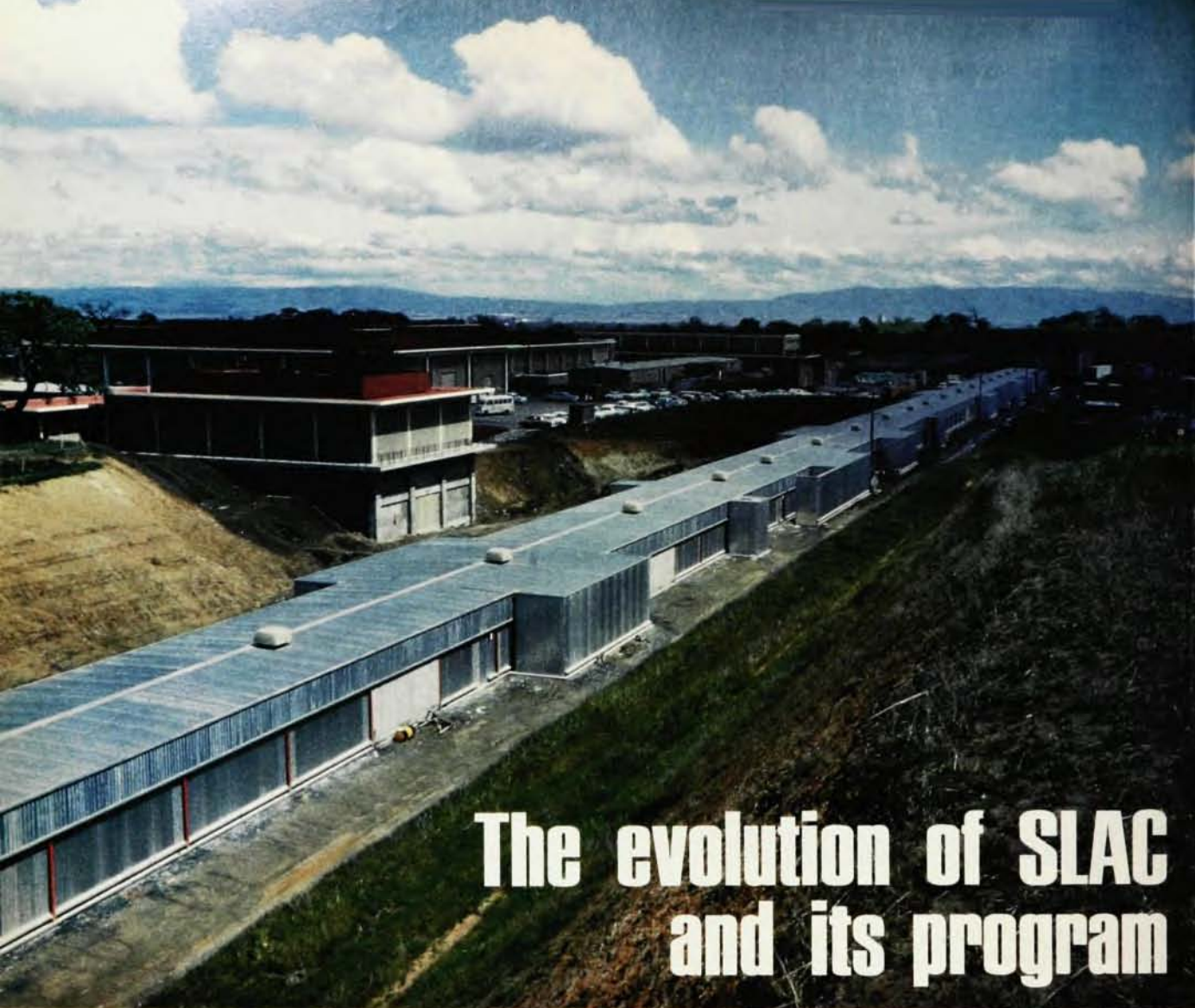




The evolution of SLAC and its program

In the two decades since construction began, the Stanford Linear Accelerator Center and its electron accelerator, two miles long, have made many fundamental contributions to particle physics.

Wolfgang K. H. Panofsky

In 1957 Stanford University proposed to the Federal government the construction of a two-mile-long linear electron accelerator. Originally called the "Monster," the machine was to be very much larger than any that had been previously carried out under the aegis of a single university. It was also an electron rather than a proton machine, and thus not in the mainstream of particle accelerators at the time. However, it seemed clear that there was useful physics to be learned from elec-

tron scattering, and the proposal for "Project M" was well received. Renamed the Stanford Linear Accelerator Center a few years later, the proposal was accepted and funded. Ground was broken in 1962, and the accelerator produced its first beam in 1966.

Since then, the laboratory has made many fundamental contributions. In the late 1960s, results from SLAC uncovered the grainy "parton" structure of the nucleon and, in the mid 1970s, experiments at SLAC found a whole new family of mesons, the psions. Detailed exploration of these new phenomena, at SLAC and elsewhere, has generated persuasive evidence for the quark structure of hadronic matter.

The tau lepton, the third charged lepton (after the electron and the muon), was found at SLAC in 1975. And in 1978 an experiment at SLAC demonstrated interference between electron scattering amplitudes of electromagnetic and weak origin, thereby giving strong support to the Weinberg-Glashow-Salam theory unifying the electromagnetic and weak interactions.

