or stand by until attended to. Sometimes this is expressly for efficiency reasons, but generally it is because the production process does not allow for cooling and reheating, for forgetting and then relearning. Storage room may not be available, so products must move through continuously. Reactions, as in chemical plants, are almost instantaneous and cannot be delayed or extended. In loosely coupled systems, delays are possible; processes can remain in a standby mode; partially finished products (tail assemblies or students) will not change much while waiting.

2. The sequences in tightly coupled systems are more invariant. B must follow A, because that is the only way to make the product. Partially finished products cannot be rerouted to have Y done to them before X; Y depends upon X's having been performed. Though it may be expensive, in an aircraft assembly line a door or seat or the radio controls could be added later if there were a disruption. This is not the case for a nuclear or chemical plant, or a pharmaceutical production line. In a trade school the sequence of courses is carefully prescribed and more

tightly coupled than a university where it does not matter much when the science or math or language requirement is met.

3. In tightly coupled systems, not only are the specific sequences invariant, but the overall design of the process allows only one way to reach the production goal. A nuclear plant cannot produce electricity by shifting to oil or coal as a fuel; but oil plants can shift to coal and vice versa with little or no reconversion. While continuous processing plants (tightly coupled) can vary product mixes and production volume within limits, they cannot make significant changes in the way the product is made. However, manufacturing plants (generally loosely coupled) may, even on a temporary basis, eliminate heat-treating by substituting another metal; may shift from unit to batch production; contract out for subassemblies; install or remove robotic devices; substitute plastic for metal, and so on. In this sense, there are many ways to produce the item. But for tightly coupled systems such as dams, chemical plants, power grids, bakeries, and recombinant DNA technologies, there is little flexibility. Loosely coupled systems are said to have "equifinality"—many ways to skin the cat; tightly coupled ones have "unifinality."

4. Tightly coupled systems have little slack. Quantities must be precise; resources cannot be substituted for one another; wasted supplies may overload the process; failed equipment entails a shutdown because the temporary substitution of other equipment is not possible. No organization makes a virtue out of wasting supplies or equipment, but some can do so without bringing the system down or damaging it. In loosely coupled systems, supplies and equipment and manpower can be wasted without great cost to the system. Something can be done twice if it is not correct the first time; one can temporarily get by with lower quality in supplies or products in the production line. The lower quality goods may have to be rejected in the end, but the technical system is not damaged in the meantime.

Recovery from Failure

Coupling is particularly germane to recovery from the inevitable component failures that occur. One important difference between tightly and loosely coupled systems deserves a more extended comment in this connection. In tightly coupled systems the buffers and redundancies and substitutions must be designed in; they must be thought of in advance.

In loosely coupled systems there is a better chance that expedient, spur-of-the-moment buffers and redundancies and substitutions can be found, even though they were not planned ahead of time.

Since failures occur in all systems, means to recovery are critical. One should be able to prevent an incident, a failure of a part or unit, from spreading. All systems design-in safety devices to this end. But in tightly coupled systems, the recovery aids are largely limited to deliberate, designed-in aids, such as engineered safety devices (in a nuclear plant, emergency coolant pumps and an emergency supply coolant) or engineered safety features (a more general category, which would include a buffering wall between the core and the source of coolant). While some jury-rigging is possible, such possibilities are limited, because of time-dependent sequences, invariant sequences, unifinality, and the absence of slack. In loosely coupled systems, in addition to ESDs and ESFs, fortuitous recovery aids are often possible. Failures can be patched more easily, a temporary rig can be set up, a crane moved over, a pancake landing made without power. There may, fortuitously, be plenty of space separating a burning subsystem from other systems; it may be possible and relatively harmless to flood an area to put out a fire, though no designer planned for that. Tightly coupled systems offer few such opportunities. Whether the interactions are complex or linear, they cannot be temporarily altered.

This does not mean that loosely coupled systems necessarily have sufficient designed-in safety devices; typically, designers perceive they have a safety margin in the form of fortuitous safety devices, and neglect to install even quite obvious ones. Most of the safety devices required after inspections by the Occupational Safety and Health Administration (OSHA) are fairly obvious items, such as railings, rough treads on stairs or passageways, emergency switches to shut off equipment, and requirements that the coupling for hoses be arranged so that an explosive gas such as hydrogen cannot be inadvertently supplied to an area that needs an inert gas. Even in loosely coupled systems there is still enormous room for improvement in safety features.

Tightly coupled systems are not completely devoid of unplanned safety devices. In two of the most famous nuclear plant accidents, Browns Ferry and TMI, imaginative jury-rigging was possible and operators were able to save the systems through fortuitous means. At TMI two pumps were put into service to keep the coolant circulating, even though neither was designed for core cooling. Subjected to intense radiation they were not designed to survive; one of them failed rather quickly, but the other kept going for days, until natural circulation could be established. Some-

TABLE 3.2
Tight and Loose Coupling Tendencies

Tight Coupling	Loose Coupling
Delays in processing not possible	Processing delays possible
Invariant sequences	Order of sequences can be changed
Only one method to achieve goal	Alternative methods available
Little slack possible in supplies, equipment, personnel	Slack in resources possible
Buffers and redundancies are designed-in, deliberate	Buffers and redundancies fortuitously available
Substitutions of supplies, equipment, personnel limited and designed-in	Substitutions fortuitously available

thing more complex but similar took place at Browns Ferry. The industry claimed that the recovery proved that the safety features worked; but the designed-in ones did not work. In both accidents the critical safety features were disabled.

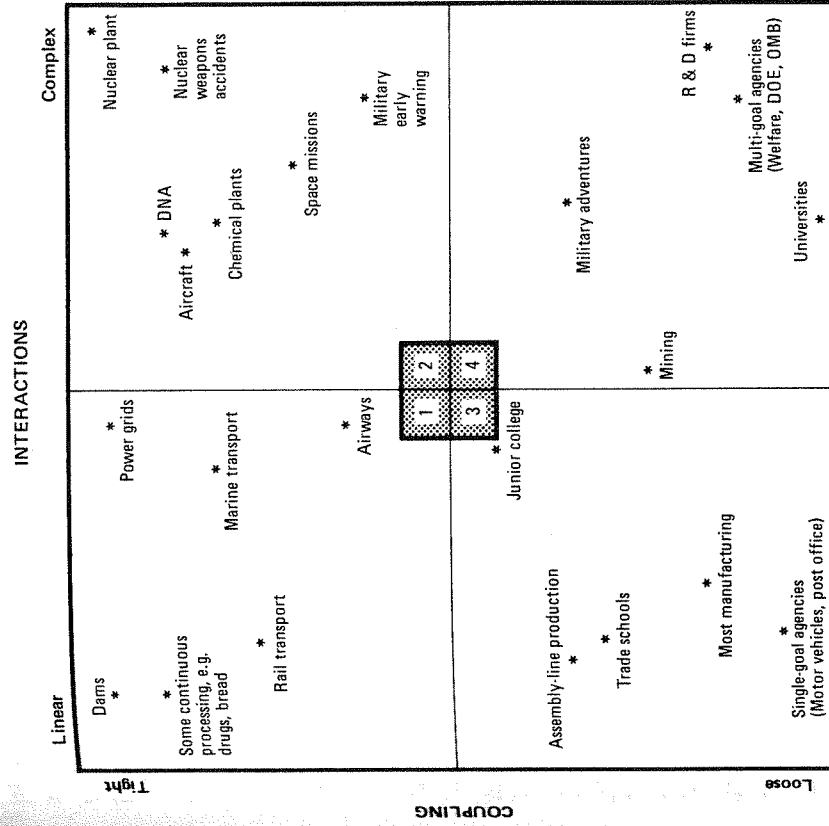
What is true for buffers and redundancies is also true for substitutions of equipment, processes, and personnel. Tightly coupled systems offer few occasions for such fortuitous substitutions; loosely coupled ones offer many.

The characteristics of the two systems are summarized in Table 3.2. These are tendencies; no system has, for example, absolutely invariant sequences. Nor is a system likely to have all of the various characteristics.

The Organizational World According to Complexity and Coupling

Figure 3.1, the Interaction/Coupling Chart (I/C chart) puts interaction and coupling together in a two-variable array. The placement of systems is based entirely on subjective judgments on my part; at present there is no reliable way to measure these two variables, interaction and coupling. One may well quarrel with the placement in the figure. One good reason for a quarrel is that there is no precise specification of just what constitutes the system. Consider marine transport. On the high seas the system includes the ship, radio communication, the weather, and perhaps one other ship. But as the ship enters a crowded channel, we have not only the weather, but bank effects (the suction created as the ship passes close

FIGURE 3.1
Interaction/Coupling Chart



to the underwater bank of a channel), tides and current flow, wrecks and rocks, bridges, tows and other ships, a crowded radio channel, and navigation lights mixed in with the lights of highways, industrial plants, and distant towers. The high seas and the channel systems are very different. With "flying" we mix commercial airlines with recreational or sport flying. Mining includes strip mining, which should be closer to manufacturing, and underground mining. These ambiguities are a problem, but for the sake of a general argument I think the rough placement of vaguely defined systems is still useful.*

*One serious problem cannot be avoided, but should be mentioned: To some unknown extent it is quite possible that the degree of coupling and type of interactions have been

By combining our two variables in this way, a number of conclusions can be made. First, it is clear that the two variables are largely independent. Examine the top of the chart from left to right. Dams, power grids, and nuclear plants are all roughly on the same line, indicating a similar degree of tight coupling. But they differ greatly on the interaction variable. While there are few unexpected interactions possible in dams, and not that many in power grids, there are many in nuclear plants.

Or, looking across the bottom, universities and post offices are quite loosely coupled. If something goes wrong in either of these, there is plenty of time for recovery, nor do things have to be in a precise order. In the post office, mail can pile up for a while in a buffer stack without undue alarm; people tolerate the Christmas rush, just as students tolerate lines at fall registration. But in contrast to universities, post offices do not have many unexpected interactions—it is a fairly well laid out (linear) production sequence without a lot of branching paths or feedback loops. This is not true of universities. Here there are multiple functions—teaching, research, and public service, for example—and they can interact in unexpected ways. Indeed, they are expected to; synergistic interactions are desired, though “negative synergy” is also possible.

For example, let us suppose that a university reviews the record of an assistant professor of sociology, and decides that not enough has been published to warrant promotion to associate professor and the tenure or job security that goes with it. It is all quite straightforward, since research and publication are clearly specified as performance criteria for faculty and are outputs of the system. But imagine that the faculty member is a very good teacher, and the students protest the action. Imagine also that dismissal will threaten a public service program that she ran. Not only do the students protest, but some influential community members protest. To convince the students that the department is not indifferent to teaching quality, it has to institute a teacher evaluation program, and this now threatens the status of the senior faculty (associate or full professors with tenure) who have not had to pay much attention to this criterion before. Furthermore, the local churches send a delegation of church officials to protest the demise of the service program, and in the process,

inferred from a rough idea of the frequency of system accidents in the various systems, rather than derived from analysis of the properties of the systems independent of the nature of their failures. That is, if there are few accidents caused by air traffic control, that “must” mean it is not highly complex and tightly coupled, and then evidence for that conclusion is sought. Since the analytical scheme evolved from examination of many systems, there is no way to avoid this possible circularity. The scheme would have to be tested by examining systems not included here, as well as by collecting data based upon more rigorous concepts and verifying the placement of the systems that are included here.

demand more say in university personnel practices. The dean now has three problems on her hands: the department has to publish more (to justify firing the assistant professor), spend more resources on teaching to mollify the students, who have now made teaching a part of the formal personnel evaluation, and face a mobilized external group that demands some say in personnel decisions. This is an example of complex interactions; in contrast to the linear post office, the university is high on complexity.

But note that we are not likely to have a “system accident.” Because universities are loosely coupled there is ample slack to limit the impact of this one personnel decision on other areas. Senior faculty members can be quietly reassured that student evaluations will be processed slowly, and evidence of the unreliability of evaluations will be discreetly passed about. (In one case that I witnessed, the university couldn’t find the \$2,000 needed to keypunch and process a student-initiated teacher evaluation program, the university’s first.) Community members will be promised more access to the Student Activities Council, but realistically they will not be expected to attend often. The university research foundation is quietly told to favor sociology applications in the next cycle of grants, or joint appointments are handed out to make the sociology program look more productive. No interruption of output is experienced. The event remains an incident, with damage limited to the “part”—the teacher. While interactive, recovery is easy because of loose coupling.

The post office and the university, then, are similar with regard to coupling; both can recover from upsets because sequences are not that inevitable, there is slack in resources, substitutions are possible. But the post office is largely limited to linear interactions, while the university is full of many potentially complex ones that can reach unexpectedly into other parts of the system.

Educational systems are represented in three places in the I/C chart: universities are loosely coupled and complex, as we have seen; trade schools are somewhat loosely coupled but very linear; and junior colleges are near the middle of both dimensions. Trade schools have a well-defined curriculum and fairly explicit set of subjects students are expected to learn. This permits linear interactions to dominate and there are few surprises. However, the cafeteria style of offerings allows students to retake courses, making it somewhat more loosely coupled than junior colleges, and trade schools are also staffed by part-time, often temporary instructors who come and go with the changing student demands. Junior colleges are a bit more tightly coupled because of sequencing within programs (blocks of courses) and because of personnel inflexibility. They are

Petrochemical Plants

less linear than trade schools because they undertake more functions, in particular, socialization and value impregnation (though not nearly as much as universities do). Neither of these institutions would be expected to have system accidents that result in stopping the output of subsystems; if such subsystems as courses or community programs do fail, it would be because of component failures—lack of students, instructors, or community support.

Conclusion

This completes our major analytic work in this book. We have accomplished quite a bit: defined accidents, distinguishing them from incidents; defined various classes of victims in a way that allows us to better assess catastrophic potential; defined system and component failure accidents; and defined our two key concepts, the types of interaction (complex and linear) and the types of coupling (loose and tight). These variables have been laid out so that we can locate organizations or activities that interest us and show how these two variables, interaction and coupling, can vary independently of each other.

