

UNCLASSIFIED

Defense Technical Information Center
Compilation Part Notice

ADP023054

TITLE: Then and Now: Flight Research in the Second Half of the 20th Century

DISTRIBUTION: Approved for public release, distribution unlimited

This paper is part of the following report:

TITLE: SAFE Journal. Volume 34, Number 1, Fall 2006

To order the complete compilation report, use: ADA458046

The component part is provided here to allow users access to individually authored sections of proceedings, annals, symposia, etc. However, the component should be considered within the context of the overall compilation report and not as a stand-alone technical report.

The following component part numbers comprise the compilation report:
ADP023051 thru ADP023054

UNCLASSIFIED

Then and Now: Flight Research in the Second Half of the 20th Century

Curtis Peebles

On a morning in the first decade of the 21st century, a research aircraft and its chase planes wait at the end of the runway. Once everything is ready, they take off and climb into the clear blue sky. The research pilot then begins the first test point as the chase planes and ground controllers keep watch. The carefully choreographed flight plan is carried out at the planned speeds, altitudes, dynamic pressures, angles of attack and sideslip. The successful flight is the result of more than 50 years of advances in flight safety. And "flight safety" means not only survival equipment, but also flight planning, test procedures, simulations and a vast database of aerodynamic knowledge and experience. When the mission is over, the airplanes landed, and post-flight debrief completed, the research pilots, engineers and support personnel leave the NASA Dryden Flight Research Center, located on Edwards Air Force Base, California, by driving down Lilly Avenue.

Named for Howard Lilly, the first NACA research pilot killed in the line of duty, it is a reminder both of how much has been learned and the price paid for it. Today, few people remain who experienced that time, when the facility was limited to a single hangar with an attached lean-to for office space and a few makeshift dorms as housing. This was a time when a trip to Los Angeles required a long bus ride on winding two-lane mountain roads. Most important of all, it was a time when the pilots and crewmen were flying into the unknown.

The two decades following the end of World War II saw a revolution in aviation technology. Every aspect of aviation design and technology

would change during this period. Pilots found themselves flying, on a daily basis, at speeds two or three times faster than they had during the war. The development of new aircraft, driven by Cold War rivalry and improvements in aircraft performance, came at a rapid pace. An environment characterized by rapidly changing technology, ever-greater speeds and altitudes and aerodynamic and engineering unknowns put test and research pilots into situations for which they were not prepared.

The result was a loss rate that today would be entirely unacceptable. Between 1947 and 1967, spanning the first two decades of supersonic flight, a total of 107 pilots, aircrew and passengers were lost in crashes. The losses came in 69 accidents, which included those during research missions, cross country flights and proficiency hops.

Test pilots in the late 1940s found themselves flying at high speeds with life support and survival equipment not significantly different than the gear worn in open-cockpit aircraft during the 1920s and 1930s. When Air Force Capt. Charles E. "Chuck" Yeager exceeded Mach 1 for the first time, he was wearing a standard-issue flight suit, boots, oxygen mask and parachute. In the 1940s, pilots still wore leather flight helmets, which did little more than keep their earphones in place. To protect himself should the X-1 go out of control, Yeager scrounged up an Army tanker's helmet that was heavily padded and provided limited protection. He wore this on his Mach 1 flight. It was not until the end of the 1940s that the first fiberglass hard-shell flight helmets were issued.



Charles E. Yeager in the late 1940s. He is wearing a standard flight suit, seat parachute, and holding a hard shell helmet. The flight suit was made of cotton and would burn. The later nylon flight suits were quickly removed from service, as they would melt and stick to the pilot's skin, causing severe burns. The hatch in the side of the X-1 was useless for an in flight escape, as it was forward of the wing. *Air Force photo*



Charles E. Yeager and Arthur "Kit" Murrey and the X-1A. Both are wearing the early partial pressure suits. These provided protection should the cockpit depressurize, but were very uncomfortable. *Air Force photo*

These were not the only shortcomings in pilot equipment. Before World War II, anything

above 10,000 feet was considered "high altitude." During the war, this definition was greatly expanded. Reconnaissance aircraft began reaching altitudes higher than 40,000 feet. Pilots could see the curvature of the Earth's horizon below, while above the sky was a deep blue-black. Above this altitude, the atmospheric pressure was too low for human survival, even with a supply of breathing oxygen. It was necessary for the cockpit to be pressurized and for the pilot to wear a pressure suit.