Early history

The history of electron accelerators at Stanford University started with the brilliant contributions of William W. Hansen. There has rarely been a physicist like Hansen, who combined physical insight with superb analytical

Wolfgang K. H. Panofsky, director of the Stanford Linear Accelerator Center since its inception, will step down from that post next August.

The Stanford Linear Accelerator's

10 000-foot-long klystron gallery, seen here shortly after its completion in 1965, runs directly above the earth-covered accelerator housing. We are looking upstream toward the coastal hills, with the original control center in the left foreground.

power and mechanical skills. The accelerators he built, starting in the late 1940s, made great contributions to physics. For example, Robert Hofstadter and his collaborators used them to establish the electromagnetic dimensions of the proton and the neutron and also of heavier nuclei. Moreover, inelastic electron scattering and various tests of the electromagnetic behavior of muons and pions set the stage for further discoveries in physics. In consequence, the proposal to construct the "Monster" was well received and eventually led to approval in 1962 to proceed with the construction of SLAC at Stanford University under the auspices of the Atomic Energy Commission.

The Mark III accelerator, constructed at Stanford under the leadership of Edward L. Ginzton and first operated in 1951, eventually reached 300 feet in length—30 times smaller than the proposed SLAC linac. Thus the actual creation of SLAC was a very large "leap" and required the answers to many human dilemmas, administrative difficulties, technical problems, and, above all, questions in physics.

All prior projects of Stanford University, including the construction of the earlier electron linear accelerators, were carried out within the framework of the regular departmental structure of the University. Although the W. W. Hansen Laboratories of Physics formed the umbrella laboratory under which the Mark II and Mark III electron accelerators were built, the individuals responsible were members of the regular departmental faculties of Stanford. Also the Mark II and Mark III accelerators were designed to be research tools intended for the use of Stanford faculty, staff and students; the participation of outside visiting scientists was incidental. By contrast, it became clear from the outset that a machine costing above \$100 million (at a time when a million dollars really was a million dollars!) would have to be a National Facility, that is, it should be accessible to any scientist on the basis of the quality of his proposed undertaking, without preference for Stanford people.

Electron vs. proton accelerators

At the same time, we were also fully aware of the fact that the SLAC machine was a maverick in the then prevalent pattern of US high-energy physics. At the time of the SLAC proposal in 1957, and even at the time of

groundbreaking in 1962, the main thrust of American high-energy physics depended on proton accelerators, primarily the Bevatron at Berkeley and the Cosmotron and the Alternating Gradient Synchrotron at Brookhaven. Only a small number of high-energy physicists shared Stanford's enthusiasm for electron machines. However, competition for funds then was not extremely intense. Therefore, although few physicists then intended to use SLAC, there was general acquiescence, even if not outright support, by the entire scientific community for the construction of the Stanford accelerator. One can speculate whether SLAC would ever have been built if today's financial climate had prevailed in the 1960s. One can only surmise what insights in physics would have been lost, or at least greatly delayed, had SLAC not been approved. Because few non-Stanford physicists were likely to be willing to commit large fractions of their scientific careers to plan for the new machine, we had to take the initial responsibility for planning for physics research with the machine when it was completed and made nationally accessible.

There was another important difference between the research planning for SLAC and the national pattern, which centered around proton machines: The technical nature of doing physics with an electron accelerator having a high intensity and low duty cycle required that most of the experimental program be "facility-centered." Rather elaborate detection and counting devices would have to be constructed to serve a succession of scientific experiments. In contrast, a large number of excellent experiments then being done with proton accelerators were more of the "building block" type. The participating physicists constructed experiments with relatively small components—counters and their associated electronics, shielding blocks, small magnets, and so forth. An elaborate central "facility" was not needed to conduct experiments, only to provide the beam.

The technical reasons for this difference are twofold:

► Because the linear accelerator beam is "on" for only a small fraction of the time (that is, it has a low duty cycle), it is very difficult to do experiments in which time coincidence is a primary signature identifying each event. When many counters look directly at a target exposed to an intense but low-duty-cycle beam, almost all events appear to be "in coincidence" as seen by the different counters; thus some pre-sorting of events is required.

► The particle-production cross section of interest for electron machines tends to be small, while at the same time an intense spray of electrons,

positrons and x rays constituting an electromagnetic "shower" is generated in a narrow cone in the forward direction. This highly intense cone has to be isolated from the devices that are to detect the events of interest.

It became evident at an early stage in the planning that the Stanford linac could only become a tool for excellent particle physics immediately after turn-on if we created a very strong in-house research staff. This group would have to put a large part of its scientific skills and careers "on the line" to design the equipment for exploiting the electron beam once it became a reality. The leaders of this research staff therefore had to be regular members of the Stanford University faculty, because attracting the necessary talent would only be possible if the leadership were composed of "first-class citizens" on campus. This new faculty was set up as a separate structure to prevent a major imbalance in the professorial mix within the Stanford physics department. At the same time, we assured the outside physics community of full and equitable access to the SLAC facilities, and we set up the necessary advisory committees and other administrative machinery to make sure that this assurance would correspond to reality.