We now have the basic tools to proceed through a variety of systems and examine them in terms of the type of accidents they have, and their catastrophic potential. We shall start with chemical plants, which resemble nuclear power plants in their system accidents, though not in their consequences, yet which illustrate a comparatively old and stable technology.

Petrochemical plants produce far fewer headlines and much less controversy than nuclear power plants. They have been around for about one hundred years, so ample engineering experience exists, and the public is familiar with the sight of their gangly towers and squat storage tanks. Fires rage uncontrolled in them at times, and a few operators might be killed, but we do not have protest marches, a Chemical Regulatory Commission, scientific panels and conferences that are covered by the media, or a search for alternatives to our plastics and gasoline. It is a low-profile industry—deliberately, as we shall see. It is also quite safe; the fires and explosions kill few operators and few bystanders.

Yet the chemical industry is far from an unpromising subject for accident analysis; in fact, it provides some of the best examples of system accidents we shall come across. It is quite tightly coupled, and has many complexly interactive components. We are concerned with it primarily because it illustrates the presence of system accidents in a mature, well-run industry that has a substantial economic incentive to prevent accidents. None of these factors are true of the nuclear power industry; it is new, poorly run, and short of catastrophic accidents that would close all plants overnight, accidents by and large only generate costs that can be

Exotics: Space, Weapons, and DNA

This final data chapter deals with three very high-tech systems, and that is about all that links them together. One has almost no catastrophic potential (space missions); one has the ultimate catastrophic potential (nuclear weapons); and the third, recombinant DNA research and production, or DNA for short, has hardly begun, but could well develop in a direction that would be second in catastrophic potential only to nuclear war. Each system contributes to our argument in different ways, providing further evidence that system accidents are inevitable in complex, tightly coupled systems.

The first section, on space missions, will elaborate in a clear and dramatic fashion a point touched upon in the sea stories and the accounts of nuclear and chemical plant accidents: the attempt by the great designers of complex systems to do without the lowly operators. In the space program the operators were hardly lowly, since they were experienced, well-trained test pilots, the Waldo Peppers and the great Santinis of the space age. But they were treated as lowly operators until the designs came apart and their managers became befuddled, and only the pilot-turned-astro-

naut could save the missions. The space missions illustrate that even where the talent and the funds are ample, and errors are likely to be displayed before a huge television audience, system accidents cannot be avoided. I have argued throughout this book that we should give all risky systems more quality control and training than we do, but also that where complexity and coupling lie, it will not be enough. We gave the space missions everything we had, but the system accidents still occurred. This is not a system with catastrophic potential; the victims are first-party victims. Catastrophic potential resides in most, but not all, complex and tightly coupled systems.

The weapons systems have hardly been starved for funds or engineering talent either, but accidents abound. I think they are system accidents, though it is hard to tell from the limited public data. In any case, the catastrophic potential is infinite in the case of nuclear weapons. DNA involves genetic engineering, and we have had almost no experience with genetic engineering on a production basis (the most threatening part of that endeavor) and very little on a research basis, so we have no accidents to explore. But a scenario is imaginable, and if we are to be prudent as a society, scenarios should be explored.

We will start out on a relatively positive note by examining the space missions. The centerpiece of this part of the chapter will be an account of an extraordinary mission, Apollo 13, which commenced with a system accident, and ended with a recovery that dramatically illustrates the most exemplary attributes of both humans and their machines. The event will tell us something further about complexity and coupling: the recovery was possible because ground controllers were able to make the system more linear and more loosely coupled, and to put the operators back into the control loop that rarely included them.

No such encouraging lessons come from the section on nuclear weapons and early warning systems. We will not dwell on "the fate of the earth," that is, the destructive power of nuclear weapons, but on the limits of human capabilities and the even narrower limits of organizational capabilities. There is much to fear from accidents with nuclear weapons such as dropping them or an accidental launch, but with regard to firing them after a false warning we reach a surprising conclusion, one I was not prepared for: because of the safety systems involved in a launch-on-warning scenario, it is virtually impossible for well-intended actions to bring about an accidental attack (malevolence or derangement is something else). In one sense this is not all that comforting, since if there were a true warning that the Russian missiles were coming, it looks as if it would also be nearly impossible for there to be an *intended*

launch, so complex and prone to failure is this system. It is an interesting case to reflect upon: at some point does the complexity of a system and its coupling become so enormous that a system no longer exists? Since our ballistic weapons system has never been called upon to perform (it cannot even be tested), we cannot be sure that it really constitutes a viable system. It just may collapse in confusion!

Finally, there is the recondite, novel, and thrilling field of recombinant DNA research. The potentials here for human benefit appear to be more extraordinary than all the other technologies put together. The potentials for human disaster are equally unprecedented, and rival that of nuclear holocaust—if there can be a question of rivals where extinction is involved. With breakneck speed we (that is, primarily the oil companies) are proceeding down a thoroughly unknown path, without brakes, without headlights, in search of undreamed amounts of private profits. The system accidents may, indeed, already have occurred, giving new meaning to Yeats's oft quoted line, "What rough beast, its hour come round at last, slouches towards Bethlehem to be born."

Space Missions

Industry Errs

The conquest, if it be that, of outer space illustrates a normal learning curve for technological activity. The rockets stopped blowing up and achieved remarkable reliability. But the history of space missions also illustrates that there remains an obdurate residual propensity for system accidents.

In Chapter 2, I argued that every industrial activity exhibits organizational failures, incompetence, greed, and some criminality. The space program certainly performed better in these respects than the nuclear industry. I know of no criminal activity. But even this vastly favored program exhibits enough incompetence and organizational failures to remind us once again that nothing is perfect. A major contractor, North American Aviation, was known as "Brand X" in the trade—a derogatory term indicating "cheap." The abundance of funds and cost-plus contracts may have encouraged such instances as a workman throwing the remains of his lunch into a costly automatic chart-recording instrument; a sealing device ruined by gouging it with a screwdriver while trying to force it

into place; a booster rocket shutting itself off and leaving Walter Schirra and his *Gemini 6* on the pad because a dust cap had been left on a pressurizing line in one engine, even though a quality control inspector had approved the work as done.¹ No, McDonnell Douglas, the vendor in this case, was not duplicating its profit-maximizing DC-10 efforts; these examples come from facilities swarming with government inspectors in a program that, while profitable for contractors, would not generate the long-range profits that could come from aircraft sales.

The Mercury part of the space program, involving only orbital single-astronaut flights, was reviewed by NASA in a 440-page document that *New York Times* reporter John Finney called a "remarkably harsh indictment of American industry."² Among the problems: spare parts that were 50 percent defective; capsules with more than 500 defects; batteries with holes in them; vital electronic parts improperly soldered; valves improperly installed, leading to attitude-control problems on flights; dirty gas pressure regulators; and contaminated oxygen and water for breathing and drinking. The contractors denied they performed anything but superbly, but did not respond to the specific charges.

After the tragic fire on the launch pad that killed three astronauts, an Apollo inquiry board was appointed, and it in turn drew upon the research of twenty-one panels of experts—1,500 overall—when it prepared its report. Six of the eight board members were NASA employees, so there was some fear of a whitewash—NASA was investigating itself and its prime contractor, Rockwell International. However, the report, 3,000 pages long, was described in the press as fairly scathing. Among the findings: "Adequate safety precautions were neither established nor observed for this test." "The over-all communication system was unsatisfactory." "Deficiencies existed in command module design, workmanship and quality control, such as: . . . regulator failures, line failures and environmental control unit failures. . . . Coolant leakage at solder joints has been a chronic problem. . . . Deficiencies in design, manufacture, installation, rework and quality control existed in the electrical wiring."³ A wrench and other "noncertified equipment items" were left in the command module at the time of the test; 113 "significant" engineering orders were not accomplished when the module was delivered and records did not match up. Rockwell protested most of the conclusions, and one of its officers speculated that astronaut Grissom might have kicked the bundle of wires thought to have caused the spark, thereby installing operator error into even this accident—which occurred while the three astronauts were patiently reclining in their seats waiting for the next part of the test.

The most probable cause was our old friend, a bundle of wires whose insulation had rubbed off. Unfortunately, in a pure oxygen environment, that meant a flash fire.

Hardly noticed in the wake of that accident was another one, four days later, in a module simulator where the effect of near-pure oxygen atmosphere on blood-forming organs in rabbits was being investigated. Two enlisted men were killed when a work-lamp wire shorted and ignited the atmosphere. Like the launch-pad tragedy, this one was most likely a component failure accident.

The Gemini flights, a prelude to the Apollo mission, were near perfect, though there were close calls. But another prelude, the Ranger flights, designed to survey the moon, had five out of nine failures. *Ranger 6* failed, incidentally, for a reason that should be familiar to us—a safety device. In order to make sure the television cameras would come on to take pictures of the moon's surface, there were redundant power supplies and triggering circuits. According to a Babcock and Wilcox engineer, a short in a safety device (a testing circuit) depleted the power supplies by the time the Ranger reached the moon. The engineer notes that the more redundancy is used to promote safety, the more chance for spurious actuation; "redundancy is not always the correct design option to use."⁴

Most of the failures of the space program have not been death-dealing, and if they were, they were limited to first-party victims—the astronauts or technicians. However, in three cases of failures with plutonium power packs, the risks are potentially catastrophic, since plutonium is perhaps the most deadly substance known to humans. One of these power packs was retrieved successfully from the Santa Barbara Channel off California—barely off California; had it hit any place on land, whether near the city of Santa Barbara or not, the consequences would have been catastrophic. The second went into the Indian Ocean and is still there. The third was in a navigational satellite sent up in 1964 that failed to achieve orbit when its rocket engine failed. It reentered the atmosphere over the Indian Ocean and distributed 1 kilogram of plutonium-238 about the earth. By 1970 it was estimated that about 95 percent of it had settled on the ground or the earth's waters. The accident was estimated to produce a three-fold increase over the amount of plutonium contamination produced by all atmospheric nuclear weapons testing.⁵ This received almost no publicity, in contrast to the breakup of a Soviet nuclear-powered satellite in 1978 and another one in 1983. The first public mention of it may have been in a 1967 item in the journal *Science*.⁶

The industrial problems continued right up to the present-day shuttle program, though in the shuttle's defense, it has not had the ample fund-

ing of the moon missions. The engines failed or exploded frequently, the pesky tiles kept falling off in shipment or merely from being bumped by a forklift truck, the cost overruns mounted, the expected cost of each flight more than doubled, and the scheduled first flight was delayed ten years. After two failures of spacesuits in a November 1982 flight, NASA convened another panel of experts, and the conclusion was not much different from the panel studying the launch pad fire: "Egregious oversights" by the prime contractor, United Technologies. In one case two small plastic pins were missing, creating a leak in the spacesuit; the inspection sheet indicated they were in place, and the employee's supervisor had signed it. Worse yet, a stray steel chip was found in an exhaust vent of the suit's oxygen supply system—it could have caused an explosion of the suit, which could have blown a hole in the side of the shuttle.⁷ Tight coupling abounds in these systems. There can be little room for the all-too-human construction errors that will appear with the best of vendors, and certainly not for those we have been detailing.

But that is a minor point; more important is the appearance of system errors, and the means of recovery from failure. The space program provides a particularly insightful glimpse into something we have touched on before and will bring us closer to a major organizational issue: If there can be inexplicable interactions, who can best cope with them—the operators, or the design/management team? Some variants of organizational theory, reinforced by democratic values, generally assert that those closest to the disturbances in a system are in the best position to act upon them. Decentralization is the recommendation. But other variants and much of engineering logic assert that those close to disturbances cannot act fast enough, or with enough comprehension, to cope; therefore, designers should try to eliminate as many human tasks as possible and give them to machines, and managers should have the means to tell the system and/or its operators just what to do. Centralization is the recommendation. The conflict between these views runs through the space program. The regular appearance of mysterious interactions only escalated the conflict. We will consider this in the next section.

Twentieth-Century Taylorism

If everything in the DEPOSE system components could be perfect, we could drop the "O," the operator. The designer would be the operator; he or she would turn the system on and it would run, and the owner could turn it off when she wanted. The first suborbital and orbital flights were conceived of that way: the "astronaut" was not really needed, and went along as a test subject. (The term "pilot" was rejected; it implied control,

or piloting. An astronaut was a "star voyager," a more passive term.) Indeed, the first suborbital flight and the first orbital flights had no one aboard; the second in each category had a chimpanzee aboard—Ham, and Enos. As with the astronauts, the chimpanzees had been extensively conditioned not to blow it. The chimps manipulated dials that had nothing to do with the ship; it was just to test reaction times. For the human passengers, it was almost the same, and at one time it was felt that they should be given tranquilizers of some sort to make sure they did not interfere. (I am drawing here on the immensely entertaining, and exceptionally perceptive book by Tom Wolfe, *The Right Stuff*. He will be our guide for this section.⁸)

Think first of the "great designers," scientists and engineers who plan, or design, these spaceships. There are two basic sections, the rocket and the pod at its tip. The rocket has to go up, abort if anything is wrong, depositing the pod safely by parachute, or, if the launching goes well, shut off and detach itself from the pod at the proper moment, allowing the pod to hurtle like a cannonball into space. The "great designers" are also responsible for the ground control system that monitors the sensors in the rocket and the pod and intervenes if anything goes haywire. Within the rocket and pod they have put automatic systems that come on or shut off without anyone doing anything.

Next think of the ground controllers, or middle management—dozens of them sitting in front of keyboards and television screens with headsets on. They run the control and monitoring system that the designers have created. Finally, there is Ham, or Alan Shepard as the case may be, in the pod. Ham does not interfere in the system at all; Shepard is allowed to play with the thrusters that keep the pod from tumbling or turning, after the capsule is free of the rocket, or before that, to punch the abort button if there is an emergency. The abort button too could be automated. The hierarchy was clear: designers, controllers, and subjects. Project Mercury was supposed to be a scientific enterprise; astronauts were part of the test.

