

Fermilab's superconducting synchrotron strives for 1 TeV

The first high-energy superconducting synchrotron in the world, the Tevatron, at Fermilab, accelerated protons to 512 GeV last July. By the time of the 12th International Conference on High-Energy Accelerators held at Fermilab in August, the Tevatron had reached 700 GeV. Since October it has been used for fixed-target experiments at 400 GeV. This month Fermilab plans to raise the energy to 750 GeV or so and then run experiments at that energy until July. During the summer, the lab plans to repair marginal magnets and director Leon Lederman hopes the Tevatron will by November indeed be worthy of its name, accelerating protons close to 1000 GeV or 1 TeV. Meanwhile, an antiproton source is being constructed; in 1986 the lab expects to have a $p\bar{p}$ collider ready for experiments with as much as 1 TeV in each beam.

Overall cost of the Tevatron, including R&D, upgrading the experimental areas, building new beam lines for fixed-target experiments, and construction of the $p\bar{p}$ collider, is estimated at \$300 million.

While Fermilab was showing a superconducting synchrotron could work, CERN was showing a $p\bar{p}$ collider could work, spectacularly well—yielding both the W (PHYSICS TODAY, April 1983, page 17) and Z^0 (PHYSICS TODAY, November, page 17) last year.

Meanwhile, Brookhaven, too, had at long last developed a reliable superconducting magnet design for its $p\bar{p}$ collider, CBA (née Isabelle). All of these successes combined to spell doom for the Brookhaven project: Last July the High Energy Physics Advisory Panel recommended that the Brookhaven project be terminated (PHYSICS TODAY, September, page 17). Instead, HEPAP declared, the US should shoot for the stars, so to speak—and start work on a Superconducting Super Collider with 10–20 TeV protons in each beam. One of the key arguments in favor of the SSC, according to the report from the HEPAP subpanel (headed by Stanley Wojcicki), was that the Tevatron was being commissioned. “This is the first high-energy accelerator resulting from



Helium flow (represented by yellow diagram) is monitored at the control panel of the Fermilab central helium liquefier facility. The central facility plus 24 satellite refrigerators provide 24 kW of cooling for the Tevatron's niobium-titanium superconducting magnets.

more than 20 years of superconducting magnet R&D, which included commercial development in the US of high-quality superconducting cable. Unit costs are therefore well known. This technology is now of age.”

The Tevatron ring, located a few feet below the old main ring, contains 990 superconducting magnets—774 superconducting dipoles, 216 superconducting quadrupoles, and, in addition, 204 spool pieces (black boxes containing correction coils and instrumentation). The magnet's coils are made of niobium-titanium, which is kept at 4.8 K by circulating liquid helium and liquid nitrogen around the coils. The cryogenic system now in operation at Fermilab is orders of magnitude larger than any previous system, having 24 kW of cooling at 4.8 K. Helen Edwards, who was in charge of starting up the Tevatron, pointed out to us that the Fermilab cryogenic system is itself as complex as a standard accelerator.

Because the Tevatron “will have beam loss, you have to be ready for it,” Edwards says. There had been some concern that extracting protons for

fixed-target experiments would be difficult because some of the protons would hit the magnets, causing them to quench. In recent operation, the Tevatron has had one or two quenches a day. The quench protection system continuously monitors voltage across the magnet cells at about 250 points and detects any resistance greater than that needed to produce 0.5 V across a cell. Once this indication of a magnet going normal is received, the culprit is heated uniformly, to protect any local spots from overheating; the current is shunted around the magnet; the main power supplies are bypassed and dump resistors are activated, so that the total energy stored in the ring (which will be 350 MJ at 1 TeV) is extracted in 12 seconds. Generally the Tevatron can then be started up again. But if it is necessary to remove a magnet, that procedure takes about two weeks, Edwards told us: You have to warm a string of 32 magnets (four cells) to essentially room temperature to avoid water condensation. Once the defective magnet is replaced, you must be sure to remove gases other than helium

from the cryogenic space to avoid freezing. So far no dipole or quadrupole magnet has failed in operation, but a handful of spool pieces have been damaged and needed replacing.