We also had to convince the Stanford community that the Monster was not a threat to regular academic values. We designed the link between SLAC and the balance of the Stanford community to be *intellectually* tight but *administratively* separate so that SLAC would not overwhelm the existing administrative machinery of the university. In other words, SLAC would operate under general policy set by the university, but its actual operation would be almost entirely autonomous. This method has worked out well in practice. We then proceeded to negotiate a contract between Stanford University and what was then the Atomic Energy Commission. This negotiation resulted in a contract that fully preserved academic values and policies, and that totally delegated to the university the responsibility for managing the SLAC program.

Construction of the linac

The most essential step in building SLAC as a laboratory was, of course, the construction of the 2-mile linear accelerator itself. Here I will give only a few technical details. The job was under the direction of Richard B. Neal, and he deserves the primary credit for the construction being accomplished on schedule, without too many surprises, within budget, and to performance standards exceeding the original goals set by the proposal. The detailed story of the construction of the 2-mile machine is documented¹ in the "Blue

Book," in which people who contributed to the subsystems of the 2-mile accelerator describe the technical characteristics and history in their respective areas.

One of the first decisions made during the construction project was that building a separate prototype for the basic accelerator was not necessary. This may seem foolhardy in view of the extrapolation by a factor of 30 that was involved, but we made use of the fact that a linear accelerator is in fact linear—a small section of it can function while the larger part of it is still under construction. We therefore awarded contracts for the first 800 feet of the accelerator separately and managed to obtain a beam from this section while the rest of the machine and, in particular, the experimental target areas were far from completed. This saved the money that a separate prototype would otherwise have consumed, and it also raised our confidence that we had made no fundamental design errors.

From the point of view of electron orbit dynamics, the extrapolation by a factor of 30 in energy and accelerator length is actually minor. The accelerating fields in the disk-loaded waveguide structure of the linac do not significantly affect the radial momentum of the electrons, while the longitudinal momentum of the particles grows. Thus, even in the absence of external focusing, the orbit excursion grows only logarithmically, by a factor of 3.4. This corresponds to the fact that the relativistically contracted length as observed from the electron's frame is only 3.4 times longer for the SLAC linac than for the 100-meter Mark III machine. Focusing requirements to confine the beam are thus moderate and alignment standards are not severe.

Notwithstanding these comforting facts, the matter of stability of the machine, especially in earthquake country, received a great deal of internal and external attention. Thanks to the efforts of many seismic experts, the chosen site, which placed the injector only one-half mile east of the San Andreas fault, was analyzed in great detail. Careful design and construction practices allowed us to hold the earthquake risk to standards assuring the safety of people and minimizing the damage to physical facilities.

Both in-house talent and outside industrial contractors worked on the construction of the accelerator. The principal civil engineering for the accelerator was handled through an exceedingly capable outside architect-engineering-management firm directed by John Blume, whose help had also been crucial with the earlier seismic studies. They were responsible for the

Construction. After the klystron gallery was completed in March 1965, several of the thick-walled tunnel sections that house the SLAC accelerating structure and the beam switchyard

components were still visible before being covered with dirt. SLAC's first scientific discovery, the fossil remains of a 14-million-year-old

Paleoparadoxia, a large aquatic mammal long extinct, is evidenced by the dark patch in the left-foreground hillside. US Highway 280, then under construction, crosses over the klystron gallery.



design of all accelerator housing, beam switchyard, and building and target areas; they also managed the construction. (See the photo above.) The actual construction work was done by a large variety of contractors. Here our experiences were highly varied, ranging from contractors who did an excellent job to those who had to be persuaded every step of the way to perform as promised. We were, of course, obligated under government rules to award each item of construction to the lowest bidders, unless we could prove that they were unable to do the work, which is as difficult as proving that someone is sane when attempting to murder a President!

Let me illustrate this problem with just one example. We had received bids for a major electrical job. Our construction manager, a 75-year-old gentleman working with Blume's firm, said "Don't award the job to the lowest bidder." I asked why, and he said "Because he's a son of a bitch." The AEC manager, the late Larry Mohr, replied, "That doesn't disqualify him in the eyes of the AEC." So the job was awarded to the low bidder, and indeed our construction manager turned out to be correct.

The accelerator itself, of course, involved an enormous amount of engineering and construction of prototypes for separate components and subsystems. Feeding power from the klystrons to the accelerator required very

complex waveguide plumbing. We decided to mock up a prototype consisting of a single klystron feeding an accelerator section through the actual waveguide system. To provide for adequate shielding in the linac, the klystrons feeding power to the machine are located 25 feet above the accelerator. The mockup therefore had to be constructed as a tower containing the klystron and its supplies, placed above the accelerator section at ground level. The easiest way to install the waveguide feeds from the upper story of the tower down to the accelerator was by helicopter, a method we also used in the actual accelerator construction. As it happened, the location of this tower prototype was next to the Stanford football stadium, and it also happened that the lowering of the waveguide by helicopter was done on a Friday, the day before a critical football game. The Stanford football coach was practicing some very secret formations in the stadium at the time, and he saw the helicopter as a spy operation by Saturday's opponent! He therefore cancelled the practice, and when he was informed what the actual situation was, sent a strong letter of protest to SLAC. As you can see, building an accelerator on a university campus has its singular difficulties!