Actually, for the first suborbital flights, this was not inappropriate, since all that was involved was firing a capsule into the air like a mortar shell and having it come down in the right place. It was a bit more complex for the orbital flights, and immensely more complex for the moon missions. Still, the planet flybys—Jupiter, Saturn, Mars—were without on-board operators, and the middle managers at Mission Control did very well indeed without them. Did we need operators for the first orbital flights and the moon missions? They might have been nice to have there to handle emergencies, but they were expensive. Much of the

complexity of the system that could create emergencies was the result of habitat and retrieval requirements once humans were aboard. Though the human subjects were not expensive in labor terms (they lived on their military salaries of around \$200 a month and some perks), they were expensive to keep alive up there and to recover in one piece.

The question of their necessity is not easily answered. If we were only curious as to what happens to someone in a pod in space, we went to a lot of trouble to find out. If humans were, as one conference of scientists and engineers put it, "added to the system as a redundant component," they were still very expensive ESDs.⁹ The unmanned probes of far space got along fine without them. If humans were to be significant nodes in the loop of design, equipment, supplies, and ground control, exercising some judgment, doing a bit of piloting in a "spacecraft" rather than a "capsule," there were no signs of that plan in the first flights. It is not clear that any policy or consensus existed on the matter. What Wolfe suggests is that in selling the space program to the public, the public response caused that decision to be made. The first seven astronauts were instant heroes after they were selected. They were to ride on top of rockets that were always blowing up, and come down in a large can in the ocean, and to do that, they must be heroes. When the seven found they were instant heroes, they began trying to dictate or at least influence the terms of their node in the loop. As we shall see, they demanded some modification in the pod and some control over it.

The importance of the question goes to the organizational heart of all high-risk systems: If they are risky, to operators and other potential first-party victims, as well as to the immensely expensive investments, why not eliminate the operators? A package could be landed on the moon and ordered to send back pictures, perhaps even samples. Some argued this would be safer and cheaper. Operators are unreliable, as well as alive. Most people who think about these things, as we have seen at length, feel that operator error account for 50 to 90 percent of the failures in complex systems. If the environment is as hostile as outer space, and the error rate of operators so great, why risk the operator of the system?

But, what if the designer errs, or the environment is uncooperative, or the builder leaves out a valve ring or makes a faulty switch? Can the Middle Managers in control recover? The evidence suggests that the answer is no, because the goal is not to explore space by shooting sensors into it. If that were all, we could get by with fly-by missions or soft landings of suitcase laboratories. But if we were to occupy outer space, control it, police it, use it to spy on the Soviets and shoot thunderbolts at them, that would require humans. Furthermore, we might even mine

space someday, or move there. Any of these possibilities requires more than a recording cannonball. The national pride of landing an American on the moon also played its part.

The military-political aspect was there from the start and probably overriding. After the assassination of President Kennedy, the success of the Apollo project to put a "man" on the moon by 1969 was resoundingly reaffirmed by the new president. "I don't want to go to sleep by a Communist moon," President Johnson said, and that may have been all that was needed to be said. Occupying the moon required more than the moon fly-bys; those basically required good aim and some mid-course corrections, and toggle switches that turned cameras on and off. Even the orbital flights and the moon probes involved frequent failures that only operators aboard the vessels could recover from. Designers, builders, and controllers were not enough, it became clear.

The first astronauts were selected from volunteers who were test pilots. According to Wolfe this was expediency, not foresight. NASA needed a pool of applicants that could be assembled immediately after the Russian *Sputnik* insolently beeped through our heavens in October 1957. Test pilots would already have security clearances, evidence of physical and engineering skills, and could be given immediate orders to report for duty. While radar observers were probably the best trained for the passive skills needed to observe and be observed, test pilots seemed to be a convenient, accessible, and not inappropriate pool. The officials got more than they bargained for, because test pilots were not likely to be passive observers.

Test pilots, especially ones in the military services, had to be extraordinary human beings. Wolfe describes in gripping detail what they went through in their test pilot training. They generally had engineering backgrounds from college, received extensive technical and engineering training in flight schools and in their testing work, were extraordinary pilots, often aces from the most recent war (Korean, in this case, though some went back to World War II), and were unequivocally fearless. They were testing rocket planes and their forerunners, trying to break through the sound barrier of 660 to 760 miles per hour (the figure depends upon altitude, temperature, and the like). They lived dangerously on the ground (more Navy pilots, whether test pilots or not, died in automobile crashes than air crashes) as well as in the air. Most important of all, test pilots were in charge of their craft. Wolfe describes it thus early in his book:

To take off in an F-100 at dawn and cut in the afterburner and hurtle twenty-

five thousand feet up into the sky so suddenly that you felt not like a bird but like a trajectory, yet with full control, full control of Five Tons of thrust, all of which flowed from your will and through your fingertips, with the huge engine right beneath you, so close that it was as if you were riding it bareback, until you leveled out and went supersonic, an event registered on earth by a tremendous cracking boom that shook windows, but up here only by the fact that you now felt utterly free of the earth—to describe it, even to wife, child, near ones and dear ones, seemed impossible.¹⁰

The complexity and coupling of the ships these people flew were intense. The engines were unproven, or were pushed beyond design limits; test pilots' job was to break through the envelope and create and feel out new envelopes, or to "punch holes in the sky," riding on top of a "candle" that eventually had three-quarters as much thrust as the Redstone missile that powered the first suborbital flights. All kinds of things went wrong, and recovery was often impossible. Even after bailing out (if being ejected upwards at 100 miles per hour by an explosive charge can be called that), there was room for the unexpected interaction. The ace Chuck Yeager bailed out only to find himself on fire. His seat was ejected with him by rocket propellant. The rocket propellant was still burning after he separated from the seat. Both he and the burning seat fell through the air. It was close by, above him. When the parachute opened, and slowed up Yeager, the seat passed inside the plastic parachute shrouds, threatening to burn and sever them. This did not happen, but it did collide with him; as his chute slowed, the burning seat crashed into his visor. The burning propellant stuck inside the helmet and set one side of his face on fire. It filled the helmet with smoke and more fire as the seal melted in the oxygen hose, sending pure oxygen into the mask to fuel the flame. Ripping the visor off, he set a glove on fire and burned the flesh inside. He landed near a highway, and a passing youth lent him his knife to cut away the smoldering glove. He survived.¹¹

Instantaneous decisions had to be made in rocket planes when they flamed out at 30,000 feet and tumbled in free falls like a brick, or lost the thrusters that stabilized them, or malfunctioned on takeoff. Eventually test pilots would fly to over 50 miles altitude at speeds up to Mach 7, or seven times the speed of sound, encountering and solving new problems at every step. The military's rocket plane program was cancelled in 1963 when the race to the moon was in full tilt, indeed, on the day that Yeager was burned by his seat after ejecting from a F-104 fighter with a huge rocket bolted to it. But rocket planes surfaced again in the reusable space shuttle that first flew in 1982.

The military test pilots that were accepted for the first space probes

were in for a rude shock. They were to be passive blotters absorbing the new conditions, monitored with electrodes and rectal thermometers, without even a proper window to look out of, only two small portholes on each side. The first American in space (the Russians had long been there; and were to stay ahead in everything for years), Al Shepard, did virtually nothing except watch the panel and be observed through a camera and wires and thermometers. In fact, much of the training of the astronauts consisted of conditioning—getting used to sensations in simulators that created high-gravity forces, learning to keep their hands off the switches, getting used to the sounds of the rocket motors. People with test pilot experience were not likely to easily accept this role. When they got an unexpected chance to act, they took it.

The “lab rats,” as they called themselves during the long period of testing and conditioning at medical and psychological labs, decided to protest and insist on having some control over the flight. Their chances of succeeding were not good; humans are fallible, but it was presumed that designers are not. Even if builders goofed, the designers had redundant components to take care of the problem. The middle managers at ground control would have the most complete and accurate information, and have control over much of the spacecraft, either directly or through the automatic devices the designers had put in. But by now the seven astronauts selected for the Mercury missions were national heroes, and this gave them their opportunity. Their status in the project was rising. They asked for a window to look out of, and got it. They asked for an emergency hatch to get out after splashdown, and got it. (It was installed for the second suborbital flight; before that, other than climbing up the slim neck of the capsule, they had to wait until mechanics on the recovery ship unbolted the hatch from the outside.) They demanded control over the rocket itself, should it malfunction. They did not get this. They demanded complete control of the reentry procedure, over the hydrogen-peroxide thrusters that would control the attitude, pitch, and yaw of the pod in space. They wanted to establish the angle of attack manually during reentry, and fire the retro-rockets themselves, without the use of an automatic system. This one they won partially; they were given a manual override system, but the automatic one stayed.

The Operators Err

The initial flights suggest that pilot control by these grossly overqualified subjects was a mixed blessing as far as the designers and controllers were concerned. As Wolfe put it, not only did our rockets keep blowing up, but our boys kept botching it. Alan Shepard had little to do on the

first short suborbital flight, but before lift-off he had his periscope out, and used the filter because of the sun. He forgot to remove the filter before it retracted (automatically, of course) for lift-off. He would not be able to remove it later, so he tried to remove it while it was retracting, but accidentally brushed the Abort button—fortunately not enough to activate it. He stopped trying to remove the filter, and thus was unable to report sighting various constellations and stars, or observing the orange band at the edge of the earth (events we know would take place because balloons had been up there already), since the filter made it all black and white. This may have contributed to his delay in starting the reentry countdown. Late in starting the countdown, he gave up trying to manually control the attitude of the craft (a task the astronauts had won control over) and put it on automatic. But he forgot to turn off the manual system, and thus precious fuel dissipated from the capsule. In this short flight it did not matter, but it would in a long one.

It mattered soon enough. Scott Crossfield, in the second orbital flight, played around so much with the thrusters on his flight that he ran out of time and power to properly align his module for reentry—a dangerous business, since one could skip back into space forever, or plunge into the atmosphere so sharply as to burn up; a few degrees off and that would be it. Furthermore, the automatic control system would no longer hold the capsule in position for reentry. (It had malfunctioned once before, with an ape, and would malfunction once again with John Glenn). He switched to automatic but forgot to change another switch, and for ten minutes was eating up fuel from both systems. When the time came to fire the retro-rockets, which would slow up the ship, he was behind, and the capsule was not positioned correctly; in addition he was late in hitting the firing switch. He was off by nine degrees in his angle. He had to release his parachute early, and by hand, since the automatic system was out of fuel. He overshot the target by 250 miles, and for forty minutes the impression created on television was that he was dead. Walter Cronkite was crying.¹²

These fabulously trained test pilots, asked to do a few maneuvers and make sure they set some switches properly, were not perfect. The most serious failure of all was in the second suborbital flight, by Gus Grissom. He lost the capsule. It qualifies as a system accident in our terminology, with multiple operator and procedure errors, and has the added interest that a safety feature started it.

Gus Grissom was the second man in space, flying a suborbital lob into the Atlantic near Bermuda just as Shepard had done. This time the “module,” as it was called (it was later to be the “capsule” and then,

finally, the "spacecraft"), had a significant modification demanded by the former test pilots—an emergency escape hatch with explosive bolts to be fired if the module did not right itself in the water and began to leak. There was no need to use it; the flight was perfect and Grissom busied himself with various housekeeping chores after landing, telling the helicopter to wait a few minutes. He may have been nervous; his pulse rate was significantly higher than that of Shepard's throughout the flight. After finishing his chores, he radioed for the helicopter to hook on, saying, as it had been practiced many times, "Okay, latch on, then give me a call and I'll power down and blow the hatch, okay?" The helicopter said yes, "will give you a call when we're ready for you to blow." But as the helicopter swung low with its long shepherd's crook to catch a lifting cable, the hatch blew. Out popped Grissom, since the capsule was filling with water, and he started swimming fast. The helicopter pilot was surprised but not worried; the astronauts had practiced for just such an emergency, and the pressure suit was more buoyant than any life preserver. The astronauts had even seemed to enjoy playing around in the water in what amounted to a large, form-fitted boat. So the helicopter pilot quickly sought to recover the capsule, which was taking on water and sinking.

Grissom swore ever after that he had no explanation for the hatch blowing—it just blew. Subsequently, extensive testing under all kinds of conditions in a sister capsule could not make it blow, and explosive bolts had been used for years in fighter aircraft without any mysterious activations. (Still, there were all those contractor errors we cited early in this chapter.) After Grissom had armed the explosive charges, as required, the button could have been pushed in error or inadvertently while he was moving about in the capsule.

Grissom exited into the ocean and started to swim. But it just so happened that since he had landed safely and was about to be picked up and carried to the carrier deck by the helicopter, he had unplugged the oxygen inlet to his suit. Thus, his form-fitted boat had a large hole in it, and it was some time before he realized that instead of bobbing merrily, he was sinking. He was less buoyant than normal anyway, because the pocket of his suit was filled with souvenirs, such as two rolls of dimes and models of the capsule. He managed to reach down and close the inlet valve, but took a lot of water into his lungs in the process. In the choppy sea, the spacesuit with its water pulled him under as waves came by. Still, there was nothing to worry about because there was a second helicopter up there in case there was any trouble. It was watching him; in fact, he could see someone taking moving pictures of him. He waved frantically, and

they waved back, assuming he was jubilant about being the second man in space. It had all been practiced so many times it was familiar. The helicopter hung there, as Gus swallowed more and more water and struggled to keep his head up, because the pilot was keeping watch over the other helicopter that was being dragged into the ocean by the sinking capsule. Aboard that other helicopter the warning lights were lighting up the control panel as the ship tried to carry 25 percent more than its rated lifting capacity, and finally the pilot cut the capsule loose; the helicopter's wheels were under water and the engine might blow. Then the back-up helicopter turned to Gus and he was hauled, incoherent, aboard. He remained incoherent for some time, furiously grabbing at life preservers in the cabin of the helicopter.¹³

The failures were trivial—accidental blowing of the safety hatch; open oxygen inlet valve; the preoccupation of the first helicopter with the capsule; the misreading of signals from Gus to the second helicopter; perhaps an overweight suit; a choppy sea. For any single eventuality there is an Emergency Safety Device, a redundancy, a back-up, a planned bit of overcapacity. But when they came together, we almost lost the astronaut. Had he not been there, of course, all would have gone fine; in fact, it would have been fine if the "safety" hatch had not been installed.