Lederman (who was named Fermilab director in 1978 after Robert R. Wilson resigned over Washington funding restrictions on the lab) told us that now that Tevatron has operated, "We think of running the accelerator 48 weeks a year but keeping the machine cold forever," essentially, except for repairing individual cells. "Every time you warm up a piece of accelerator you run some risks of creating links or other mechanical damage. A superconducting magnet is like a 21-foot wet noodle. The wires have to be positioned just right, holding the position to a few thousandths of an inch. When you cool down, you get a shrinkage of about $\frac{3}{4}$ inch per magnet. That produces terrific mechanical stress."

Early in February, Fermilab had a site-wide two-hour power failure, providing a severe test of many fail-safe procedures. The magnets slowly warmed up, the helium boiled off benignly, and various computer controls behaved rationally, Lederman told us. As the experts rushed into the control room to monitor the recovery of their systems, the scene was reminiscent of the canonical World War II movie about a torpedoed submarine.

The vacuum system for the superconducting magnets is also much more complicated than in a standard accelerator. The system has an insulating vacuum jacket with four isolated inner volumes, each of which must be leak-tight to the insulating vacuum.

Magnet design. Although the possibility of a superconducting magnet system had been discussed as early as 1970, it was not until the old main ring operated at 200 GeV in 1972 that R&D started on superconducting magnets. At that time, no one was sure that a superconducting magnet would work in the ramping field of a synchrotron because of the high energy losses produced by eddy currents. When the Tevatron operates at 1000 GeV, the magnet will be raised to full field in 20 seconds, stay there for 20 seconds, and then drop to zero in another 20 seconds. If Fermilab had decided to pulse the machine more frequently, it would have needed a still larger refrigerator to remove eddy-current heat. For every watt of heat loss at 4 K, you need 750 watts from power lines at room temperature to pump out that watt of heat, Alvin Tollestrup pointed out to us. "The reason you win is, although a refrigerator is very inefficient, it's not nearly as bad as a copper magnet is."

Fermilab did the magnet R&D in-house, designed the appropriate machine tools, built a factory to manufac-

ture superconducting coils, and then built a factory to complete the magnet assembly. The basic ingredient of the cable is filaments of niobium-titanium. The lab purchased 50 tons of the Nb-Ti alloy and then distributed it to commercial manufacturers. The alloy is first made into a 3-mm-diameter rod about a meter in length; the rods are a meter in length; the rods are inserted into a hexagonal copper tube; then 2100 of these tubes are combined into a copper cylinder with 1-inch thick walls. This billet is heated, put into a large press and extruded. Afterward the initial extrusion is drawn down through dies until a strand 0.027 inches in diameter is produced. Now the original Nb-Ti rods are in the form of 2100 filaments 10 microns in diameter, encased in copper. Then 23 strands of this material are twisted to form the final cable, known as Rutherford-style cable. The next step, called keystoneing, is to roll the cable so that it has a trapezoidal cross section, to allow the turns of the coil to assume naturally a circular shape. Originally the cable was filled with lead solder, which tended to flake when the cable was keystoneed; these flakes would subsequently cause shorts in the completed magnets. To overcome this problem, the cable was insulated by wrapping it with a combination of Kapton film and epoxy-impregnated glass tape.

The alloy was made by Teledyne Wah-chang, the reduction to strands by Intermagnetics General, Magnetic Corporation of America, Oxford-Airco and Supercon, and the cabling by New England Electric Wire Co. Fermilab coordinated and supervised quality control. The entire accelerator used \$25-million worth of superconductor, counting some wastage. Richard Lundy, who was in charge of the Fermilab magnet factory, believes that if the Superconducting Super Collider used a design similar to the Tevatron magnets, they could be built by industry. "If the design is more radical, I'd have reservations," Lundy says. SSC could require 10 000 to 20 000 magnets, a major industrial project with high risk involved, he feels.

From 1975 to 1978, Tollestrup recalls, Fermilab "looked bad to the physics community because of all the problems we developed with the magnets. But each time we were reviewed by DOE or our trustees, it'd be a different problem." He points out that although Al McInturf and William Sampson at Brookhaven had started R&D on superconducting magnets before Fermilab got into the business, "We, in a sense, educated the community to magnet problems. So when Brookhaven ran into trouble, many people knew how to criticize them."