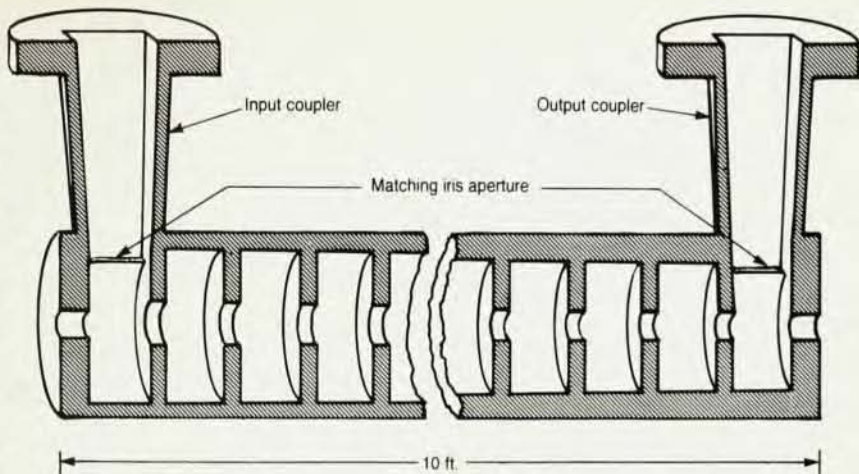
Dealing with industry

Our experience in dealing with industry was mixed, ranging all the way

from absolutely superb performance to some disappointments, not only in connection with civil construction but also with the highly technical items. We contracted with a research and development firm for the development of a prototype for the modulator to supply pulsed power to the klystrons, but half a million dollars later this resulted in a very unsatisfactory design. We then built our own prototypes for the modulator at SLAC, under the direction of Carl Olson, and procured the 245 modulators as a straight fabrication job; this procedure resulted in first-rate units, thanks to an excellent contractor, and wound up saving a lot of money.

The performance of klystrons is absolutely crucial to the success of the SLAC accelerator. Our early experience in making our own klystron tubes for the Mark III accelerator had been mixed. At certain times our tubes performed well, but there were periods where the yield of in-house production slumped and physics work almost came to a halt while we examined why performance was so poor. In view of this early history, we decided to play it safe and build up both an internal capacity to produce klystrons, under the direction of Jean Lebacqz, and also to contract with industry—in fact, with two different firms—as an additional source of supply. The first two contractors chosen were unable to perform to the required standards, and two new contractors successfully bid for the job. As a result, SLAC's inventory of klystron tubes came partly from our own efforts and partly from industry. The photograph below shows samples of klystrons from our several sources of supply.

Having an in-house capacity for building klystrons turned out to be a wise move for a number of reasons. During the early days, when the first contractors had difficulties, one of them apparently let his problems be known to Congress. I was asked, as a



The constant-gradient accelerating structure chosen for the Stanford linear accelerator. By appropriately choosing the dimensions of successive cavities, one progressively decreases the group velocity in such a way that the electric field builds up to a constant value as the energy flux in the entering wave is depleted.

witness during Congressional testimony, whether it was true that the klystron specifications that we required industry to meet were physically impossible to attain. I replied that we met these specifications with our own tubes, and that ended the dialog. In other words, having a "yardstick" operation in-house was the most powerful lever we had to assure good performance.

As time went on, the mean lifetime of the klystron tubes grew to above 20 000 hours, so the total replacement rate dropped to only about 5 tubes per month. This is insufficient to be economically attractive to industry; by mutual agreement we have therefore phased out the industrial suppliers, and now all new klystron tubes at SLAC are "homemade."

In our case, the best balance between internal and industrial work turned out to be somewhat further in the direction of building up an in-house capability than is customary at other US laboratories. This applied not only

to the case of klystrons and modulators, but also to such diverse items as magnets, various types of detectors, electronic components, and so forth. The balancing of internal and industrial suppliers continues to be a delicate issue, but clearly SLAC would have been in very serious trouble indeed, and in fact may not have survived at all, if we had not had the opportunity to pitch in with our own forces to construct vital components when necessary.

The accelerating structure

The construction of the accelerating structure required two separate decisions: the geometrical design of the structure and the choice of fabrication method. The cylindrical disk-loaded waveguide structure permits control of both the phase velocity of an electromagnetic wave with a longitudinal field component and also of the group velocity (the velocity with which the energy front propagates into the structure). The phase velocity has to be equal to the speed of light because the electron speed approximates that velocity to within one part in a thousand already after the first few feet of the accelerator. The group velocity is controllable by choice of disk aperture and, to a lesser extent, by the other relevant dimensions.