The Designers Err

But operator error is only part of the picture. There is also design or equipment error. As with any new, complicated system, even those subject to the most rigorous testing possible, there were bound to be failures. In the first suborbital flight with a living being aboard—the chimp named Ham—the rocket climbed at a slightly higher pitch than planned, and it came down at the same angle and missed the target by 132 miles. Ham almost died, since water seeped in during the two hours of bobbing in the ocean.¹⁴ In the first orbital flight with a chimp, the electrical system malfunctioned and the flight had to be cut short. The automatic control system could not keep the capsule positioned properly, and used so much fuel that control feared they could not orient it properly for reentry (a pilot could have, it seems).¹⁵ They brought the capsule down off the California coast, rather than the Florida one. The designers were doing no better than the hot test pilots.

The automatic control system also malfunctioned in John Glenn's first flight, and the indicators of the angle of the capsule with respect to the earth were off. He took over control, and eventually made a dramatic entry using visual sighting to position the capsule for the most critical part of the trip.¹⁶ There also was a spurious warning that control picked up

that a "landing pack" attached to the capsule might have been lost. To Glenn's outrage, they kept this from him for hours, asking him only to check this and that. They would not tell him, the pilot now, and not a blotter for sensations, what was wrong, only what to do. Though he had to override the automatic system for positioning, reenter by visually lining up the horizon, and had virtually no fuel left for corrections, he was still not treated as the pilot. Ground control was in charge, and would remain in charge throughout the space program.

Rationalization of a system replaced operators; it is the province of designers and managers. This process has a long history. But rationalization allows more complexly interactive systems to be built, and the failure manager, the operator, comes back in. This too has a long history, but the few years of the space program illustrate both tendencies. Initially, radar operators who would act as redundant sensors were all that was thought to be required. The first shots made it clear that more skill was needed, and in addition, the lionization of the astronauts gave them something akin to union power. The Gemini and Apollo programs made increasing use of the operator; we went from hollow cannonballs fired into orbit to spacecraft, still with manager control, but more reliance upon operators.

With the space shuttle, the operator has still more control. An automatic landing system there, already in place, must be monitored, and the need for override is expected to be frequent. Indeed, the operator's role is likely to be larger in the shuttle because the designers can't really figure it all out. The shuttle is designed with more redundancy than any previous spaceship, according to one journalist.¹⁷ But this very fact, he notes, makes it more complicated than any other, and makes the spaceship's workings more mysterious and unpredictable, even to those who designed it. The director of flight operations is quoted as saying that the magnificent "architecture" of the craft makes it all the harder to learn to use the system. The numerous bugs in the first space shuttle flights testified to that.

We should not overemphasize the role of the operator. The rocket and spacecraft were full of automatic controls that in no conceivable way could be handled by a pilot, just as in nuclear and chemical plants and aboard aircraft. The operator must come last, the designer first. Managers generally have more information than operators and a wider purview. But despite their admitted contribution to error, operators make a contribution to recovery from other's errors or unexpected environmental conditions. The more linear the system, the less important this is, and the distant probes are more linear systems that moon- or earth-orbiting

missions with landings. The more complex the system, the more important this vulnerable but exceedingly flexible component, operator, is. We will see that dramatically in the case of *Apollo 13*.

Apollo 13

In April 1970, the fruition of ten years and 10 billion dollars was at hand: *Apollo 13*, defying superstition by its very designation, was launched into a trajectory to the moon, equipped with a landing module for a moon landing. Two days into the flight, as they neared the point of no return, halfway to the moon, trouble developed. For the next four days, with full coverage from Houston Mission Control, the specter of three daring astronauts possibly becoming stranded in space and slowly asphyxiating gripped this and other nations. The recovery from the accident was truly heroic, and a technological marvel. The accident's origins, of course, were mundane. Our source here is Henry S. F. Cooper's intense and dramatic book *Thirteen: The Flight That Failed*—technological writing at its peak.¹⁸

Two weeks before the lift-off, a routine prelaunch test was made of the craft, including a test of the two oxygen tanks. These are large nickel-steel alloy spheres designed to withstand 900 pounds of pressure per square inch. They carried liquid oxygen, which is kept at an extremely low temperature, -297°F . Liquid oxygen is unstable, and the tanks have a capped dome with pipes and electrical wires going into it to run fans for stirring the viscous liquid, heaters to expand it and force it out if it doesn't flow freely, and all the associated gauges and switches. The two tanks, along with the hydrogen tanks, contain the elements for fuel cells that generate the electricity to drive virtually all the key systems in the spacecraft—guidance propulsion, power for communication, power for removing carbon dioxide, air to breathe, and so on. At the end of the test the engineers had difficulty getting the liquid oxygen out of the tank. So they turned on the heaters and fans in the tanks, and left them on. There is, of course, a thermostat for the heater, designed to turn it off once the temperature reached 80°F . But the thermostat is designed to work on a 28-volt power supply utilized in the spacecraft; a 65-volt power supply was being used in the test. This would make no difference for most of the tests. But the thermostatic switch was designed to operate at low temperatures, as well as on low power. The high voltage probably prevented it from shutting the heater off at 80° , and the subsequent heat (estimated to have approached $1,000^{\circ}\text{F}$.) not only fused the switch, but burned all the insulation from the wires after they entered the tank.¹⁹ The system was unplugged after a few hours' time without notice of the high heat.

There would be no sign of the failure of the switch or the insulation, and no occasion to check it; the tank had already passed inspection.

There are three modules involved in the spacecraft: the service module, where the tanks and other equipment and supplies reside; the command module where the crew resides and where the controls are; and the lunar module, a small, disposable craft used only to go a short distance from the spaceship to the moon's surface and back again. When the spacecraft was part way to the moon, John Swigert, the command module pilot, had trouble with the two hydrogen tanks; a caution light went on in the command module and, of course, at ground control. He also had difficulty with the quantity gauge for one of the oxygen tanks (the one that had been damaged during the testing). It is possible that the trouble might have been unrelated to the damage; the tanks are tricky, and that is why there are fans inside to stir up the liquid and heaters to build up pressure to force liquid oxygen out. Mission Control ordered the fans turned on so that the quantity, or level, gauge would function correctly. An electrical arc occurred because of the uncovered wires, and this created heat.

A short like this is conceivable, so there was a caution light warning of a short on the instrument panel (along with 250 other gauges). Since the hydrogen and oxygen tanks both empty out into the same place, the fuel cells, the system was designed in such a way that the caution light for the hydrogen system preempted that of the oxygen system. One does not need both, so to speak. But since the hydrogen one came on first, that is what they focused upon, and there may have been a problem there—but again, probably not, since false alarms are common in these highly interactive and tightly coupled systems.

In time the arc heated the oxygen tank sufficiently to blow off the cap. Other materials in the service module caught fire from the arc and, finally, one of the six bay covers (hatches) on the module blew off. The ship had lost its major source of oxygen—and support of life.

Though the ship was two days from earth, the explosion was seen on the earth. Next door to ground control an experiment was going on which involved tracking the spacecraft with a radio telescope. The image of the craft—just a dot by now even on the powerful telescope—was displayed on a television screen. The experiment was ending; the ship was too far away by now, and there were a lot of blips and other interference that registered on the screen. Then they noticed a bright spot on the monitor which grew to the size of a dime and then disappeared. It was the venting oxygen reflecting light from space. They attributed it to a problem with the monitor, and in any case, they had no operational

connection with the ground controllers next door. Had the explosion occurred several hours earlier, it is conceivable they might have interpreted it as a malfunction of the craft and inquired, but probably not.

The explosion was also felt aboard the craft. There was a small jolt, and a distinct bang—not the sort of thing you expect in a voiceless void. Swigert attributed it to one of the other astronauts bumping into the hatch as he came in from the lunar module, but he wasn't sure. Mission Control got a brief interruption of telemetry, but that was not uncommon. Besides, ground control was now preoccupied with another problem; in addition to the hydrogen caution light, a warning light appeared, signalling trouble with the electrical system. The astronauts saw it too, of course, and we know they were worried because their pulses, routinely monitored by management, went from around 70 to 130. Management's pulses are not monitored, so we do not know what they felt, but they made light of the problem and were slow to be convinced that anything was really wrong.

As noted, the hydrogen and oxygen tanks fed a reaction to three fuel cells; these provided electricity for running radios, controls, et cetera, and also provided water for the system—for cooling essential parts, and for drinking. Finally, they provided oxygen for breathing. Only one of the two oxygen tanks exploded, but the second one was leaking; a pipe had been damaged by the explosion. In addition, there is a common-mode connection at the point of the fuel cells, such that the oxygen in the undamaged tank could find its way out through this path. Vital life-support systems were draining away; the redundancy of two tanks was defeated by design.

Aside from losing the source of air and the means of removing carbon dioxide, the loss of electrical power meant that the astronauts could not control the attitude of the spacecraft; that attitude was now upset by the oxygen venting from the service module. The spacecraft had to rotate every twenty minutes to even out the burning temperature on the sun side, and the supercooled temperature on the shade side. It wasn't doing this. In addition, the erratic motion threatened to put the spacecraft's inertial measuring unit used for navigation in a particular attitude or position known as "gimbel lock," from which it could not be moved, and which would render the guidance system ineffective. It would be like losing your compass. Finally, the motion interrupted communication with control—it was difficult for control to know which of the four antennas to tune to, as the spacecraft rotated erratically. For part of the time there was simply no communication.

For seventeen minutes no one could conceive of the source of the

increasing number of failures that were occurring. One cannot see the service module, or inspect it. There is no gauge that reads, "Oxygen tank explosion," only gauges of pressure and temperature and quantity. The warning light for the oxygen tank was not even on, since the hydrogen one was on and preempted it. As with Three Mile Island, there was at least one gauge that might have pointed to the problem—the oxygen quantity gauge—but there was no gauge to point to its significance. The astronauts noted the gauge, but disregarded it. The warning signals pointed to the hydrogen tanks and the electrical system, and they knew that the radios kept going out, the roll could not be controlled, two of the three fuel cells were losing power but not the third, and so on. None of these signals made any cumulative sense, but it seemed the problem was clearly electrical.

The search for the problem was conducted with the well-worn assumption we have been exploring in this book: since the system is safe, or I wouldn't be here, it must be a minor problem, or the lesser of two possible evils. Just as reactor cores had never been uncovered before, or that ship would never turn this way, or they would never set my course to hit a mountain, the idea that the heart of the spaceship might be broken was inconceivable, particularly for the managers and designers in Houston, all gathered together for the historic flight. Cooper puts it well:

...they felt secure in the knowledge that the spacecraft was as safe a machine for flying to the moon as it was possible to devise. Obviously, men would not be sent into space in anything less, and inasmuch as men were being sent into space, the pressure around NASA to have confidence in the spacecraft was enormous. Everyone placed particular faith in the spacecraft's redundancy: there were two or more of almost everything.²⁰

As with Three Mile Island, the extensive training was for straightforward problems of component failures, and not for inexplicable interactions of three failures (the problem with the hydrogen tanks, occurring just when there was an indication that the oxygen needed stirring, and a design that masked the subsequent failure of the oxygen tank). During the course of pre-flight simulations the ground controllers, who themselves rehearsed the flight many times since they would essentially control it, grew impatient with the sometimes bizarre and implausible accident scenarios that were programmed into the simulator. The previous Apollo missions had rather minor and straightforward single component failures, while the scenarios they were now being given sometimes seemed ridiculous. By the time of *Apollo 13* they had won their point, and the astronauts had joined them in the effort to restrict trials to "one-

point" or "two-point" failures. "Four-point failures" or "way-out disasters" (what we would call system accidents) had been dropped from the simulation runs.²¹ It may not have made any difference; the idea that a short would occur in an oxygen tank and not be signalled was probably inconceivable even to those who designed the bizarre simulations. The chief pilot, Swigert, said afterward that "nobody thought the spacecraft would lose two fuel cells and two oxygen tanks. It couldn't happen. If somebody had thrown that at us in the simulator, we'd have said, 'Come on, you're not being realistic.'"²²

Faced with the increasing number of malfunctions, the managers set about following standard operating procedures. The most likely explanation was a telemetry failure. False readings from the gauges aboard the spacecraft were to be expected, and had occurred in the past. They proceeded to check instruments, or have the astronauts check them. Then the next most likely event would be a minor failure or problem with one or more parts of units. This was carefully checked out. Even after it was clear (first to the operators, note, then finally to the managers) that there had been an explosion, the managers hoped that the mission might still go ahead and there could be a moon landing. It was not until about forty-seven minutes after the accident, twenty-eight minutes after they knew something quite serious was wrong, that ground control completely gave up on the notion that a lunar landing might somehow still be possible. This hope, or production pressure, had operational significance: the managers were reluctant to take any action to contain the increasing damage, or to test the system to see what damage had been done, if that action would mean an end to the possibility of a lunar landing. It might be called, after Tom Wolfe, "Houston Right Stuff."

Diagnosis

The actual discovery of the failures exhibits further parallels with the system accidents we have been covering in this book, giving us deeper insight into the mental models that operators (astronauts or controllers) must construct. When the warning light came on indicating electrical system problems, the managers found that one of the two main electrical buses fed by the fuel cells was suffering a significant loss of power. (A bus is a kind of junction point where electricity gets distributed; it can be shut off, breaking the circuits.) The equipment powered by this fuel cell—approximately half of the command module's equipment—was also beginning to fail. But the warning lights in the capsule and on the ground went off, signifying normal operation. Both the astronauts and the controllers concluded a minor malfunction had cleared up; the astro-

nauts' report of a jolt and bang receded from consciousness. Then another anomaly appeared in the bandwidth of the radio waves carrying information from the capsule. The ground controllers surmised that if there had been a temporary minor electrical problem in the bus that serviced the antenna, this must be a related problem. Now that the electrical problem seemed to be solving itself, communications would clear up momentarily as well. In fact, the radio disturbance was caused by the panel from the service module flying off into space and striking the antenna, but no one dreamed of such an explanation at this point.