The X-1 series and the Douglas D-558-II Skyrocket were the first aircraft to reach altitudes at which a pressure suit was required. These early suits were adequate but very uncomfortable. The early versions were called "partial pressure" suits. They resembled a tight-fitting flight suit made of heavy fabric and had tubes running down the pilot's arms and legs and that literally squeezed the pilot. Should the cockpit lose pressure, the tubes would inflate, drawing the fabric even tighter. This protected the pilot against the effects of depressurization.

The suit was tight even when depressurized. It lacked a cooling system. The pilot's own body heat would build up, causing him to perspire and leaving him soaked in his own sweat. It was not unusual for a pilot to lose several pounds during a long flight. Despite its shortcomings, however, the suit soon proved its worth. On August 25, 1949, Air Force test pilot Frank K. Everest Jr. was making a flight in the X-1 when he lost cockpit pressurization at 69,000 feet. This event marked the first operational use of a partial pressure suit to save a pilot's life.

Aircraft escape systems of the period also were deficient. Pilots of early U.S. and British jets had to open the canopy and climb out in an emergency, the same procedure as had been used since World War I. The pilots of the X-1 series aircraft faced the same problem. The original Bell Aircraft X-1 had a hatch in the right side of the cockpit through which the pilot would enter once the B-29 launch aircraft took off. The hatch was useless for escape during an in-flight emergency, however. It was located directly in front of the wing's leading edge, and the pilot would be struck if he tried to jump. The

later, second-generation aircraft, the X-1A, X-1B and X-1D, were fitted with a conventional canopy design but the pilot still had no choice but to make his escape by jumping over the side.

During World War II, a few Nazi German fighters had been equipped with crude ejection seats. Following the war, ejection seats began to be fitted into production jet fighters. The ejection seats of this era had limited operating envelopes. An ejection at low altitude and/or low speed would not allow the pilot enough time to open his parachute before hitting the ground. Problems also existed at the upper end of the performance envelope; at supersonic speeds, an ejection would subject the pilot to a violent wind blast.

Given the shortcomings with ejection seat technology, and the difficult environments of high-speed/high-altitude flight, some argued the obvious: a different approach was needed. This took the shape of a capsule system. The D-558-I, D-558-II and the Bell X-2 all featured a nose capsule designed to separate in an emergency. Pilots for both the Air Force and the National Advisory Committee for Aeronautics (NACA) who flew these aircraft, however, had little faith in the nose capsules' usefulness in an emergency.

NACA pilot Stanley Butchart had this to say about the D-558-I capsule: "They had a piece of paper showing us the speed and altitude envelope where you would be safe to get out. You got out of those things by pulling one handle which dropped the nose of the machine off -- then another handle that would release your little back rest and you kind of crawled out the back. That's not much of a way to get out of an airplane when you're in trouble. The envelope was rather restricted too as far as speed and altitude. When you stop to think of it, [at] the higher speeds, and you drop the nose off, you're going to get a very big negative g as you come out of there. So that restricts you as to how fast you can be going and still use that escape method. We would look at that and kind of throw it in the back of the desk and go on about our work."

NACA pilot Scott Crossfield, who flew both the D-558-I and -II, was even more critical, noting: "...this is the way to commit suicide, to keep from getting killed. They never did have the development on them that they should have had, and they weren't any good anyway. If you could make a capsule that was good enough to live through the emergency, you might as well fly it and throw away the airplane." The shortcomings he and Butchart saw with capsule escape systems became a tragic reality during the early years of high-speed flight.

Howard Lilly was killed in a crash of a D-558-I on May 3, 1948. Lilly had just lifted off the Edwards lakebed when witnesses saw a large section of the fuselage skin separate from the aircraft, followed by smoke and flames. The D-558-I wallowed in flight for a few seconds, then began a left yaw and roll and dove into the lakebed. Lilly was killed on impact. The crash investigation showed that the compressor section of the plane's jet engine had failed. This sent chunks of the compressor housing and broken turbine blades out the side of the aircraft, cutting the control cables and fuel lines. Flying too low for the capsule to be of any use, Lilly never had a chance.