By spring 1975 Fermilab was testing

one-foot-long magnets with a $2\frac{1}{2}$ -inch bore. The one-foot magnets were installed in the tunnel and beam transported through them. Some full-scale models had been built earlier, but they did not work. Then director Wilson decided the magnet bore should be raised from $2\frac{1}{2}$ inches to 3 because some errors in the magnet's field are a function of distance from the cable and a larger bore would provide a greater region where the field was uniform. That required higher current density in the magnets; so the number of strands per cable was increased from 17 to the final value of 23.

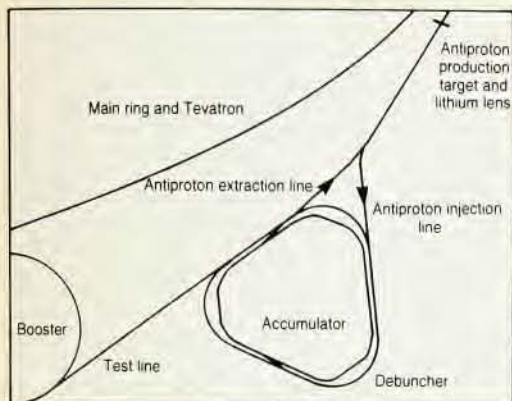
Then a series of test magnets were built, with 3-inch bore and one-foot length, which encountered a number of problems. When the magnets are operated, Tollestrup explains, there are 6000 pounds/inch trying to blow the magnet apart and 1000 pounds/inch trying to compress the winding azimuthally. How to support such forces? At first, to support the coils the experimenters inserted a porcelain ring; then the coil was overwrapped with two spiral steel bands, one clockwise and then, to counteract the torque, a second band was applied counterclockwise. However, when the magnets were energized, they tended to untwist. The solution to the force problem was to develop a stainless-steel collar to contain the forces. But a new problem arises because of the differential contraction of the coil and collar when they are cooled, the coil shrinking more than the collar. The solution was to make the coil much bigger than its final size and force the collar over the coil with a large press. The net result is that the collar continues to squeeze the coil when it's cold. The collars are made reproducibly from laminations stamped with a precision die from a stainless-steel sheet.

While the one-foot models were being built, the experimenters realized that to develop a full-scale magnet, they would need to build a full-scale factory. Fermilab built its own presses, machines for winding, wrapping, bending and collaring. As much research went into making the factory as into making the magnets, says Tollestrup, who headed the magnet developers.

In 1977 the Office of Management and Budget upgraded the Tevatron to a construction project, but $\frac{1}{6}$ of the ring continued to be called an R&D project supported by operating funds. By now R&D is considered finished, but calling $\frac{1}{6}$ of the cost R&D "permitted us to make mistakes," Lederman told us.

Richard Lundy, who had been put in charge of the magnet factory in mid 1979, told us that by then some of the dipole magnets worked well at the design field of 4.3 T and with adequate quality; that is, in a cylinder of one-inch

Construction of antiproton source at Fermilab is evident in this 1 February photo. Service building being built in foreground is above the accumulator and debuncher rings. Target hall is visible on the way to the main ring (upper left). Protons accelerated to 120 GeV in the main ring will hit the target and produce antiprotons. Then the antiprotons will be cooled and accumulated in the debuncher and accumulator rings until their density is high enough for $p\bar{p}$ collisions.



radius, the non-dipole fields would be 10^{-4} the strength of the dipole field. At that time the factory was turning out ten dipole magnets per week and had already produced 100. In January 1980, the magnet production line was halted for six months while a new problem was solved.

In the design used at that time, to compensate for the shrinking of the coil diameter when chilled, G-10 fiberglass supports were squeezed at room temperature, Lundy explained. But the necessary pressure often broke the supports and at the same time introduced a frictional force in the longitudinal direction along the cryostat. As the magnet was cooled down, it would rotate. At first the experimenters hadn't realized they had a problem because they hadn't thermally cycled the magnets many cycles. The solution, found after six months' effort: "smart bolts." Instead of the original fiberglass support, the designers used an external adjustable bolt, which acts like a spring, allowing a pressure to be maintained on the coil structure as it shrinks when cooled. As a fringe benefit, Lundy explained, they could detect any unwanted quadrupole fields and, while the magnet was cold, the smart bolts could be adjusted to shift the magnet direction as required.