Still within a minute of the accident, the astronauts had noticed that the quantity gauge for the oxygen tank had gone off the scale on the high side, but no one was drawn to this source as the cause of the problem. After all, they had been having trouble with the oxygen gauges right along; that was why the oxygen tank had been ordered stirred in the first place.

But now the other main bus was beginning to fail; ground control was beginning to get confused. Half the warning lights on the relevant control monitor were lit up. It was still hard to know where to begin to try to isolate the electrical problem. Two of the three fuel cells were also discovered now to be without any power whatsoever. Since the fuel cell that was still functioning drew its oxygen from the same tanks as the dead fuel cells, ground control saw no reason to suspect the oxygen tanks. To preserve the functioning of the equipment aboard the command module, reserve storage batteries that were designed to provide power during re-entry were hooked up. Eventually, as the scope of the disaster dawned on everyone, these batteries were disconnected to preserve power for re-entry to earth, but at this early stage there was still hope of going forward with the mission.

Finally, ground control ordered the astronauts to begin reading out loud all the gauges in the command module that had anything to do with the electrical system. When they got to the gauge for the pressure in the oxygen tanks, they discovered that the pressure gauge was reading zero, though the quantity gauge had gone off scale on the high side. One of the astronauts in desperation went to a window and saw a cloud of vapor venting from the service module. The astronauts now understood why they had been having trouble controlling the attitude of the spacecraft, and with the help of the pressure gauge reading, they finally had a grasp of the problem. Even though the astronauts reported their observations immediately, the managers went through one more check to determine whether all of the malfunctions including, presumably, the visual sighting of vapor, might still not stem from an instrumentation problem. The

astronauts understood the dimensions of their problem thirteen minutes into the accident. It was another four minutes before ground control finally agreed and ordered the crew to start turning equipment off in the command module to ease pressure on what remained of the electrical system. Even at this late hour, however, ground control was still trying to keep its options open for completing a lunar landing. They carefully avoided turning anything off that would definitely end all possibility of completing the mission.

This difference between the analysis of the operators and management is revealing. The complexity of the system was more baffling to the controllers, even though they possessed more information than the astronauts, because their monitoring of it was all somewhat indirect. The astronauts *felt* a jolt and this remained an important part of their analysis throughout the process of trying to track down the problem. They were much more inclined to take the problem seriously, less inclined to look for an instrumentation failure. The astronauts finally looked out the window and *saw* gas venting from the service module. This left no doubt in their minds, but ground control went through one last attempt to find instrumentation failure that would mean production would continue. The inability to see and feel the system directly was a severe handicap for ground controllers despite the great sophistication of their electronic monitoring system. The only ones on the ground to actually see the failure, through a radio telescope hooked to a crude television monitor, did not know what they saw, and were not monitoring the operations.

The accident allows us to review some typical behavior associated with system accidents: (1) initial incomprehension about what was indeed failing; (2) failures are hidden and even masked; (3) a search for a *de minimus* explanation, since a *de maximus* one is inconceivable; (4) an attempt to maintain production if at all possible; (5) mistrust of instruments, since they are known to fail; (6) overconfidence in ESDs and redundancies, based upon normal experience of smooth operation in the past; (7) ambiguous information is interpreted in a manner to confirm initial (*de minimus*) hypotheses; (8) tremendous time constraints, in this case involving not only the propagation of failures, but the expending of vital consumables; and (9) invariant sequences, such as the decision to turn off a subsystem that could not be restarted. All this did not just take place with a few high-school graduates with some drilling in reactor procedures, or a crusty old sea captain isolated in his absolute authority, but happened with three brilliant and extremely well-trained test pilots and a gaggle of managers (scientists and engineers all) backed up by the "Great Designers" themselves, all working shifts in Houston and wired to the

spacecraft. As great as all these folks were—and we shall see shortly how truly great they were—the complexity and the coupling of the system defeated the mission, and nearly led to the “final accident,” the loss of the ship in space.

Recovery

The resources at hand in NASA for recovery did make a great deal of difference. Though the damaged equipment and dwindling supplies could not be replaced, nor the severing of subsystems undone even if they wanted to do so, and even though the subsystem failures were final, the system itself was partially intact, and an unexpected redundancy, one not designed-in but fortuitously available, existed—the little, disposable lunar module.

NASA had four complete teams, each with dozens of experts, available to staff its ground control system on a 24-hour basis. They were all trained in the system, familiar with it, and indeed some had helped design it. Such resources are not available to any other high-risk system in our society. NASA was able to free up about forty experts to concentrate on working out solutions to get the astronauts home safely; they were free of routine flight management duties. In addition, almost every step they devised could be quickly and realistically tested in a very sophisticated simulator before it was tried out in the capsule—something rarely available to high-risk systems. With these resources, they were able to work out completely new and heretofore unimagined procedures for recovery.

NASA had a completely dead command module with no regular oxygen supply and no regular electrical system halfway to the moon. An emergency oxygen supply of limited capacity was still available, along with batteries for reentry maneuvers, for powering up the command capsule for the last hours of reentry. But by now the ship had passed the point of no return and would have to spend days traveling out to the moon and using the moon's gravity to swing it around and head back to earth. The three astronauts would have to spend four days living in the lunar module, designed for two persons and for just a brief jump from the command module to the moon and back. The lunar module would have to establish trajectories and maintain attitude control for the whole spacecraft, while keeping the operators alive. Then, just before reentry, the lunar module would be discarded, since it had no heat shield and would burn up upon entering the atmosphere. The operators would have to return to the command module, bring it back to life after four days of

dormancy, and have enough fuel and oxygen and electrical power to manage the difficult reentry problem.

All the designed-in redundancies of the spacecraft now were ridiculously limited. The limited oxygen supply in the lunar module would now have to last three people almost four days. Water would have to be rigidly conserved (and dehydration eventually led to reentry errors). The battery-generated electrical power was severely limited (that is one reason they transferred to the lunar module; it would use less). The lunar module did not have enough carbon dioxide filter canisters to cleanse the air for so long a period, but the command module canisters did not fit the lunar equipment (and who would have thought that this would ever be needed). Eventually, plastic bags for waste and electrical tape made a reasonable connection possible. The electrical system was designed so that the command module would support the lunar module; now it had to be reversed somehow so that the lunar module could power up the command module. No one knew if it would be possible. These were just some of the inconceivable problems facing the designers and managers and operators.

The operators of a spaceship work from checklists, just as ordinary humans do when they go Christmas shopping, move their households, pack for a ski trip, or study for a final exam. The checklists utilized by the operators of spaceships, however, take a team of engineers about three months to prepare, and could take three or more hours just to read off to the operators. It is not a matter of simply looking to see if switches are in the correct position; it involves entering values in the computers, running tests and subroutines, setting and resetting valves and switches. It configures the craft through a series of sequences. Faced with an unprecedented situation, teams of engineers had to create, in hours, new checklists—piloting directions for the operators. For example, all the checklists for powering up the lunar module were based on the assumption that it would be receiving power from the command module. Finally, a checklist was found that involved starting up the lunar module from its own batteries, but it was long and complex and would require two hours for the astronauts to complete. There was only fifteen minutes of power left in the command module, not two hours. A checklist was improvised on the spot based on the designers' and managers' intimate knowledge of the detailed workings of the system.

In a sense, what was involved was the design of a substantially new system based upon the components of an old and no longer functioning system. And along with creating a new system, a whole new set of proce-

dures for operating it also had to be devised. Because of the uncertainties of such a project, in one sense the new system would be more complex than the original one. While it would involve interactions that had never been contemplated before and were therefore not well understood, it would also require a drastic simplification of the system. Complexity was reduced as well as safety devices. The moon-launching part of the mission, scientific experiments, television coverage, on board comforts and even automated devices were discarded, leaving them with a more single purpose, linear system. Some of this loosened the coupling of the system, but the complexity was greatly increased by the lack of redundancies and slack. Everything had to be conserved. To save fuel they continued the trajectory to the moon so that its gravitational force could be used to reverse the direction of the craft. This also gave controllers time to drastically simplify re-entry procedures and check lists, because the astronauts became error-prone with the lack of sleep, water and heat. At each point, controllers chose the option that gave them the most opportunities for intervening by themselves and for the operators to make corrections at subsequent stages. Thus, at each point decisions were made that would loosen the coupling of the system. As the flight approached the stage of reentry, some of the pressure on the consumables began to ease as reserves being held back for various contingencies were no longer necessary as crisis points were passed. So, as the flight got closer to home, some redundancy in the consumables reappeared, and the system, in effect, became less tightly coupled than it had been at earlier points in the recovery process.

The new system created by the accident also suddenly had the potential for a catastrophic accident that the old system did not have. Inside the lunar module there was a container of radioactive fuel, including some plutonium, for powering an experiment that was to be left on the moon. Now it was coming back home in the little lunar module. A similar canister in an earlier space shot had burned up in the atmosphere and scattered plutonium in the upper atmosphere. This canister had been designed to remain intact through reentry, but it might burst upon hitting the ground. In calculating where the command module would eventually be recovered, the controllers also had to calculate where this radioactive fuel canister would hit after the lunar module burned up in reentry. One early trajectory for the recovery of the command module would have brought the radioactive fuel down in a heavily populated area of Madagascar. The representatives of the Atomic Energy Commission at NASA lobbied against that course. (It is not clear what they would have said if it had been a lightly populated area of Madagascar.) In

the end, the canister appeared headed for deep water off the coast of New Zealand. But the astronauts made a last minute error in attitude control; it would have been difficult and dangerous to try to correct. Fortunately, the new location for the casket was the Indian Ocean, and the AEC representatives were satisfied.

As it turned out, the lunar module was made into a much better "life-raft" than anyone had expected it to be, and it dragged the dead command module around the moon, back to earth, and successfully powered up the command module for a smooth reentry. It was a job no one ever imagined the lunar module would have to perform. It is an example of indigenous safety devices unplanned by designers, but just there, that are common in linear, loosely coupled systems but not in complexly interactive, tightly coupled ones.

The recovery from the failure was brilliant. Cooper captures it just as brilliantly in his long, detailed, and detached narrative of this phase of the mission. It was a technological triumph, but almost an inverted one. Cooper plays with this idea in the following:

The supreme achievement of American technology had broken down utterly. All that was left was a spacecraft whose very complexity made it harder to handle, plus a group of flight controllers and three astronauts who were themselves products of the vast bureaucratic machine that had produced the malfunctioning spacecraft. On the face of it, this might appear to have made it all the more difficult for them to get outside the situation and impose their will on the wayward spacecraft. However, the accident had also demolished most of the technological appurtenances, such as checklists and flight plans, which substitute a sort of delayed time for immediacy, and also much of the automatic equipment aboard the spacecraft which performed tasks that earlier mariners would have performed for themselves. Now the flight controllers and the astronauts were no different from any other sailors facing disaster at sea. They would do a lot better by themselves than their elaborate paraphernalia had done by them.²³

He is correct when he links the controllers and the astronauts together as mariners; in the recovery the distinction between manager and operator, indeed, between designer, manager, and operator tended to disappear. With *Apollo 13*, more than any other flight, the astronauts returned to their old test pilot status—professionals in a cooperative, extraordinary endeavor that made little of rank and control. The conditioned subject, the chimp, the redundant device, the man who had only to sit on top of a piece of fireworks with a thermometer up his rear to achieve hero status, disappeared forever. But so, hopefully, has the notion of the "Omniscient Designer," and the "Omnipotent Manager."

Early Warning Systems

Since Hiroshima and Nagasaki the ultimate in catastrophic accidents has been accidental nuclear war, an accidental determination of the fate of the earth, a holocaust. Headlines tell us of steady technological advances in missiles and the defense against them. The defense is largely more missiles or faster ones or quicker counterstrikes. Missile gaps continue to be debated; the 1960 presidential election was thought by many to be decided on the basis of a missile gap when missiles were few in number and overwhelmingly in our hands. As the gap has narrowed, the number of missiles has increased to thousands. (There are between twenty and thirty thousand U.S. nuclear weapons in the United States and abroad.) Complicating the picture is the increasing number of armed conflicts around the globe, their increasing ferocity and durability, and the increasing number of combatants with access to nuclear weapons. The nightmare of *Dr. Strangelove* has not gone away; in fact, that scenario, bizarre as it was in 1964, is pastoral in its loose coupling and linearity compared to scenarios that can be imagined today.

Dr. Strangelove was a hilarious but chilling film by Stanley Kubrick concerning our Strategic Air Command and national defense posture. In terms of system accidents it runs like this: General Jack D. Ripper (the humor is broad, to say the least) is a crazed SAC bomber commander who launches his squadron of B-52s with their nuclear bombs against the Russians. (Psychopathic behavior has been left out of all the accident scenarios in this book, but we know it exists. We will include it this one time, as failure #1.) Only General Ripper knows the secret code that will recall the bombers, and he will not reveal it. (Failure #2; a procedure error, or procedure risk. It resembles the rumored ability of submarine commanders to fire nuclear missiles without presidential orders under some circumstances, such as, of course, a loss of communication with our president.) The army is forced to attack his base, but (failure #3) the secret recall code dies with General Ripper in the attack (a common-mode failure). But there are still recovery possibilities, because, in those halcyon days, it takes a few hours for the B-52s with their missiles to reach Russia. (It is ten minutes now for submarine missiles, less than eight minutes for Pershing II missiles about to be based in Europe; none of our medium or long-range missiles can be recalled, or even destroyed in flight if fired in error. The system is immensely more tightly coupled today.) A lucky guess and the combined efforts of the Russians and the Americans make it possible to recall or shoot down all but one of the

SAC bombers. But one is left, giving us failure #4. (One is reminded of the reconnaissance airplane that did not get the message during the Cuban missile crisis to keep clear of Soviet territory; it strayed over the border when we were "eyeball to eyeball" with the Russians, considering an air strike against Cuba, and commencing a blockade. It complicated the negotiations no end.²⁴)

The slack in the system had removed all but one potential source of catastrophe, it was thought, but it just so happened that the Russians had installed a "Doomsday" device that would destroy the world automatically if even a single nuclear bomb went off. The "Doomsday" machine was installed to deter nuclear attacks once and for all, but it would only work if we knew of it. Failure #5 was that the Russians had armed it but were waiting for an opportune moment to announce its existence to the world. Once operational, there was no way to disarm it, for its effectiveness depended upon neither the U.S. nor the Soviets being able to defeat it. (The closest counterpart today is "launch-on-warning," a posture threatened by both sides that would send irretrievable missiles to the enemy if it looked as if the enemy were launching an attack. That presumption of an attack will preoccupy us in this section.) At this point in the film, with appropriate patriotic music, no further human intervention is possible. The last remaining B-52 from the squadron, rejecting radio messages to return because the messages lacked the proper code, reaches its target, and the world goes radioactive. In the final scene U.S. officials are left discussing the possibility of preserving at least the U.S. version of civilization underground in abandoned mine shafts. The air force chief of staff warns the president of a possible "mine-shaft gap."