The X-2 also was fitted with a capsule system. Designers envisioned that in an emergency, the X-2 pilot would separate the capsule from the rest of the vehicle's fuselage. A static line would deploy the stabilization chute attached to the back of the capsule. This would maintain the capsule in a nose-down attitude, and begin slowing it to a terminal velocity of 120 miles per hour. Once the capsule had slowed and was below an altitude of 10,000 feet, the X-2 pilot would jettison the canopy, climb out of the cramped cockpit and open his parachute.

Many found fault with the concept. Everest shared the low regard of his NACA pilot counterparts toward capsule escape systems, noting that the separation would subject the pilot to a negative 14g acceleration. As the pilot's pressure suit helmet was nearly touching the X-2's canopy, he would almost certainly be knocked unconscious. Everest viewed the capsule design to be unsatisfactory and would

never have used it except in a dire emergency due to the extreme g forces to which he knew he would be subjected.



Air Force test pilots Iven Kincheloe (standing) and Mel Apt with the X-2 at Edwards AFB. The small size of the capsule meant that Apts' knees were even with the cockpit railing. When the bulk of a partial pressure suit was added, the pilot could barely fit into the aircraft. In the event of an emergency, an X-2 pilot would have to climb out of the falling capsule, and parachute to a landing.
Air Force photo

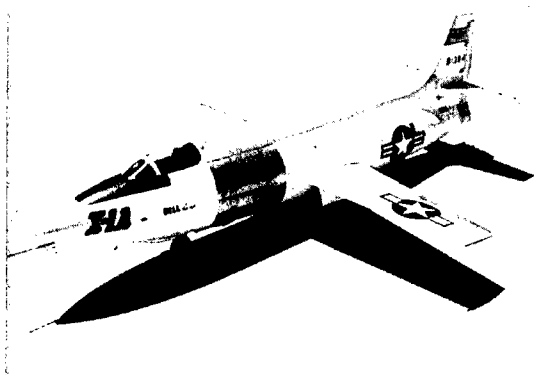
Beyond the issues of antiquated survival equipment and the deficiencies in escape systems, there were much more basic differences between flight safety conditions then and now. Simply put, present-day safety procedures and risk assessment concepts would not have been understood by pilots and engineers of the 1940s and 1950s. This is due in part to the fact that they lived in a different time, when aircraft accidents were far more common and the risks both less understood and more acceptable. Procedures at that time were geared to propeller aircraft but were being applied to flight testing of high-speed jets and supersonic rocket planes. The more significant differences, however, were with the tools, knowledge and experiences we have today but which had yet to be developed in the early years of the jet age. Just as the technology of aviation had to evolve to meet the demands of the postwar era, so too did flight planning, training procedures and data processing.

In the 1940s and 1950s, there were two competing philosophies of research flight planning. The first, which was based on systematic, incremental speed and altitude build-up, was favored by the NACA. Typically, speeds during NACA research flights would be increased by only a 0.1 Mach number on each flight. This approach resulted in an extraordinarily good safety record for the NACA. Between its founding in 1917 and Lilly's death in 1948, no NACA pilot was killed during a research flight. The NACA approach was to collect the most complete and accurate data possible, with the time required to collect that data a secondary consideration. Setting speed or altitude records was not an issue.

The Air Force favored an alternative philosophy that valued speed over thoroughness and reflected the rapid advancements of aviation technology in this period, the demands of the Cold War and the political imperative of attaining speed and altitude records. Flights were made with jumps in Mach numbers of 0.5 or greater. This expedited flight test approach reflected the Air Force's operational focus, as compared to the NACA's research priorities. The Air Force was in the midst of being transformed into a service branch based on jet-powered and supersonic aircraft. It needed flight-ready aircraft as soon as possible. If flight characteristics had shortcomings, the thinking went, these problems could be corrected in later production versions of the aircraft.

At the NACA Flight Research Center (now the NASA Dryden Flight Research Center), goals could be shaped by input from other NACA centers such as the Langley Research Center in Virginia or the Ames Research Center in northern California, based on the centers' research activities. Input also might come from military services, or from contractors or from Dryden center chief Walt Williams. There were no designated flight planners in the 1940s or 1950s. Project engineers did their own mission planning, and would establish whatever procedures they thought necessary to obtain the data they sought. Equipment needed to carry out a mission would be fabricated in a machine shop. The engineers would indicate what they

needed, and technicians would devise ways of building it. No paperwork was needed.