A structure that is held constant for each individual section between power feeds is easier to manufacture than a varying structure, but the accelerating electric field diminishes exponentially with distance as the energy flux diminishes through dissipation in the walls. Such a uniform, or "constant-impedance," structure thus generates less desirable acceleration and electrical-breakdown characteristics. Alternatively, by appropriately choosing the di-



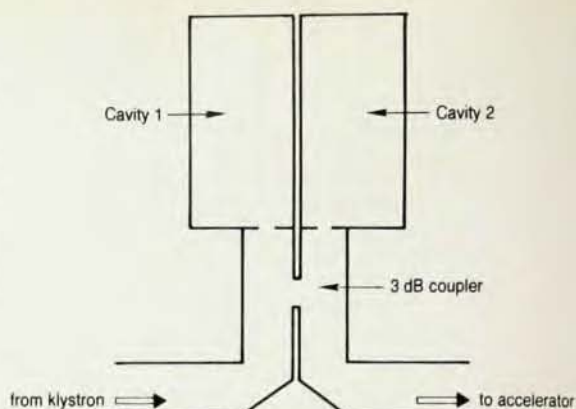
The 250 klystron tubes that supply the rf power for SLAC, have, over the years, been provided by the laboratory's in-house manufacturing facility and four commercial producers.

mensions of each successive cavity, one can progressively decrease the group velocity in just such a manner that the electric field will build up to a constant value as the energy flux in the entering wave is depleted. This results in the "constant gradient" structure shown in the figure on page 37. This was the structure chosen—a major departure from previous practice. As it turned out, this choice was fortunate, not only because it permits us ultimately to sustain larger overall acceleration (greater than 20 MeV/m), but also because it leads to greater stability at high beam intensities.

The choice of fabrication method is also a critical item. The shrink-fit fabrication methods of the Mark III accelerator would not have been suitable, because the iris disks in the accelerating structure had gradually become loosened over the years. New approaches were necessary. We developed two alternative fabrication methods and kept both of them going for over a full year as viable candidates for large-scale production. One technique was to separate copper disks by aluminum spacers, and then to electroplate a thick layer of copper on the outside. The aluminum spacers were then dissolved with lye. The other technique was to braze disks and rings (shown in the photograph below) together in a specially constructed vertical hydrogen furnace. Both methods led to satisfactory prototypes; however, we abandoned the first method because the fabrication time was quite long, and thus the interval between discovering possible defects and taking corrective action was too long.

Actual fabrication using the brazing and hydrogen-furnace technique be-

The SLAC beam energy was recently raised above 30 GeV by coupling a low-loss cavity to the waveguide linking each klystron to the accelerator. This cavity accumulates power from the klystron. When the phase of the drive to the klystron is switched, this stored energy combines with the power from the klystron, thus briefly doubling the output.



came quite a novel undertaking. It is a repetitive job, but it must be carried out to extremely high precision, and mistakes can be very damaging. The rings and spacers must be machined, annealed, then finish-machined, stacked, and finally brazed together in an atmosphere of hydrogen, using a primitive-looking furnace. All this was done by starting with about two million pounds of oxygen-free, high-conductivity copper, and the actual work was done largely by part-time workers (mostly housewives) responding to this special opportunity for steady, part-time employment lasting several years. It speaks well for the quality of that operation that not a single one of the 200 000 brazed joints in the 2-mile accelerator has developed a vacuum leak during the more than 15 years of steady operation.

Beam breakup

Commissioning the finished 2-mile accelerator was generally less difficult than had been anticipated. In fact,

obtaining a beam in the 2-mile structure was not significantly more difficult than in the older 300-foot machine; the first beam through the full machine was obtained in May 1966. However, there was one unanticipated difficulty: The beam intensity was limited by an unexpected "beam breakup" phenomenon. We had predicted beam-current limits due to the so-called backward wave instability, which would be incurred in any one accelerator section. We had not, however, anticipated the breakup caused by regenerative beam-wall interactions between successive sections. The basic physical phenomenon responsible for the problem was soon diagnosed: If a bunch of electrons within the beam travels somewhat off-axis, it excites asymmetric electromagnetic fields in the structure, and these deflect succeeding bunches even more off the axis. The resulting instability grows both in time and in space along the axis. Happily, the choice of the nonuniform, constant-gradient structure greatly mitigates this instability, because only a small portion of each section matches another. Nevertheless, the problem initially limited the attainable beam intensity to about $\frac{1}{3}$ of the design value. The cure, in the form of small structure deformation and increased magnetic focusing strength, was implemented in a number of relatively small steps that cumulatively brought the beam up to the originally predicted value.

In the original proposal, we had conservatively predicted that the energy of the accelerator would be between 10 and 20 GeV. This caution was prompted by the concern about anticipated klystron performance. Actually, the device exceeded 20 GeV early in 1967, and, with gradual growth in klystron power, the beam energy has crept up by several GeV above this value.

A subsequent modification of the accelerator, begun a few years ago, has now raised the beam energy above 30 GeV. In this scheme, a low-loss cavity is coupled to the rectangular wave-



Copper iris disks and rings were brazed together into ten-foot sections in a specially constructed hydrogen furnace to form the SLAC linac accelerating structure. The thin foil rings in the foreground are the brazing material that was placed between the components to be joined.

guide linking the klystron to the accelerator, as shown in the figure on the opposite page. This cavity accumulates the microwave power from the klystron for a short period of time. The phase of the drive to the klystron is then switched, and the stored energy in the cavity combines with the continuing power flow from the klystron in feeding the accelerator. Thus, for a period of time shorter than the usual pulse, the power flow into the accelerator is nearly doubled.

Experimental facilities

In certain ways, the original 1957 proposal was a more visionary document with respect to construction methods and human and administrative problems that it was with respect to the technical arrangements needed to meet the requirements of physics research. Perhaps this is not surprising, considering the fast pace of high-energy physics research and the decentralized initiatives guiding the research program.