Today, the ridiculous accident scenario is not as ridiculous as it once seemed. With Soviet satellites watching for signs of missile launches, we have had a dropped wrench setting off a Titan missile (fortunately, it went only a few hundred yards), and a rumor of an accidental launching below the Canadian border that landed a missile in Canada. But there is so little information about accidental missile launchings or accidents that came close to causing launchings that I cannot even speculate on the matter.

This is also true of the matter of "broken arrows"—accidents with nuclear weapons, such as dropping them during ground transport or loading, dropping other objects on them, or dropping them accidentally from airplanes as in the Palomares, Spain incident in 1966 (the cost of cleanup operations was well over \$50 million; 5,000 barrels of the plutonium-contaminated soil are buried in lovely South Carolina).²⁵ There have been between "over twenty-seven" (official count) and one hun-

dred twenty-five (International Peace Research Institute count) of these. In the official list, fourteen of the twenty-seven accidents involved the explosion of the detonating device (which in one case created a crater in Texas 35 feet across and 6 feet deep). But the detonating device must go off with extreme precision to explode the warhead; it consists of several explosive charges surrounding the nuclear core. It is thought that it is extremely unlikely that an accident could produce just the right pattern of explosions (indeed, we are not sure an intentional detonation can do it very often). However, we have probably had a few tries at it. In one case, five of the six interlocking safety devices on a bomb failed, and according to Dr. Ralph Lapp, head of the nuclear physics branch of the Office of Naval Research, only one switch prevented a 24-megaton bomb from destroying much of beautiful North Carolina.

One may easily imagine a scenario of commonplace industrial failures and military mishaps leading to such accidents; one need not even invoke the concepts of complexity and coupling. Even though a nuclear explosion is extremely unlikely, a broken arrow can spread deadly plutonium about, and the earth has been scraped in Spain, Greenland, North Carolina, Indiana, and probably elsewhere to remove the deadly contamination. While we no longer keep B-52 bombers aloft with armed nuclear weapons on a continuous basis, they do fly about and the weapons—unarmed—are transported by air routinely in cargo aircraft. (A helicopter carrying nuclear bombs once made a forced landing at Coney Island.) Their plutonium is available for distribution in an accident. While the destructive potential to third- and fourth-party victims is not as great as that entailed by an accidental detonation of one of the thousands of existing warheads, the probability of spreading plutonium about in, say, a 2-mile wide swath 25 miles long, contaminating earth and people alike, is much higher. The most destruction, of course, would come from an accidental launch of nuclear weapons, but that has the lowest probability of all. Let us see why.

We will focus upon the most credible example of accidental attack—false alarms. The matters of terrorism and insane commanders will not be considered.

The Warning

In a huge series of caves carved out of rock underneath Cheyenne Mountain in Colorado, the NORAD (North American Aerospace Defense Command) early warning command center waits for signals that the Russians are coming. It is not, though one might imagine it to be, an uneventful life. While the Russians are reasonably inactive, except for

some test firings of missiles and space shots, everything else seems to be going off all the time. When something does go off, a "missile display conference" is called. In 1979, according to a Senate report, there were 1,544 missile display conferences. In just the first six months of 1980 there were 2,159 of them—over ten a day.²⁶ The enlisted men and women who staff the monitors, thus, are reasonably occupied.

The conferences are not much, though they must keep the managers busy there; they are telephone conferences, calling the duty officers at three other command centers to verify the accuracy of warnings. A "missile threat assessment conference" is more serious, though the telephone is again the medium. In this, persons higher than the duty officers at four separate command posts are called. There were seventy-eight threat assessment conferences in 1979, about one every five days, and presumably two or three times as many in the first six months of 1980. After a missile display conference and then the missile assessment conference we come to the third level of missile conference, one that has never been called, the "missile attack conference." The president of the United States is in on that one.

What keeps these outposts of our defense busy are primarily atmospheric disturbances that produce an infrared "signature" somewhat similar to that produced by an actual missile launch that is detected by our satellites. There are also flocks of birds to contend with, and space shots and testing, as well as miscellaneous anomalies in the atmosphere or the equipment. With so many warnings, the system probably has its responses worked out quite well. This was true on November 9, 1979, when the monitors indicated a massive Soviet attack. It was clear that it was not an anomaly; it showed both land-based missiles and submarine-based missiles were involved. In fact, it closely conformed to what the Pentagon anticipated the Russians would do if they attacked. It was a quiet day, internationally, which may or may not have contributed to the skepticism. The attack showed up simultaneously on monitors in the Colorado headquarters of NORAD, the National Command Center in the Pentagon, Pacific Headquarters in Honolulu, and "elsewhere."²⁷ A thousand Minuteman Intercontinental Ballistic Missiles (ICBMs), capable of hitting targets in Russia, were placed on low-level alert. Ten tactical fighters took off, but what they would have done is not clear to me; they cannot shoot down incoming missiles. But after six minutes the alarm was certified as false. That might seem speedy, but Russian submarine-based missiles will allow us only eight to ten minutes to mobilize our counterstrike, and in addition to arming and preparing to fire the Minuteman missiles the president has to be called, since only he or she can order a missile

launch on such a warning. Over half of this time was used up before the alert was cancelled. But actually, it was suspected to be false within two minutes.

The realistic character of the false alarm was understandable: a training tape that simulated an expected Russian attack had been loaded on an auxiliary computer for routine purposes. Somehow the signal found its way into the active, on-line alert system. (One can visualize the enlisted personnel that monitor the screens saying, "Boy, that looks just like the real thing; just like what they trained us for," and pushing all the buttons in sight. Equally likely, however, would be this reaction: "That's unreal; nothing ever happens the way the training program says it will.") How did the managers know it wasn't real? They checked the reports from two independent sources, the satellites and the early warning radar systems—a matter we shall come to shortly.

Less than a year later, at 2:26 A.M. on June 3, 1980, the Strategic Air Command received an indication that two submarine-launched missiles were on the way; it came from NORAD headquarters. SAC called to confirm, and NORAD was unable to confirm, even though the message came from their computers. The SAC duty officer took the precaution of ordering all B-52 alert crews to board their airplanes and start their engines. It is essential that these bombers, armed with nuclear bombs, be aloft if there is a strike, or the Russian missiles will presumably have wiped out all their bases. It would take the B-52s hours to reach Russia, but they would be above the holocaust in North America. Shortly thereafter the SAC display indicated no missiles, and no other parts of the warning system indicated any either, so the B-52 crews were ordered to shut down their engines.

A few minutes later the SAC display monitor, receiving messages from NORAD headquarters, indicated Soviet land-based missiles on their way to targets in the United States, and a little later the monitors in the Pentagon showed submarine missiles coming in. The monitors do not "picture" these events, as a radar screen does, but only display numbers in the appropriate categories, indicating number of missiles, direction, etcetera. The duty officer at the Pentagon convened a missile display conference, the lowest level of alert involving a conference telephone call, and then went on to the next level, a missile threat assessment conference bringing in more officers. The commander of NORAD said that there was in fact no threat, and one minute later the SAC alert was terminated. In the meantime, the Pacific Command airborne command post actually sent bombers into the air from its base in Honolulu. The whole episode lasted three minutes. Since the Soviets have the capacity

of monitoring the activity of our airbases and may have the capacity of monitoring telecommunication activity, if not the actual content of the messages, one wonders if they went on alert with their own missile crews and bombers.

The cause of the false alarm could not be determined. But NORAD knew it was false because neither the satellites nor the radar picked up signals of land or submarine-based missiles being launched or penetrating our perimeters. NORAD kept the system in the same configuration for the next few days to see if it would reappear, or so they testified in Congress.²⁸ Three days later the identical alarm recurred. SAC crews again started their engines, and again the alarm was determined to be false in less than three minutes. NORAD switched to a backup computer, and began the search for a possible malfunction. Eventually they found it. It was a defective, tiny silicon computer chip (cost, 46 cents), not in the computer itself that was set aside, but in the "multiplexer" that routes messages to the various command posts. This multiplexer sends a message to the command posts on a continuous basis to confirm that the channel is open and usable. However, this message was in the same format as that used to indicate a real attack. Why it was the same format is not clear. It seems unlikely that this was an inadvertent similarity, though that is possible; more likely, the form of the message was similar to insure that that form could be sent. The message has a space to indicate the number of missiles, the number given is zero in the routine, test message—we have no missiles today. Due to the malfunction of the chip, this zero changed to a "2" and then apparently to other numbers; it sent this data out to some, but not all of the command posts. The available testimony does not indicate what information other than "2s" for the number of missiles was also sent.

NORAD changed the routine message format so that it no longer resembled an indication of an actual attack. It also corrected another oversight, resembling the PORV warning light at the Three Mile Island plant. NORAD had been aware of what it told its equipment to send out, of course, but it had no monitors at its headquarters to show what actually was sent out. (At TMI, the operators were aware of what the automatic system had told the valve to do, but had no way of knowing what it actually did.) NORAD operators had no way of knowing that, without intending it, the signal was saying, "We have two missiles for you today." NORAD installed monitors, and presumably three new shifts of enlisted men and women now watch these and compare them with the signal that is supposed to go out. Now that this little anomaly has been discovered and corrected, one wonders what the next one will be that no

one thought of, but was so obvious once it was discovered. The opportunities seem endless since NORAD is such a vast and complicated command, control, and communication system resting on radar and satellite stations that apparently transmit ten pieces of garbage—false warnings—a day.

Nevertheless, false information is not dangerous if it is not credible. Our early warning system has two major checks upon the credibility of computer errors of this type, and one on the credibility of sightings of missiles and, fortunately, these are reasonably decoupled. The instances we have reviewed were caused by false information generated at NORAD headquarters and sent to the command posts. But each post, and NORAD itself, has its own monitors showing what the radar and the satellite systems are seeing. Checking these, and seeing no indications in either one that bore any relation to the quantitative information from NORAD, the officers knew that it was a false alarm. That is why the president was never called. The system must work very fast, and in these three cases it did—within three minutes or less it was known or strongly suspected that the signals were false.

The two major sensors also detect missiles through different methods, and are not linked to each other. (Satellites pick up the launch, radar the incoming flight, so they are actually not two measures of the same event.) Land-based Soviet ICBMs would be detected first by the Ballistic Missile Early Warning System (BMEWS) strung across northern Canada and Alaska and, closer to impact, by the Perimeter Acquisition Radar Attack Characterization System (PARCS). (The elaborate and uncommunicative names are not reassuring; for me, the names and acronyms only reinforce a notion of unmanageable complexity.) PARCS is near Grand Forks, North Dakota, and is reported to be capable of a very precise identification of exactly what kind of missiles are involved and what their "multiple independent reentry vehicles" are aimed at. Of course, the Soviets have not allowed us to test this capability. For submarine-launched missiles off either the Atlantic or Pacific coast, the satellites are the first line of detection, followed by a special radar system called Pave Paws (I will not burden you with what that stands for). For those coming from the Gulf of Mexico, we have an older radar system.

The command centers can check all these. Now it is possible that two or three of these sensors could malfunction simultaneously, or nearly so, even though they are independent of each other and use different detection methods. But it is highly unlikely that the malfunction of the satellite would produce the signature that would just happen to match the malfunction of the BMEWS radar, and that there would also be a very

similar malfunction of the PARCS radar. Not only would they all have to malfunction, but they would have to produce compatible malfunctions. With multiple and independent sources of information, the more detailed the information, the more unlikely an error. Information in industrial plants is generally crude (the valve is open or closed) or singular (the temperature is X). An independent source can mistakenly confirm such values because they are simple. If, however, the information has many parameters—number of missiles, trajectory, speed, and size—then an incorrect confirmation is much less likely.

But we do not know how similar the information must be to be credible. If our anti-submarine warfare satellite spotted a foreign submarine in the Gulf of Mexico, and if a satellite suggested three missiles were fired in the Gulf during a thunderstorm, but it was an ambiguous signal because of the weather, and it just so happened that the radar scanning the Gulf malfunctioned in such a way as to suggest two or five missiles were coming (the malfunction might have been associated with the storm), and if it seemed the missiles would arrive too quickly to wait for PARCS in Grand Forks to verify their arrival, and one of the four command centers had all this information at hand, the duty officer might just send the airplanes up, arm the Minutemen, and call for a second- and third-level alert conference. The three missiles suggested by the satellite might well "fit" with the two or five suggested by the faulty radar, especially since, in an attack, missiles would be fired off in rapid sequence and the number of missiles in the air would change rapidly as seconds passed.

Suppose the Soviet satellites also picked up the highly suggestive atmospheric signal from the Gulf and within three minutes put Cuba on the alert for a nuclear attack from the United States. We would notice that immediately, and it would reinforce the suspicion that Soviet submarines were attacking. We could even suspect that PARCS would not work that well when the chips were down, and thus was not tracking the missiles from the Gulf, and start trying to communicate with our submarines. The Russians would note that, and it would reinforce the partial information they had. They might even see our hardened silo covers roll back in the Midwest. There would be no time for a hot-line call, and we would not expect the Russians to tell the truth anyway, nor they us.