The X-1A after being turned over to the NACA. One of the first modifications made to the aircraft was the addition of an ejection seat. NASA photo

After each flight, data would be worked up and analyzed for any indications of dangers ahead. If any were suspected, the flight might be repeated to be sure the warning was valid. The flight plan would include “off-ramps” – contingencies such as alternative flight plans or emergency procedures – so the pilot would know ahead of time what to do if problems arose during the flight.

In the Air Force, there was, officially, a chain of command regarding test flights. In late 1953, when Yeager reached a speed of Mach 2.44 in the Bell X-1A, Gen. Al Boyd, commander of the Air Research and Development Command at Wright-Patterson Air Force Base in Ohio, had ultimate responsibility. The “approving official” was Brig. Gen. J.S. Holtner, commander of the Air Force Flight Test Center (AFFTC) at Edwards, who had the go/no go authority for the flight. Everest was Yeager’s immediate superior and would have signed off on the flight as well.

The reality, as Yeager recalled some three decades later, was more causal. He wrote: “By now these rocket research flights were so routine that [Capt.] Jack [Ridley] and I were on our own, pretty well free to do our own planning and flight profiles with neither NACA nor the Air Force looking over our shoulders. General Boyd, for example, was back at Wright, taking charge

of a missile development program. And the NACA guys now had their own test flight program and could care less about ours.”

In practical terms, both these approaches meant that those who planned the flight were also the ones assessing the mission’s safety. Without the tools and procedures of today, however, they had little to go on. One major source of data was wind tunnel testing performed during development of the aircraft. When Yeager and Ridley planned the high-speed flights in the X-1A, they knew based on wind tunnel data that the aircraft had reduced stability at speeds in excess of Mach 2.3. The issue facing them was not the ability of the X-1A to reach high speeds; with a turbopump and increased fuel supply the aircraft could easily exceed Mach 2. Rather, it was the issue of stability. This was something they felt could be dealt with. As speed increased, they reasoned, Yeager would simply avoid any rapid control movements. Ultimately, during this time, decisions about the level of risk, and whether or not it was acceptable, were based on engineering experience as well as analytical and emotional judgments.

But this was also a time when many unknowns lurked to trap the unwary pilot. The flights faced aerodynamic phenomena that had not yet been experienced due to new aircraft designs and the speeds they were capable of reaching. The mass of conventional piston aircraft was evenly distributed between the fuselage and wings. The demands of high-speed flight altered this in the post-war generations of aircraft, in which the mass of the engines, fuel and other equipment was now concentrated in the fuselage.

Yeager and other X-1A personnel were not aware of another, far more dangerous threat. In December of 1953, “inertial coupling” was only a theoretical concept. One reason the concept was unknown was its subtlety. Inertial coupling was triggered by a combination of an aircraft’s mass distribution, speed and roll rate. These change over the course of a flight as fuel is burned and the aircraft is operated at different altitudes, performing different maneuvers. As long as these factors are *not* at the critical values, no adverse effects will occur and the

aircraft will behave normally. If the conditions *are* met, however, the results are catastrophic.

Yeager was the first to encounter the phenomena, on his Mach 2.44 flight in the X-1A. The first sign of trouble occurred as the aircraft began its speed run at 76,000 feet. Yeager noticed the X-1A beginning a slow roll to the left. He responded by applying aileron and then rudder to stop the roll. The aircraft did not stabilize but began a more rapid roll to the right. When Yeager attempted to counter this, the X-1A abruptly reversed direction into a fast left roll. Yeager shut down the rocket engine and the X-1A tumbled out of control at a speed of Mach 2.44. The aircraft made several complete rolls in one direction followed by several in the other. Yeager said later that during one roll he was looking at the Palmdale area, and then on the next he could see the mining town of Boron. From his position to the north and west of Edwards the two areas were 45 degrees apart.

Unless Yeager got the aircraft back under control, he would not survive. The X-1A lacked an ejection seat. During the violent tumbling, Yeager's head hit the canopy so hard that his helmet cracked the Plexiglas. The battering caused him to black out several times. With Yeager incapacitated, the aircraft fell some 50,000 feet, slowing to subsonic speed in an inverted spin. Yeager revived and was able to recover the X-1A first into an upright spin, then into normal glide flight. Despite being groggy from the tumble, at low altitude and without a chase plane Yeager was able to land successfully on the lakebed.