Magnet production. Once the magnet factory staff realized smart bolts would solve the problem, Lundy told us, production resumed. No major design changes were needed after that summer of 1980. A year after dipole production began, they started making

quadrupoles, a much easier task because of the previous experience. The Fermilab philosophy of producing the magnets was: Don't worry about making a perfect coil. Let the magnet factory adjust the field of each magnet with shims as needed to compensate for the systematic variations introduced in production. As each magnet was produced, its properties were measured and the information given to the factory to control the size of shims used in subsequent magnets. This procedure eliminated systematic errors, but very careful quality control was needed to limit random errors. Finally, accelerator theorists determined the optimum set of magnets to put together so as to compensate for field errors. So, Lundy says, "we did the overall installation in a patchwork fashion."

Installing the magnets in the tunnel began in June 1982 and continued until last May. "We were in a rush to get the machine running again," recalls Peter Limon, because it had not been available for physics for over a year as it was. Most of the time they ran two shifts, five or six days per week. "We took people from all over the lab and even from 'Rent-a-Technician.'"

Edwards believes that starting up the Tevatron was relatively trouble-free because of the long testing period that preceded it—first a single magnet, then two, then a string of them, learning to assemble them, make them leak-proof, and so on. After the great celebration last 3 July when 512 GeV was reached, the accelerator reached 700 GeV in August. In October, phys-

ics experiments at 400 GeV began. The accelerator has also stored beam with a lifetime of about 24 hours at relatively low intensities at 400 GeV.

When Fermilab was using conventional magnets to produce 400 GeV, the repetition rate was 12–15 seconds, including a spill time (flat top) of 1 second in which data taking occurred; the accelerator produced more than 2×10^{13} protons/pulse. With superconducting magnets, at present the Tevatron produces about 10^{13} protons/pulse but eventually is expected to more than double that figure. However, to avoid the risk of magnet quenching, the Tevatron is operated with a gradual ramping, a long flat top and a gradual drop; in 400-GeV operation, the repetition rate is 39 seconds, including a 15-second flat top that is advantageous for many experiments. As the energy is raised closer to 1000 GeV, Edwards told us, the machine will probably be operated with a 60-second repetition rate. She notes that although the Tevatron doesn't produce as many protons per hour for fixed-target experiments as the old main ring, it *does* produce much higher energy protons and allows the accelerator to run more weeks per year because power consumption is much lower. And, she adds, "There was no way we could afford running the old main ring dc as a collider. That dominates the whole Tevatron thing."

This spring the accelerator experts will start trying to operate the Tevatron with a long, slow spill and at the same time extract some of the beam in 1-millisecond pulses, a tricky operation

that increases the risk of magnet quenching. However, a fast extracted beam is needed for neutrino experiments, scheduled to start next year.

The proton-antiproton collider project is well underway, according to John Peoples, who heads the project. The old main ring will receive a batch of protons from the old booster, accelerate it to 120 GeV, compress the batch in time and then the protons will hit a target to produce antiprotons.

Every 2 seconds the target is expected to yield 8×10^7 antiprotons. These are sent to two rings—first a debuncher and then an accumulator. Although the CERN collider at present has only an accumulator, CERN plans to add a debuncher called the Antiproton Collector. The Fermilab debuncher allows a much larger momentum spread initially so that more antiprotons can be collected in the end. However, Peoples points out, two rings take longer to make and are more complicated. By stochastic cooling, the debuncher is expected to reduce beam emittance in both transverse planes by a factor of three in 2 sec. Then the antiprotons

will be transferred to the accumulator, where stochastic cooling is to be done in the longitudinal plane. After 10 000 pulses of the Tevatron, the experimenters expect to stack 6×10^{11} particles in the accumulator. So in two hours, Peoples says, "we should have enough antiprotons to do colliding beams: 2×10^{11} antiprotons."

To produce collisions, 2×10^{11} antiprotons will be removed from the accumulator and injected into the main ring, accelerated to 150 GeV and injected into the Tevatron as three bunches traveling counterclockwise. Earlier, three bunches of protons will have been accelerated in the main ring and injected clockwise into the Tevatron. Both particle beams will then be simultaneously accelerated close to 1 TeV and stored. For a long beam lifetime, Fermilab will need a very reliable Tevatron.