I am sure that military experts could point out many flaws in this hypothetical scenario, and I truly hope they can. But without a security clearance and a year to study the system I cannot be sure that this scenario, or some other, is completely out of the question.

It is also conceivable, though very unlikely, that a malfunction could reverse the path of communication between the sensors and NORAD.

We have seen that NORAD can generate false information on its own. Given the complexity of the system, it is not inconceivable that the failure that produces this might also produce a reverse flow, wherein the false information ended up on the satellite and radar screens through some devious route. After all, a training tape played to a packed house in the on-line NORAD system.

To summarize, the detection system—NORAD and the sensing systems—is moderately complexly interactive. Linearity is introduced because the subsystems are independent of each other, not proximate spatially, nor subject to many common-mode components (though there are some). Some loose coupling is present because recovery is possible from some low-level failures—the B-52s need not take off, or they can be recalled—and final action depends upon finding the president in time and convincing her or him. Some loose coupling is built into the system inadvertently (and inadvertent recovery aids are particularly valuable because they can cover contingencies designers did not think of). The PARCS system was originally designed as part of an anti-ballistic missile system to shoot down incoming missiles aimed at our ICBMs. That turned out to be impossible, but the system is now available as a check on the other sensors. The satellites went up to provide earlier warning than radar, but since the satellites are independent of the radar systems, they provide some check.

But note a final consideration. Each side attempts to complicate the job of the other side. Thus, incoming missiles can fire decoys that might just confuse PARCS and Pave Paws and BMEWS; aware of this possibility, these folks might mistake a confusion of geese and sun spots for the confusion the Russians were trying to spread. Then there is the disturbance produced by nuclear explosions in near space, which would disrupt most of our communications; could a blackout over an interlocked power grid extending from Colorado to the capital be misinterpreted as an electromagnetic pulse from a Russian weapon fired from one of their satellites? The defense system is one that grows on complexity. We are not merely adding safety devices to buffer a failure in the DEPOSE components; the “failures” we must guard against are those that an enemy (supremely clever and resourceful, it is assumed) is actively promoting. Each side mindlessly makes it less possible for the other to rest assured that it does not seek its total destruction. Thus, we confront another source of error that no other system we have studied has to cope with. Ironically, this new source of error is found in that system which could decide the fate of the earth.

The Response

The response system, once the credibility of a warning is thought to have been established, is apparently so complex as to provide considerable safeguard against an accidental war; it is hard to imagine that the system would actually work. Of course, this means that it provides little defense against an actual attack. Since there is really no “defense” any more, only retaliation and a possible mine-shaft gap, this may be a blessing. The World Wide Military Command System, of which NORAD is a part, operates so poorly and intermittently, despite billions of dollars, that we may suspect that its strategic weapons system is as failure-prone as its detection system. An admiral recently remarked that the strategic weapons response system is so complex, because of all its safeguards such as two-person control, multiple-command requirements, and confidences, that he wondered whether we could ever manage to launch our nuclear weapons even in the face of a clear threat.²⁹ The details of the presumed complexity of the system are not available for analysis; one may comfortably assume, however, that it is not a very linear system.

The response system is more tightly coupled than the detection system. A ballistic missile cannot be recalled or destroyed. A submarine commander and his chief officers presumably can fire their missiles if they believe that their lack of communication with their managers is because there has been a Soviet attack. It is probably the case that there is no signal to tell them that the lack of a signal is not a signal—a problem with tightly coupled systems. There is not likely to be any recovery from the consequences of an accidental strike; the Doomsday machine is, in effect, in place. In fact, it is worse; launch on mere warning has been threatened by both sides. It is something the generous Stanley Kubrick did not anticipate.

Conclusion

In conclusion, the early warning system appears to be moderately complex and coupled, but not disastrously so. Since its catastrophic potential is extremely high, we should well wish for further independence among the subsystems and more corroborating possibilities. In particular, we could hope that one aspect of complex systems, the presence of invariant sequences, could be reduced, such as by the ability to disarm the warhead and destroy our own missile in flight. But we should note that in this system, failure to deliver may be high, and since it is certain that false alarms are far more prevalent than true ones, this ineffectiveness may be a virtue. Indeed, one wonders if it is ever worth risking the

high probability that an alarm is false by launching missiles at Russia. The output of such a response is mutual assured destruction, or MAD. The response system is more complexly interactive and more tightly coupled than the warning system (though it is hard to separate the two). Missiles cannot be recalled; submarine commanders may be out of touch but able to act on their own; missiles may go off accidentally. The complexity of the system is such that it may limit its ability to share in the destruction of the earth, but it also means that an inadvertent first strike is possible. Our missiles could go off, not because of a false warning based upon the faulty interpretation of signals from the environment, but because internally-generated signals may go off by themselves.

With weapons system we confront new complexities. First there is the peculiar matter of failures of failures equalling success. If the detection system fails and falsely indicates an attack, but the response system fails to respond, we have success—that is, we have not started a war. One does not like to hope for failures to insure success. But the probability of each of these failures, separately or jointly (greatly reducing the joint probability of both failing) remains higher, I believe, than the probability that the Soviets would make a first strike. We are more likely to fail to respond to a false alarm than the Soviets are likely to intentionally attack.

However, we have a truly interactive system here. We do not know the probability that the Soviets would either accidentally launch an attack, or launch on a false warning. Presumably they are confronted with the same complexity and coupling problems that we have (though there is some belief that their missile system is, as their conventional weapons systems is thought to be, less technologically advanced—that is, less complex and tightly coupled—than ours). The permutations of risk analyses in this area could be an endless hall of mirrors.

Finally, we have, for the first time in this book, a system where the environment is intentional and self-activating. The enemy can intervene in the operation of the DEPOSE components through decoys, deception, disruption of communications, theft of designs, even perhaps the corruption of operators. We need safety devices to guard not only against the failure of a part or unit, but also against the ability to mask that failure. This adds another dimension of complexity to something already complex enough, and greatly limits the possibility of recovering from failures before the system itself fails. There does not seem to be any end to this increase in complexity and coupling. The judgment that, as of 1983, the early warning system is only moderately complex and coupled, may be

overtaken by 1990 because of what we think the Russians are doing to intervene in our system—and what they think we are doing with respect to theirs. As noted before, it is ironic that the most fearsome system on earth should have this added burden of complexity and coupling.

Recombinant DNA Technology

With recombinant DNA processes, we will have to engage in an anticipatory analysis. The series of industries that focus on this technology are only just beginning to come into existence. There are no published descriptions of the production system in commercial DNA labs, and there are certainly no accounts of accidents in these labs as yet. However, we have a fair idea of the nature of the production system, and some people have written at length about the dangers of an accident. The system appears to be complex in its interactions and tightly coupled, but I caution the reader that I know even less about it than I do about nuclear weapon systems.

Recombinant DNA technology, or gene-splicing, is a microbiological technique that allows scientists to graft the genetic information from one organism into the cell nucleus of another. This enables the biologist to design new life forms for accomplishing specific tasks. Among the early successes of DNA have been the creation of bacteria that can produce complicated biological molecules used in medicine, which had been difficult and expensive to obtain. Human insulin, growth hormone, and interferon are among the first of these synthetic biological products that have commercial application. Much more is coming; indeed, whole sectors of the economy may be transformed. In the chemical industry, work is underway to create micro-organisms that can serve as catalysts for chemical reactions, thereby doing away with some conventional transformative processes requiring high temperatures and pressures.³⁰ Forms of bacteria might be created to consume oil spills or transform waste products into energy sources. Some biologists are speculating that agricultural production as we know it may well be transformed through the creation of new and exotic plant forms, and meat production could be transformed through the substitution of "single-cell protein cultures."³¹ Finally, this technology holds out the potential for the application of bioengineering techniques on defective human genes—a technology for

nothing less than designing or redesigning human beings. It is, in short, a process of remarkable potentials that could transform not only our economy but our sense of what is human.

The commercial potential of DNA is apparent in the venture capital flowing into the industry, and the increasing centralization of the industry taking place through cash-rich oil companies. The four original small gene-splicing companies that began with a handful of PhDs and venture capital in the mid- and late 1970s grew to a combined value of over \$225 million by 1979.³² By 1982 it was estimated that there were 350 companies involved in the recombinant procedures.³³ The larger pharmaceutical companies moved into the industry quite early, of course, and were quickly followed by chemical industries and, especially, the large oil companies. Large sums of capital have been committed to the industrial application of gene-splicing by such chemical companies as Dupont, Pfizer, and Monsanto, and among the oil companies, Arco, Standard Oil of Indiana, and Occidental Petroleum. Commercial success in the agricultural market may not come before the 1990s, but may come sooner for the medical field.

Unfortunately, with all these fantastic potentials go some fantastic risks. These industries will produce new, *living* technologies; life forms that are unique, unprecedented, and in some respects very poorly understood. In many of the proposed applications, new organisms will be introduced into the environment in massive quantities. Such quantities may produce totally unexpected interactions; there is simply nothing in our experience to go by. Once introduced, there seems to be very little and quite possibly no chance for intervention should unexpected and untoward interactions take place. As Pamela Lippe of the Friends of the Earth warned Congress in 1977:

DNA is probably the most unforgiving technique we have yet developed. Radiation decays. We can stop making toxic chemicals. But a novel organism has a life of its own; once it has escaped or been released, once it has established an ecological niche, it is out and replicating, perhaps beyond our ability to control or clean up.³⁴

The catastrophic potential of DNA is different from all other systems we have described in that it does not lose toxic or explosive substance upon the environment, but rather creates interactions between systems that were not previously linked at all, and perhaps could not be foreseen to have been linked. Once the linkage is made, it cannot be controlled by the operators. A catastrophe of literally epidemic proportions may be set in motion.

We know something about the interactions of toxic chemicals with the environment, and while these are not new organisms with unpredictable behavior, we can get some idea of the magnitude of the problem by briefly reviewing this relatively simple interaction between the system and its environment. Perhaps the most widely known example is DDT. It was the publication of Rachel Carson's *Silent Spring* in 1962 that first acquainted the public and many scientists with the chain of unexpected effects stemming from the application of DDT, endrin, dieldrin, and other related pesticides. The danger she uncovered was not that of direct poisoning, which was observable and well understood. Rather, Carson pointed out that these poisons "magnify" in living tissues. As they move up a food chain, from plants through small herbivores to a succession of larger and larger carnivores, the poisons become more concentrated in living tissues. The unanticipated interaction here was not a direct poisoning of people but rather that the food chain on which plants, small animals, and man are linked became impaired. The ramifications are still working their way through the planet in, for example, inhibition of photosynthesis in phytoplankton and the crippling of reproductive systems in birds.

The scenarios mentioned by scientists for possible accidents from the production of recombinant DNA are similar to the actual DDT accidents. They include an unexpected linkage between previously unlinked systems, limited prior knowledge of the interaction process, and indirect and delayed information on the consequences. These scenarios are *ecosystem* accidents that result from intended incursions on the ecological system.

Not all eco-accidents are system accidents in their origin. The oil spill in Santa Barbara, the leeching of toxic wastes at Love Canal, and the contamination of Seveso, Italy, by dioxin are component failure accidents. They are the result of poor design, operator failure, or equipment failure. Hooker Chemical Company knew of the danger of the toxic waste they buried at Love Canal. The drug firm, Hoffman-La Roche, in Switzerland was well aware of the danger of dioxin contamination in their plant in Seveso, Italy, and indeed plant officials were instructed to readily reimburse their neighbors for dead farm animals that continued to appear. Knowing that dioxin was a by-product of the pesticide the plant produced, they would not allow production to take place in clean little Switzerland, where their headquarters were, but instead had it produced in dirty northern Italy. When the chemical reactor exploded one weekend when no one was attending it, the safety device protected the plant by allowing the poison to blow up into the air through a stack, from

where it drifted over the neighboring community. Plant officials avoided a panic by simply not informing the community. Components fail in such accidents as these, and the risk of an oil drilling accident, of leaking drums of highly toxic waste, and of exploding chemical reactors are presumably not only anticipated, but carefully calculated risks as well.

With eco-system accidents the risk cannot be calculated in advance and the initial event—which usually is not even seen as a component failure at all—becomes linked with other systems from which it was believed to be independent. The other systems are not part of any expected production sequence. The linkage is not only unexpected but once it has occurred it is not even well understood or easily traced back to its source. Knowledge of the behavior of the human-made material in its new ecological niche is extremely limited by its very novelty.

Eco-system accidents illustrate the tight coupling between human-made systems and natural systems. There are few or no deliberate buffers inserted between the two systems because the designers never expected them to be connected. At its roots, the eco-system accident is the result of a design error, namely the inadequate definition of system boundaries.

It is only since science has learned to replicate complex physical, chemical, and biological processes in the laboratory that its actions have been so consequential for the eco-system. The frequency of unintended interventions in the eco-system are likely to increase as the keys to more natural process are discovered. At least since Rachel Carson, concerned scientists and activists have worried about potential negative interactions. They have tried to anticipate the normally unexpected. This form of ecological consciousness was at its peak in the early seventies when molecular biologists first raised the possibility that an as yet untested technique, the "splicing" together of genetic materials from unrelated species and their implanting in either an intended or unintended host, might be potentially catastrophic. An analysis and discussion of the way in which this was handled by the scientific community will be instructive in its policy implications for other complex, tightly coupled systems. But before we get to that, a brief introduction to recombinant technology is in order.

Gene Splicing

The recombinant technology, popularly known as gene splicing, involves a set of techniques that allows scientists to use special enzymes to cut into pieces the long double strands of molecules that make up DNA and then to *recombine* the pieces with the DNA of a carrier, called the

"vector." These recombined molecules are then inserted into a host where they will presumably propagate. By combining the foreign DNA into a vector that typically replicates in the host organism, scientists are able to induce the "expression" of the foreign genetic material in its new host. For example, if the foreign DNA carried the genetic information for human growth hormone, it might then be combined with a vector that replicated in some readily available host bacteria. The host bacteria would then begin manufacturing human growth hormone. This has actually been accomplished. Prior to this, all human growth hormone had to be extracted from human blood. As a result, it was rare and expensive. Its mass production in bacteria will be a tremendous asset to the treatment of childhood growth disorders.