The ability of a pilot to prepare himself for a risky flight was limited. The first computerized flight simulator used at what is now Dryden was the Goodyear Electronic Differential Analyzer (GEDA). This was an electromechanical analog computer that used different voltages to indicate the values of specific qualities such as speed, altitude or attitude. At the time, digital computers existed but were too limited in speed and capabilities to do real-time simulations.

Initially, the GEDA was used to analyze aircraft flight maneuvers. However, its ability to reproduce an aircraft's handling in real time made its value as a pilot simulator obvious. The first aircraft to make use of the GEDA as a simulator was the Bell X-2. NACA engineers Richard E. Day and Donald Reisert modified the GEDA with the X-2's equations of motion and its aerodynamic and physical characteristics. This was done not with software programs but by setting rotational resistors connected with plug-in wires. A strip chart with six or eight pens recorded the data. The mechanical nature of an analog computer is reflected in the related nomenclature. Day and Reisert were called "programmers," but an analog computer was said to be "mechanized."

The GEDA filled several roles in the X-2 program. These included not just pilot training but also obtaining aerodynamic data, extracting derivatives and flight planning. Engineers would "fly" the simulation to the next planned Mach number and then write a flight plan to obtain data at that point. The process would be repeated for each data point until the flight envelope had been fully expanded. The X-2 pilot also could practice a mission on the simulator, to become familiar both with its requirements as well as with potential dangers and to practice emergency procedures.



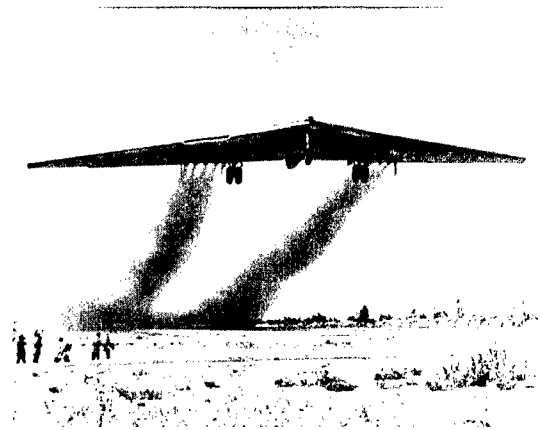
The X-1B simulator in 1958. The display is a cathode ray tube, while the cockpit is a large wooden box. *NASA photo*

The capabilities of the GEDA, however, were limited. A full simulation is called "6 degree of freedom (6 DOF)." These are movements around the yaw, pitch and roll axes; climb/dive, sideways and speed up/slow down. The GEDA lacked the capability to do a full 6 DOF simulation. It could do either a 3 DOF (yaw, pitch and roll with the other parameters fixed) or a 5 DOF for stability and control data with the speed/Mach number fixed. For the X-2 simulations, an iron pipe with centering springs was initially used as the control stick. Control position transducers translated the stick's movements into control surface inputs for the computer. (The X-2's rudder was locked at supersonic speeds, so this control surface was not included in the simulation.) The display used a cathode ray tube that showed a wing as viewed from behind. This allowed the simulator pilot to see sideslip, angle of attack and roll of the "airplane."

The newness of the simulator concept, along with the limitations of the analog computer, caused some pilots to doubt the simulator's realism. NACA engineers used data from the early X-2 flights to program the simulator. Based on this, the simulator showed that at speeds above Mach 2.4 the X-2 would become unstable due to inertial coupling when its angle of attack exceeded 4 or 5 degrees. The wings would not remain level, and aileron inputs to level the wings would exacerbate the situation until the aircraft would tumble out of control.

When X-2 project pilot Everest saw what the computer was indicating, he declared the simulation no good. At a subsequent meeting between the Air Force and NACA engineers, Col. Horace A. Hanes, director of Flight Test and Development at the AFFTC, strongly suggested to Everest that he go back and fly the simulator again. Everest did so, and was soon a believer in flight simulations. On his final X-2 flight, Everest reached a maximum speed of Mach 2.87, becoming the fastest man alive. After engine shutdown, he held the X-2's angle of attack at nearly zero degrees until the aircraft had slowed to Mach 2.2. Only then did Everest begin the turn back toward the Edwards lakebed.

While the GEDA could provide training and warn of potential dangers that loomed during a flight, a test pilot still had little in the way of outside help during a mission. Although chase planes were used, many test flights were still flown as solo missions. The loss of the YB-49, in which Capt. Glen Edwards and his four-man crew were killed, was an example of a situation in which the pilot was left to his own devices. The flight was made on a Saturday so no chase plane was available. The first indication the flying wing had crashed was when the smoke plume from the post-crash fire was reported. There was no radio distress call. As a result, the exact cause of the crash is not precisely known.