Fermilab's goal is to accumulate enough antiprotons to allow a low-luminosity test of pp collisions in June 1985. At that time, Peoples hopes that with only a single bunch of protons and a single bunch of antiprotons, the

machine luminosity will exceed 10^{28} cm⁻²sec⁻¹; the Collider Detector Facility also is expected to have a shutdown run then.

CDF, the first of two detectors for the pp collider, is under construction at Fermilab and other locations. Roy Schwitters (Harvard) and Tollestrup are spokesmen for the detector collaboration, which has about 140 individual participants; the device will cost \$40 million. A second detector (whose cost is about \$25 million), to be placed in the D0 region of the collider, has not yet been approved by DOE, but the physics itself has been approved by Fermilab. Spokesman for the D0 collaboration is Paul Grannis (Stony Brook). Both detectors are about the size of UA1 and UA2 at the CERN collider and more complex than UA1. The collision hall for CDF has been finished and the construction of the D0 area will begin in 1985, after the first pp collisions occur. The shutdown for the D0 area would end February 1986. By summer or fall of that year, Peoples hopes the first physics run with pp collisions will occur. —GBL

Relativistic treatment of low-energy nuclear phenomena

Because the binding energies of nuclei are very much smaller than their rest masses, one would not have expected relativistic effects to play a significant role in nuclear structure or nuclear scattering at modest energies. Thus the nonrelativistic Schrödinger equation has until recently been the basis for almost all calculations in traditional nuclear physics. But in the past three years, the coming together of precise new data and novel theoretical approaches has made it appear that a fully relativistic treatment is indispensable for the understanding of nuclear phenomena even at low energies.

In 1981, new capabilities at the Los Alamos Meson Physics Facility made it possible for the first time to measure in essentially complete detail the elastic scattering of polarized protons off spin-zero nuclei at energies where the free nucleon-nucleon scattering amplitudes are well known (up to 500 MeV). At these energies, one expected the elastic scattering to be well described by the impulse approximation, which calculates a complex "optical" scattering potential for the nucleus as a whole on the assumption that the incident proton is scattered by quasifree individual protons and neutrons in the nucleus, neglecting binding energies and other effects of the nuclear medium. The impulse approximation has been an important tool in the ongoing effort to understand nuclear phenomena in terms of the two-body interactions of

their constituents.

The new Los Alamos data, however, didn't seem to fit this generally accepted picture. Calculating the impulse approximation from the known free-nucleon amplitudes in the conventional nonrelativistic Schrödinger-equation formalism, the experimental groups found that they could not reproduce the three experimental functions that together give a complete description of the elastic scattering of polarized protons off spinless nuclei: the differential cross section, the left-right asymmetry as a function of scattering angle (called the analyzing power) and the spin-rotation function (which requires a determination of the spin orientation of the proton after the scattering). This apparent breakdown of the impulse approximation had not been noticed before 1981 because spin-rotation data were very sparse and because LAMPF was only providing proton beams at 800 MeV, where the free proton-neutron amplitudes are poorly known. The impulse-approximation calculations essentially involve no free parameters, but the incompleteness of the earlier data had introduced sufficient latitude to permit plausible fits.

The first of these 1981 experimental papers¹, describing the differential cross sections and analyzing power of 500-MeV protons scattered off various nuclear species by a University of Texas-Northwestern collaboration, speaks of "the breakdown of the im-

pulse approximation." "We are forced to consider the possibility of a ... fundamental theoretical inadequacy ... in the conventional application of the ... formalism." The inadequacy, the authors suggest, is the intrinsic failure of the impulse approximation to take sufficient account of the collective modifying effect of the nuclear medium on the quasi-two-body collisions. They believed they were seeing a significant and unexpected difference between free nucleon scattering and what happens inside a nucleus.

Schrödinger formalism. It now appears that the fault lay not in the impulse approximation, but rather in the Schrödinger equation itself. The Schrödinger equation is of course a nonrelativistic approximation to the real world, where the Dirac equation is presumed to provide a relativistically correct description of the scattering of protons and neutrons. Although the energies in question here are not very relativistic and the trivial effects of relativistic kinematics have always been taken into account at intermediate energies, the crucial issue that requires the relativistic treatment appears to be the spin dependence of the nucleon-nucleon interaction. Spin is, after all, an intrinsically relativistic phenomenon. Historically, the first great triumph of the Dirac equation was its explanation of the spin and magnetic moment of the electron.

The large spin-dependent effects ob-