I mentioned earlier that most of SLAC's research has to be "facility-centered." The initial complement of the equipment for experiments was selected from proposals that mostly came from the laboratory's staff. These proposals were reviewed extensively and publicly by SLAC, and funding was provided by the Atomic Energy Commission.

These initial facilities turned out to be quite different from the ones envisaged in the 1957 proposal. In retrospect, the electron scattering facilities described in the 1957 proposal were highly naive. They consisted of two large spectrometers, each sweeping a 180° arc. It was recognized in the actual design that there was little need for having a single detector sweep all the way from a forward to a backward direction, because back-scattered particles have much lower energies and are produced less copiously than those going forward. Accordingly, the best scheme is to have small-acceptance high-energy spectrometers operating in the forward direction only, and to complement these with instruments with much wider aperture and thus higher collection efficiency designed for low-energy particles scattered through large angles. Accordingly, we decided to build three spectrometers to cover the forward, intermediate and backward angles. These instruments, installed in a large shielded building (right), were the workhorses that produced the data to establish the point-like substructure of the proton and neutron in the late 1960s.

Hadron and photon beams

The 1957 proposal envisaged that, in addition to studies of primary interac-

tions using electron and gamma-ray beams, SLAC might also be a copious producer of secondary particles such as pi and K mesons. Historically, the use of secondary beams had been the sole province of proton accelerators; electron accelerators had not been competitive in this respect because the basic production cross sections, as well as the beam intensities, were generally lower.

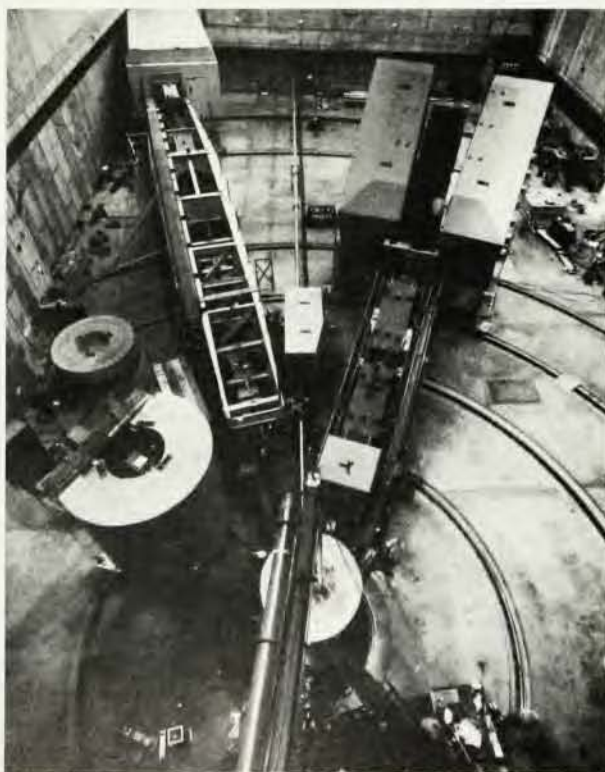
SLAC succeeded in revising this tradition for two reasons. One was that the intensities attainable in the primary beam are larger by one or two orders of magnitude than those hitherto available in other high-energy electron accelerators. The second reason is that, although the total production cross sections for secondary particles are indeed lower in electron beams than in proton beams, the secondary beams produced by very-high-energy electrons are in a very narrow forward cone. This forward concentration was predicted theoretically by Sidney Drell, and in the early days a team of SLAC physicists led by Joseph Ballam performed experiments at the now-defunct Cambridge Electron Accelerator that confirmed the essence of the theoretical prediction. Therefore, we could anticipate with some confidence that SLAC would not only be preeminent in high-energy electron and photon physics but would also be competitive in the exploitation of secondary beams composed of unstable hadrons produced on fixed targets. Accordingly, the basic structure of the research area was segregated into a complex dedicated to

studies of primary (that is electron and photon) interactions and an area dedicated to the exploitation of secondary beams.

Not only were the secondary beams produced at SLAC competitive in terms of particle flux, but some of the beams could also be designed to exhibit characteristics that were not found at proton machines. For instance, SLAC can produce high-intensity monoenergetic gamma-ray beams, whereas the photon beams from proton machines, produced by the decay of neutral pions, generally have a broad spectrum. Initially we produced monoenergetic photon beams by annihilation of positrons on atomic electrons in hydrogen; later experiments used near-monochromatic gamma-rays produced from high-energy electrons striking single-crystal targets. The most recent method, still in use, makes use of scattering high-energy electrons on visible-light photons emitted by a laser. For instance, if the 4-eV photons produced by frequency-quadrupling the light from a ruby laser collide head-on with a 30-GeV electron beam, the back-scattered photons are Doppler-shifted to 20 GeV. It is such a back-scattered monochromatic and polarizable photon beam that is currently used for experiments.