The YB-49 flying wing bomber taking off. Although a sleek and futuristic looking design, the YB-49 had severe technical and stability problems. The loss of one of the aircraft, along with Glen Edwards and his crew, remains a mystery to the lack of a chase plane on the final flight. *Air Force photo*

Chase planes supported the rocket planes' flights during launch and landing, but for the bulk of the missions, research pilots were on their own. Flight data was recorded on board using film. This would have to be developed after the airplane landed and then measured and turned into charts and graphs by the (human) computers. The only real-time data on the aircraft was supplied by radar tracking equipment in a van parked on the lakebed. The radar data showed the plane's position relative to the lakebed but was not relayed to the pilot. He had to find his own way home. No control room

existed in the 1940s and 1950s to monitor the aircraft systems in real time.

The last flight of the X-2, resulting in the death of Capt. Milburn G. Apt, underlined the limitations of flight planning procedures used in the mid-1950s, the dangers of trying to achieve speed records, the shortcomings of existing survival equipment and the technological demands inherent in the speeds and altitudes being reached. The chain of events which resulted in the X-2's crash began when Everest, the most experienced rocket pilot the Air Force had at that time, was assigned to attend Armed Forces Staff College. He made his last flight in the X-2 on July 23, 1956, reaching a speed of Mach 2.85. Everest had made all of the powered flights, and his impending transfer would leave the program without a pilot. A pair of replacement pilots had been selected in February 1956 to fill the gap. They were Capt. Iven C. Kincheloe and Apt.

Kincheloe made his first X-2 flight on May 25, 1956, reaching a speed of Mach 1.14. On his fourth flight, on Sept. 7, 1956, he reached an altitude of 126,200 feet. At this altitude, the air density was $1/250^{\text{th}}$ that at sea level, and the dynamic pressure on the aircraft had dropped close to the point where conventional aircraft controls would become ineffective. While a new world altitude record had been set, the first Mach 3 flight had yet to be accomplished.

Time was running out for the Air Force, however. NACA engineers wanted to use the X-2 to study aerodynamic and structural heating, boundary layer flow at high supersonic speeds, noise problems at supersonic speeds and aircraft handling at extreme altitudes and speeds. These studies would expand upon research being undertaken with the NACA X-1B and X-1E rocket planes.

Three more flight attempts were made by Kincheloe but were aborted before the X-2 could be launched. Apt was then selected as the pilot for the Mach 3 flight. The actual orders were that he would fly "the optimum maximum energy flight path." Apt had made training runs in the GEDA simulator and received briefings

on July 29, 1956. He was shown the procedure for reducing angle of attack to prevent loss of control. He had also made several other informal simulator runs by Sept. 24. Apt was a fully qualified test pilot with considerable experience gained in earlier inertial coupling flights in the F-100. But he was about to make his very first flight in the X-2, and was being asked to fly faster than any human had ever flown, in an aircraft known to have poor high-speed stability.

Apt would have to minimize the control movements and keep acceleration on the aircraft at 1g or below. Another difficulty he faced was that previous flights had shown the airspeed and altimeter measurements on the aircraft to be unreliable. His chances of actually reaching Mach 3, based on the simulator results, were judged minimal. Even with a full engine burn and a perfect flight path, the best speed expected was Mach 3.05. To help Apt, Kincheloe would fly as chase and coach him through the flight.

The flight was made on Sept. 27, 1956. With Apt in the cockpit, the X-2 was dropped from the B-50 launch plane. He flew the climb profile exactly, then gently pushed over into a shallow dive for the speed run. The rocket engine burned 15 seconds longer than on any previous X-2 flight. At shutdown, the X-2 was flying at 2,060 miles per hour, or Mach 3.2. The aircraft was in a 25 degree bank and a 6 degree dive. The acceleration was 1g, while the angle of attack was plus 1 degree, sideslip was 1 degree to the left and the ailerons were in a nearly neutral position. The Machmeter was still pegged at Mach 3.

It would later be speculated that Apt assumed the reading was inaccurate, and that the X-2 was actually flying slower. Crash investigators also speculated that he was worried the X-2 was too far away from the lakebed and that, unless he turned immediately, he be unable to reach it. In the X-2's tiny cockpit, he could not see Rogers Dry Lake. If this was Apt's concern it was misplaced, tragically, because the lakebed was actually in easy gliding distance.