But early in the development of these procedures controversies arose over the unrestrained application of these techniques to organisms that were clearly quite dangerous. A biochemist at Stanford, Paul Berg, was planning to put Simian Virus 40 (SV40), which causes tumors in monkeys, into *E. coli*, a bacteria that abounds in the human gut and is widely used for research. There was some evidence that SV40 could alter human cells to resemble tumor cells, and in this lay the potential danger. By slightly altering a virus that was potentially carcinogenic in humans, a new organism could be created that was even more lethal than the parent organism. But more important, by transplanting the offensive gene to a bacterial host such as *E. coli*, the far more disturbing prospect presented itself that experiments might create a previously unknown bacterial strand that, were it to escape, could conceivably cause a cancer epidemic. Robert Pollock, a microbiologist, learned about this planned experiment from Berg's student, who was visiting the Cold Spring Harbor laboratory for a summer seminar. Pollack was perhaps the first scientist to recognize the danger of creating a recombinant hybrid with unknown infectivity that might be capable of surviving in humans. His worried call to Berg touched off first anger and then concern in Berg.³⁵

Over the course of the next six months Berg raised the issue with colleagues all over the country. As a result, he postponed the SV40 experiment. By June 1973 other scientists had expressed their concern to the National Academy of Science (NAS). The NAS asked Berg to organize a committee to discuss the recombinant DNA safety issue. The NAS group convened in April 1974 at M.I.T. Their decision was to call an international conference and to request a voluntary moratorium on specific forms of recombinant experiments thought by the group to be potentially hazardous. It is worthwhile to note that such a moratorium is

probably unique in the annals of science. It was enforced through peer pressure, strengthened by the fact that the elite of the molecular biology world were its initiators.

The international conference came together seven months later, in February 1975, at the Asilomar Conference Center in Pacific Grove, California. The conferees began by trying to grade the risk of various experiments. Committees went on to assess specific controversial experiments. The outcome of the conference was a three-tier, graded series of containment measures to match the graded risk of various experiments. The intention was to put buffers between the recombinant DNA molecules and the environment. Only a failure in containment could then result in an accident. Sixteen months later, after considerable debate and conflict in the research community, the rough plans laid out at Asilomar were turned into specific NIH guidelines.

In our terms, the scientists were reducing risks by decoupling a system that might potentially be tightly coupled with the environment. Their efforts were complicated 1) by the fact that all accident scenarios were hypothetical, based on limited knowledge, and 2) by the politics involved in restricting a burgeoning and highly valuable field of scientific research, one with tremendous commercial potential as well.

In setting up the NIH guidelines scientists had to grade the potential risk of a broad range of possible experiments and then to stipulate the conditions under which specific experiments might be done.³⁶ They began by trying to identify the highest risk experiments and simply prohibit them. Researchers were forbidden to work with a list of organisms classified by the Center for Disease Control and the National Cancer Institute as either pathogenic, causing disease, or oncogenic, causing tumors. Specific genes were also forbidden, among them genes that might code for toxins, plant pathogens, and drug resistance. Scientists were further forbidden from any deliberate release of a recombinant into the environment, such as use in agricultural research. Finally, all projects were to be small scale, limited to less than 10 liters of culture. By using lower-order, nonpathogenic organisms whose behavior in the lab was well understood, the writers of the guidelines hoped to restrict the realm of uncertainty.

But elimination of foreseeable risks was only a first step. More important, by far, were the guidelines for containment. Strategies for containment were of two kinds, biological and physical. Both strategies involved graded levels of containment that were to be matched with the expected risk of the experiment being performed. The principle behind biological containment is that each subsystem ought to be a transmission barrier.

That is, viral genes, vectors, and hosts would all be altered, insofar as possible, so that they accomplish their singular purpose without being capable of other interactions. Thus, plasmids (a vector) were to be chosen for their inability (or low-frequency ability) to exchange genes between cells. Hosts, such as *E. coli*, were to be sufficiently enfeebled so that they would not survive outside the lab.

Thus in biological containment the subsystems serve double duty as functional units and as buffers. The use of these buffers constrain the limits of permissible research. When the guidelines were issued, only a small number of vectors and hosts could be used in experimentation. These limits were being questioned even before the guidelines could be published. In fact, several cases arose where researchers were discovered to be working with prohibited viruses or vectors. In general, the greatest drawback to this form of containment is that it inhibits innovation and is then likely to be resisted or circumvented where the production system is part of a competitive economy.

The other form of containment, physical containment, is based on the more traditional procedures for separating potentially pathogenic organisms from the environment. The lowest level of physical containment is often compared to the usual lab procedures of a well-trained microbiologist. Most of these procedures are aimed at protecting the subject culture from contamination by wild strains, but the buffer effect works both ways. Higher levels of containment, known as P2 and P3, include sterilization, washing of hands, sealed windows, exhaust filters, and special safety cabinets if material might be sprayed into the air. The highest level, P4, requires the changing of all clothing, showers, and negative air pressure to keep all particles inside. At the time the guidelines were issued, several universities had begun the construction of P3 facilities. Only the government had a P4 facility, the former biological warfare center at Fort Detrick, Maryland.

Within a few short years, however, the scientific community did an abrupt about-face on the regulation of DNA technology.³⁷ Faced with the prospect of stringent federal guidelines restricting the kind of research that could be done, the biologists began to reassess the risks of DNA procedures in order to calm the surge of public fear. The most publicized experiment regarding the potential dangers was the study of Martin and Rowe. They asked, How dangerous could an escaped recombinant be if it had received a gene from a lethal donor organism? To answer this question, they grafted the genetic information from a cancer virus into *E. coli* and infected mice with the new bacterial strain. What they discovered was that the new bacteria were either noninfectious or far less infectious

(by a factor of 10^9 or one in a billion) than the original cancer virus itself. Martin and Rowe felt this dispelled any fears that an escaped organism might pose some sort of health hazard. As Rowe explained it, the research demonstrated that there was nothing one could cut out of a small-pox virus which, inserted in *E. coli*, would be dangerous to work with on an open laboratory bench. "The same applies to all the tumor viruses and even the lethal Lassa fever virus," asserted Rowe.³⁸ The outcome of this research and the growing sense of familiarity with the technology led the biologists to successfully forestall the proposed federal legislation regarding restrictions on DNA research. By September of 1979 the Recombinant Advisory Committee (RAC) of the NIH had voted to exempt 80 to 85 percent of recombinant DNA research from all but the most standard of lab safety procedures. In fact there had been a rather radical change in the RAC's attitude toward risk in less than three years. The change has been characterized by Thomasson as going from protection against the worst possible case to protection against believable risk.³⁹

The question must be asked, however, whether the new evidence actually warranted such a radical departure from the risk-conscious policies that had been implemented in the seventies. There are at least some informed scientists who contend that the Martin and Rowe experiments ought not to be taken as conclusive evidence that stringent safety regulations were unnecessary. Indeed, critics pointed out that the scientific community should not focus on the fact that the recombinant bacteria were less infectious than the tumor virus itself, but rather should note that these experiments demonstrated quite conclusively that the lethal trait could be transferred through DNA techniques. In the experiments conducted by Martin and Rowe, "about half of the mice that were injected with a bacteriophage containing a dimeric (two-copy) form of recombinant DNA" did contract the polyoma infection. To some observers this finding suggested that what was notable about the experiments was that they demonstrated how DNA research could in fact create a new vehicle for the transmission of such hazardous traits.⁴⁰

There have also been other scenarios put forward suggesting that the peculiar and subtle complexities of recombinant organisms might lead to serious health hazards simply by interacting with biological systems in ways that are novel and consequently unbuffered. Perhaps the best example of a system accident scenario in this area is based on the autoimmune response in humans. Such a case requires the splicing of a protein from an animal donor into a host bacteria such as *E. coli*. Through a trivial accident, the *E. coli* could then establish itself in the intestinal tract of a laboratory worker. The recombinant would begin to produce

the animal protein in the lab worker's system. The animal protein is similar in structure to human protein. The lab worker's immune system would be triggered and would attack the foreign protein. But because of the similarity in structure, the antibodies would not distinguish between foreign and indigenous protein. The immune subsystem would thus attack healthy tissue in the lab worker. This process is called autoimmune disease. While immunologists believe the probability of such a scenario is limited, their models cannot completely reject it.⁴¹

Whoever we choose to side with in this debate, however, one fact does stand out. The special irony of the regulatory history of DNA technology is that the serious and responsible concerns of the scientific community that were registered at Asimolar in 1975 have since rebounded in such a way that there is now a special aversion and stiff-necked resistance to the imposition of any nonvoluntary regulations on DNA research. Whether this aversion would be so unchallenged and widespread if the scientific community had not been forced into a legislative corner in the late seventies is a question about which we can only speculate. It does seem likely, however, that one important legacy of the experience of having to organize and implement an intensive lobbying effort in Washington has been the growth of a pointed reluctance among scientists involved in the DNA field to openly call into question the safety of these procedures. Such a reaction is not at all surprising, and the prevalence of such an attitude has been noted by a variety of commentators.⁴² What is especially troubling about such a backlash, however, is the suggestion that certain institutional sanctions such as tenure and the availability of research funding are being used as an added incentive to discourage dissenting and specifically junior faculty from making waves. Whether we accept these allegations or not, it is notable that the current laissez-faire attitude toward DNA research clearly distinguishes the American research effort from the regulatory climate of other countries. Gene-splicing technology in Britain, for example, is constrained by a stringent set of regulations that require "medical monitoring, long-term epidemiological studies, prior review of recombinant DNA experiments and periodic inspection of research facilities as well as containment levels that tend to be higher than those in the United States."⁴³ The Japanese, now technological entrepreneurs of the first order, have nevertheless implemented a strict set of national policies which are patterned after the original Asimolar-inspired National Institute of Health (NIH) guidelines.⁴⁴ The economic projection we noted at the outset, the great interest of private, for-profit firms, and the popularity of such U.S. firms as Genetech in our stock market may have something to do with this international difference.

When this discrepancy is brought to the attention of the American research community, their reaction has characteristically been to suggest that regulatory developments are unnecessarily restrictive and likely to hinder the advance of scientific knowledge. Noting that the Japanese have shackled themselves with stringent safety regulations, Dr. Peter Farley of the Cetus firm argues that as a result the Japanese have come to be less of a competitive challenge in the field. In 1980 Dr. Farley, estimating that the Japanese were three to five years behind the Americans, suggested that Japan:

... on the other hand, is considerably less of a threat. Overall the Japanese scientists [sic] are absolutely superb microbiologists. However, they are unbelievably handicapped now because Japan has recently adopted the original very stringent NIH guidelines. This short-sighted action on the part of the Japanese regulatory agencies was received with amazement in our shop.⁴³

But as the authors of the article quoted above quite astutely point out, Dr. Farley's concern with the potentially stifling effects of controls upon his shop were contradicted when, a year later, he expressed concern that the Japanese had made considerable headway in the field and were at that point only a year behind the Americans.

In his emphasis on the competitiveness of the research, Dr. Farley gives voice to what may well be the most important and perhaps the most disturbing change in the direction of the DNA field since its inception. The rush to develop the field may have less to do with the intellectual excitement of the subject matter than with the incredible economic incentives that are now apparent. It appears that research in the field is more and more coming to be perceived as an economic competition rather than a scientific one, and some universities are making the most of the economic potential.

Though I have not spent any time in either setting, it seems quite possible that the attention to potential risks is considerably higher in the university research laboratory than in the commercial lab. University technicians work under less commercial pressure, are probably better trained, work on a variety of experiments, giving them more generalized knowledge, and are closely supervised by graduate students and professors. University labs no doubt vary in these respects, but the careers of the graduate students (who do most of the work, and who will become Ph.D. researchers shortly) and the researchers (generally, but not always, professors), do not depend upon rapid, high-volume production or finding something that simply works. Their efforts must be carefully documented, repeated, written up for professional publication, and reviewed

by their peers. The pressures to be "first" are great, but the system emphasizes care, documentation and, perhaps above all, scientific understanding or comprehension. I believe the risks of an accident are still considerable, but not as great as in a commercial lab.

In the commercial lab, economic pressures would appear to emphasize repeated trial and error to find results that work, without the probing inquiry that would lead to caution and reflection. Procedures and results are not as likely to be subjected to peer review, where errors can be detected; and there are pressures to move to commercial production as soon as possible, with an emphasis on production short-cuts and cost-cutting procedures. There are fewer buffers in this environment. There is no reason to believe that the commercial firms, under tremendous financial pressures to be the first to market the new organisms, will be any less casual or sloppy than the high-technology aerospace firms that left metal chips in spacesuits, caps off pipes, and contaminated the drinking water of the astronauts. It was easy to be aware of the implications of these failures. It will be extremely difficult to recognize the implications of failures in laboratory containment, or, once the products are marketed, in the more fearsome and unpredictable matter of unexpected interactions with a very complex and poorly understood natural environment.

Nicholas Wade argued the point in terms that will be familiar to us: "We can't even predict the behavior of totally man-made systems such as nuclear power plants, say, much less can we predict the behavior of very much more complicated systems such as *E. coli*, of which we only know about half the working parts."⁴⁶

This comparison would be rejected by the biological community; in my experience they and chemical and aerospace engineers look on engineers in the nuclear power industry as some form of bungling drop-outs from a truly scientific world. But we have seen that complex systems go beyond the capacity of engineers and designers in all areas. A measure of humility, such as attended the Asilomar conference, would seem to be in order. In contrast, researchers in both the university and commercial laboratories seem to feel that we know all we need to know about the risks of the technology. In a personal communication, Sheldon Krimsky suggests that what has emerged is a type of orthodoxy that holds that "genetic materials will either do what they are supposed to do when they are displaced, or do nothing at all." From the vantage point of the systems examined in this book, this assertion of confidence seems quite unrealistic. In our rush for scientific fame or private profit we may be preparing the ultimate accident; indeed, it may already have occurred.