It also turned out that some secondary beams produced by photons rather than heavy particles had other desirable characteristics. For instance, beams of neutral K mesons had become popular in many laboratories because they could be used to study some very

Experimental area for electron-scattering and photoproduction experiments at SLAC. The primary electron or photon beam enters at bottom, striking a target at the common pivot point of the three large spectrometers: (left to right) the 1.8-, 20- and 8-GeV/c spectrometers, each designed to cover a different range of scattering angles. They can be moved to different angles along the circular rails.



fundamental properties of the weak interaction—for example, for determining quantitatively the parameters that define the precise level to which certain symmetries in the weak interaction are violated. The main problem with proton machines for these purposes is the contamination of neutral kaon beams with neutrons. The nature of the production mechanism of neutral kaons by electrons and photons reduces the neutron contamination, so the neutral kaon beams at SLAC can be used directly in many particle detectors, even including bubble chambers. Thus SLAC in its early days became a major contributor to the worldwide activity in furnishing quantitative values of the weak-interaction decay parameters of the neutral kaon.

Beam switchyard

The much larger variety of applications anticipated in the 1957 proposal demanded a complete reengineering of the methods of distributing the beam to experimenters. This led to the design of the "beam switchyard" carried out under the direction of Richard Taylor.

The switchyard is much more than a tool to direct beams to a variety of experimenters. It is also a "purgatory" system designed to assure that each experimenter receives electrons of known energy and energy spread, and that the "optical" properties of the primary and secondary beams entering the experimental areas are known and stable. Moreover, the requirements set by the different experimental facilities for pulse delivery rate are quite vari-

able: Bubble chambers cannot handle more than a few pulses per second, matching the chamber expansion rate, while particle spectrometers can handle all pulses made available up to the 360-per-second maximum of the linac. In the SLAC switchyard, the beam first enters a pulsed magnet that deflects the beam right or left on a programmed time pattern. Moreover, the energy of the beam can also be predetermined on a pulse-by-pulse basis by activating the required number of klystrons at the correct pulse time while mistiming the rest. As a result, each beam pulse can enter one of three magnetic channels set at a fixed momentum band. In two of these channels, the energy is magnetically dispersed halfway along the path to the switchyard to permit energy selection by a means of successive cooled slits. The beam is then refocused and directed to each experimental area. The switchyard also contains targets for producing secondary beams, including hadron beams and specialized gamma-ray beams. In general, these secondary beams could be transported to experimenters outside the switchyards by methods familiar from proton accelerators.

Because of the high average beam power—a maximum value near 1 Megawatt, unprecedented for a high-energy machine—shielding and remote maintenance requirements were severe. Also slits, collimators and beam stoppers required novel designs because one had to contend not only with high beam power but also with shock stresses due to the pulsed nature of the

beam and with radioactivity in the cooling water.

The result of all these needs was a system much more complex than envisaged in the original proposal. Fortunately, it was possible to control costs of construction of the basic accelerator and of the "conventional facilities"—the beam housing, buildings, site development and utilities—to within original estimates. Thus, most of the budget contingency could be dedicated to the switchyard.

Each experimenter downstream of the switchyard can, in essence, control his own accelerator composition, time structure, energy and energy width of the beam that arrives at the experiment. Thus the technical design of the switchyard was the primary factor in making the SLAC beams available to a substantial number of simultaneous (or, more accurately, interlaced) experiments.

Bubble chambers

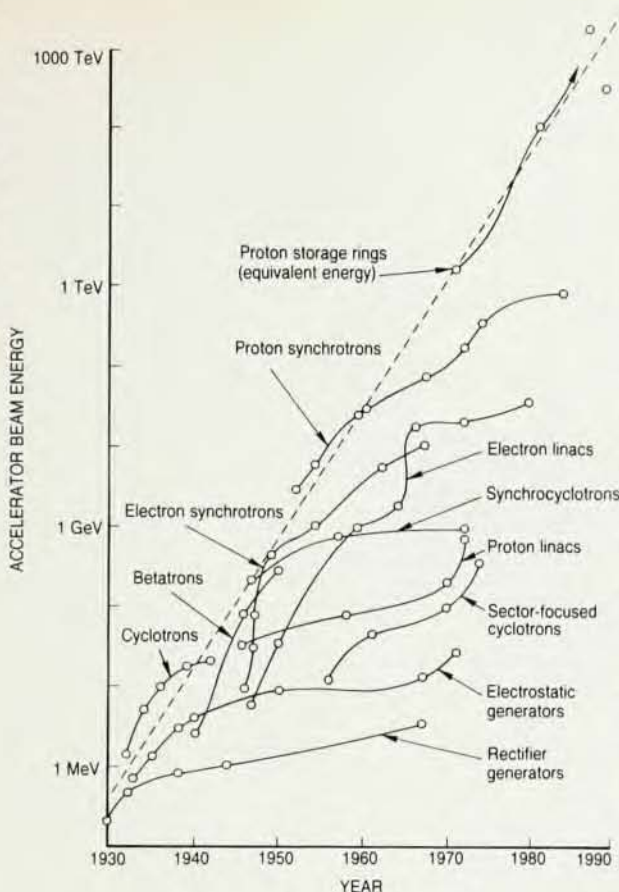
My entire story here documents the fact that the research program at SLAC became very much broader than was foreseen in the original proposal. Not only did this increased activity lead to more experimental results but at the same time it also widened the horizons of detector technology. In particular, it turned out that bubble chambers, which were not at all envisaged in the original proposal, were highly productive at SLAC. Generally, bubble chambers register in a single picture all charged particles produced during a pulse, and they can therefore handle only a few particles per pulse. Proton accelerators produce a pulse only once every few seconds, while the linear accelerator can pulse hundreds of times per second. Thus a bubble chamber can pulse more rapidly than a proton accelerator, while an electron linear accelerator can pulse more rapidly than a bubble chamber. Bubble chambers at SLAC, therefore, have much higher rates of data production than at other accelerators.