All that is known for certain is that Apt radioed, "Engine cut, I'm turning." Within 18 seconds,

the X-2 was inverted, and beginning to roll violently. Apt was thrown about the cockpit. With the aircraft still inverted, he attempted to recover but was unsuccessful. Apt triggered the capsule. As Everest had predicted, the capsule's motions during separation were so violent that Apt was knocked unconscious. He revived, but was too low to allow time to jump from the capsule and use his own parachute. He was killed on impact.

Apt had been put into a situation that ultimately proved lethal. He was flying the X-2 for the first time and was attempting to reach its maximum speed. At these speeds, Apt had to maintain a nearly zero angle of attack until he slowed to less than Mach 2.4. But the instruments he used to decide when it was safe to turn were known to be inaccurate. Nor was there an outside source for the speed data. The X-2 was known to have poor directional stability at high Mach numbers and its control system lacked any kind of electronic stability system. Indeed, the control system design used to fly at Mach 3 was little different than that of a World War II subsonic fighter. Apt was left to cope with the pitfalls of the X-2 with only his own skills to survive. Finally, in the face of all these problems, the Air Force had charged ahead to set a speed record.

When the X-15 took to the skies three years later, it ushered in a new mode of flight research. The program involved three partners – the NACA (after 1958, NASA), the Air Force and the Navy. In a break from the NACA's more limited role in earlier X-plane programs, overall technical direction was by NASA. The speed and altitude buildup would be done with a step-by-step approach rather than by making big leaps. If a flight indicated aerodynamic or control problems, the issue would be analyzed and, if necessary, the ensuing flight plans would be altered to examine it. Not until the X-15's characteristics were understood at the far reaches of its capabilities would the next step be taken.

The flight simulator had become central to the X-15 program. There were simulators at North American Aviation, the prime contractor, and at the Flight Research Center, Ames and Langley.

At the Flight Research Center, the simulator was used extensively for pilot training. These were not the informal sessions Apt had on the GEDA. The pilot for each X-15 mission underwent some 20 hours of training on the simulator. (The actual flights took about 10 minutes from launch to landing.) This simulator work involved, in addition to practicing the mission plan, running through emergency procedures and alternate mission plans should problems arise during the flight. Engineers also used the simulator to try out flight plans and understand data from earlier flights. The engineers' work with the simulators often went well into the night.

Once the flight planning for the next X-15 mission was complete, a "tech brief" was held during which engineers and the pilot went through mission objectives, go/no go criteria and research and data requirements. This was followed by the crew brief, which was usually held the day before the flight. The crew brief brought together nearly all the operational personnel and research and instrumentation engineers. The group would go through the flight plan, discuss any items remaining from the tech brief and review limitations and mission rules.

Many of the engineers who attended the crew brief would be in the control room on flight day. The X-15's on-board instrumentation radioed data to the ground, where it was displayed on strip charts. The data included dynamic pressure, angle of attack, angle of sideslip and control surface position. Engineers watched the strip charts for any indications of trouble. The X-15's position was displayed on a plotting board. If the engineers spotted anything amiss, they would report it to the ground controller, who was referred to as NASA 1. He was the only person in direct communication with the B-52 launch plane crew, the chase plane pilots and the X-15 pilot. NASA 1 and mission control would be the Earth-bound eyes and ears of the research flight. This was an advantage earlier X-plane pilots had not enjoyed. X-15 pilots would face the unknown, but unlike their predecessors, they would not face it alone.

This approach, of a pilot supported by teams on the ground that could address problems in real time, would soon find application in the emerging piloted space missions. The increasing complexity of flight research, particularly the shift to fly-by-wire aircraft computer systems, required new levels of technological review, ground test and validation.



NASA research pilot Neil A. Armstrong following a 1960 X-15 flight. He is wearing a full pressure suit. This not only protected against depressurization, but also from heating and wind blast during a high-speed/high-altitude ejection. This eliminated the need for an escape capsule. NASA photo

The hard-won lessons learned over the past half century can be seen in the procedures used today at the NASA Dryden Flight Research Center. Unlike the often casual planning of earlier times, there are now formal procedures established for each step in the development, testing and operation of a new research vehicle. Before metal is cut on a new research aircraft, its design undergoes wind tunnel testing; computational fluid dynamics computer simulations are also made. Such data has uncertainties and variables, and to determine their potential consequences computer simulations are run. These may number thousands of runs, made through different combinations and limits, to identify potential outcomes.

When a research vehicle enters the hardware stage, it undergoes ground testing. This involves several steps, from simply determining whether

an individual component works to full end-to-end tests of the completed vehicle. As computer systems are now integral to aerospace systems, the vehicle's software also must undergo extensive testing. A matrix of "fault trees" is developed, covering how a specific system failure would affect the vehicle. Any changes in hardware or software must go through an extensive review and approval process that incorporates input from configuration control and engineering review boards as well as from experts from outside the government.

If the vehicle is piloted, customized flight simulations are developed. This serves not only to train the pilot, but also to test different flight control laws for use in flight planning and to provide warning of potential dangers. The project pilot may spend many long hours, often after the normal workday ends, in the simulator preparing for a flight.



Martha Evans, NASA simulation group leader, in an early F/A-18 simulator. The visual display is wide screen and in color, while the instrumentation reflects the actual cockpit layout. NASA photo

When all this is completed, a flight readiness review meeting is held. The engineers in charge of each aspect of the project make a presentation, and then are asked hard questions about the project's status. Review board members then make their recommendations about the project, and what more should be done before the flight is carried out.

The day before a flight is scheduled, the T-1 meeting is held. It brings together the pilots, engineers, control room personnel and support staff. The group reviews the mission plan, abort criteria and other mission requirements so that all are familiar with what is expected of them. On the morning of the flight, there may be a final briefing to cover weather issues and address any last-minute changes.

The personnel now begin final preparations for the flight. The pilots climb into the experimental aircraft and the chase planes. They are dressed in fire-resistant flight suits and gloves as well as helmets. If an emergency occurs, the pilot can fire an ejection seat which will propel him out of the cockpit. The ejection seat can propel a pilot or crewman high enough for their parachute to

open no matter if the plane is sitting on the runway or at the aircraft's maximum performance parameters. As the 20 or more belts, hoses, buckles and other connections are attached, the pilot essentially becomes one with his airplane.

As the pilots prepare, the mission control room staff take their places at the consoles. They have the latest set of checklists and contingency procedures. Their video displays show color-coded diagrams of the research vehicle's internal systems – green for normal, yellow for caution and red for a malfunction. Each controller monitors a specific system, and can call an abort if the situation requires it. As today's research aircraft sit at the end of the runway, what is to come next seems almost routine. Almost.

References

- Stanley P. Butchart history interview, Sept. 17, 1997, NASA Dryden Flight Research Center History Office.
- A. Scott Crossfield history interview, Feb. 3, 1998, NASA Dryden Flight Research Center History Office.
- _____, Clay Blair, Jr., *Always Another Dawn The Story Of A Rocket Test Pilot* (New York: Arno Press, Inc., 1972).
- Richard E. Day history interview, May 1, 1997, NASA Dryden Flight Research Center History Office.
- Richard P. Hallion, Michael Gorn, *On the Frontier Experimental Flight at NASA Dryden* (Washington, D.C.: Smithsonian Books, 2003).
- J.D. Hunley, *Toward Mach 2: The Douglas D-558 Program*, NASA SP-4222, 1999.
- Kenneth W. Iliff, Curtis Peebles, *From Runway to Orbit: Reflections of a NASA Engineer*, NASA SP-2004-4109.
- Robert W. Kempel, Richard E. Day, "A Shadow Over the Horizon The Bell X-2," *Journal: American Aviation Historical Society*, Spring 2003.
- Henry Matthews, *The Saga of Bell X-2 First of the Spaceships* (Beirut, Lebanon: HPM Publications, 1999).
- Curtis Peebles, "Risk Management in the X-Planes Era D-558-II vs. X-1A at Mach 2," *Quest* Vol. 11 No. 4, 2004.
- Ronald-Bel Stiffler, "The Bell X-2 Rocket Research Aircraft: The Flight Test Program," Air Force Flight Test Center History Office, 1957.
- Gene L. Waltman, *Black Magic and Gremlins: Analog Flight Simulations at NASA's Flight Research Center*, NASA SP-2000-4520.
- Walt Witherspoon, "Casualty List," Air Force Flight Test Center History Office.