Luis Alvarez at Berkeley recognized that his famed "72-inch" bubble chamber would become very much more productive if it were moved across the bay to SLAC. This increased productivity would stem both from the enhanced data rate we have just discussed, and from the fact that SLAC could produce higher-energy secondary beams than the 6-GeV Bevatron at Berkeley. As is frequently the case,



Berkeley's historic 72-inch bubble chamber was moved to SLAC in the early 1970s, where it lived on as the "82-inch" chamber. The photo shows, shortly after the move, (left to right) Luis Alvarez, whose group at the Lawrence Berkeley Lab built the original chamber, Robert Watt, Joseph Ballam and Panofsky.

The growth of accelerator beam energies with time. For colliding-beam proton storage rings, finished, unfinished or recently abandoned, the points indicate the "equivalent" beam energy a fixed-target machine would require to achieve the same center-of-mass collision energy. Electron-positron storage rings cannot be included here in this way; the small "target" mass pushes these points way off the scale.



however, a great deal more was involved than simply moving the chamber from one place to another. Alvarez and his collaborators decided to make extensive modifications to the chamber, including a totally different expansion system. The chamber body was revised and the instrument changed from the "72-inch" to the "82-inch" bubble chamber, shown on the opposite page.

The 82-inch bubble chamber turned out to be the world's most prolific producer of exposures for the large community of high-energy physicists interested in analyzing bubble-chamber photographs. In fact, the entry of the bubble chamber into the SLAC program caused a major increase in the number of outside users. Because there were facilities for analyzing bubble-chamber photos throughout the world, the number of outside users in bubble-chamber physics has always exceeded that of in-house physicists by a large factor. More than five million bubble-chamber photographs were generated annually at SLAC for several years during the mid-1970s. In fact, the production of the 82-inch bubble chamber was so prolific that in a relatively small number of years the market for bubble-chamber photographs became saturated, because exposures at all reasonable particle energies and with all available particle

types had been made, and because the number of photographs had become so large that the limiting factor became the rate at which they could be analyzed, rather than the rate at which they could be produced.

The fact that secondary particles with electron machines are produced according to the well-understood theory of quantum electrodynamics means that searches for new unstable particles become particularly useful; if no new particles are found, one can conclude that, within the limits of available energy and expected lifetime, none exist. Such a search for new particles was carried out at SLAC by Martin Perl in the late 1960s, with negative results. It is interesting to note that it was Perl and his collaborators who later discovered the tau particle, the third charged lepton, using electron-positron storage rings.

The secondary beam fluxes at SLAC were also used extensively with other detectors, such as a very large streamer chamber built for SLAC, as well as other, more traditional, detector arrangements. All these devices have generated important physics data complementary to those produced in proton machines.

Prospects

One vexing question, raised at the time SLAC was started and which

continues to be asked today, is "How long will the life of this laboratory be?" The answer was then and still is today: About 10-15 years, *unless somebody has a good idea*. It is now 20 years after the beginning of construction, and we are again looking a decade or more ahead. As it turns out, someone has always had a good idea that was exploited, leading to a new lease on life for the laboratory. It is indeed true that full research exploitation of most, if not all, large accelerators and colliders takes about 10 to 15 years, and thus the motto relating to such machines has always been "Innovate or Die!"

Happily, new ideas have not been lacking in the environment of Stanford University. We have moved from the original proposal for the SLAC linac dedicated to electron and photon physics to the exploitation of secondary hadron beams, then to electron-positron storage rings, and most recently to the development of the SLAC linear collider—a device aimed at colliding 50-GeV electrons and positrons (see PHYSICS TODAY, January 1980, page 19).

Worldwide, as the figure on page 41 shows, we have seen a life-and-death cycle of various accelerators as the frontier of particle energy has advanced and as the kinds of accelerators and colliders needed to achieve these energies have changed. The life and death cycle of *machines* need not correspond to the life and death cycle of *institutions*, unless the size necessary for machines to remain at the frontier becomes so large that they cannot be accommodated within the boundaries of the laboratory. It is fortunate that Stanford University could accommodate a 2-mile accelerator on its own lands; thus far the additions to that accelerator—in particular the SPEAR and PEP e^+e^- storage rings and the proposed SLC collider project—also fit within the boundaries provided by Stanford University to the government under a 50-year lease. What may come after that remains an open question.

The evolution of SLAC and its program has demonstrated again that the principal contributions to physics of a new accelerator are rarely those foreseen in the original proposal. Indeed, those goals have been met, but the actual program turned out to be much richer and more exciting. Let us hope that the future will be both equally unpredictable and equally rewarding.

This article is adapted from a talk by the author, given in August 1982 at the commemoration of the 25th anniversary of the proposal for SLAC and the 20th anniversary of groundbreaking for the accelerator.

Reference

1. *The Stanford Two-Mile Accelerator*, R. B. Neal, ed., Benjamin, New York (1968). □