

# High-power multimode X-band rf pulse compression system for future linear colliders

Sami G. Tantawi, Christopher D. Nantista, Valery A. Dolgashev, Chris Pearson, Janice Nelson, Keith Jobe, Jose Chan, Karen Fant, and Josef Frisch

*Stanford Linear Accelerator Center, Menlo Park, California 94025, USA*

Dennis Atkinson

*Lawrence Livermore National Laboratory, Livermore, California 94550, USA*

(Received 2 November 2004; published 7 April 2005)

We present a multimode X-band rf pulse compression system suitable for a TeV-scale electron-positron linear collider such as the Next Linear Collider (NLC). The NLC main linac operating frequency is 11.424 GHz. A single NLC rf unit is required to produce 400 ns pulses with 475 MW of peak power. Each rf unit should power approximately 5 m of accelerator structures. The rf unit design consists of two 75 MW klystrons and a dual-moded resonant-delay-line pulse compression system that produces a flat output pulse. The pulse compression system components are all overmoded, and most components are designed to operate with two modes. This approach allows high-power-handling capability while maintaining a compact, inexpensive system. We detail the design of this system and present experimental cold test results. We describe the design and performance of various components. The high-power testing of the system is verified using four 50 MW solenoid-focused klystrons run off a common 400 kV solid-state modulator. The system has produced 400 ns rf pulses of greater than 500 MW. We present the layout of our system, which includes a dual-moded transmission waveguide system and a dual-moded resonant line (SLED-II) pulse compression system. We also present data on the processing and operation of this system, which has set high-power records in coherent and phase controlled pulsed rf.

## I. INTRODUCTION

Recently, ultra-high-power rf systems at X-band and above have received a lot of attention at different laboratories around the world because of the desire to design and construct a next-generation linear collider. For a review of these activities the reader is referred to [1,2]. These systems are required to generate and manipulate hundreds of megawatts of pulsed rf power. Standard rf components that have been in use for a long time, such as waveguide bends, directional couplers, and hybrids, cannot be used directly because of peak field considerations. Indeed, one has to reinvent most of these components, taking into account the constraints imposed by ultra-high-power operation. As with accelerator structures, these components are usually made of oxygen-free, high-conductivity copper, and operation takes place under ultrahigh vacuum conditions. Based on experimental work at X-band, we adopted as a design criterion that the peak surface electric fields should not exceed 500 kV/cm [3]. Peak magnetic field should be limited so that the pulsed surface heating does not exceed 30 °C [4].

To reduce losses and to enhance power-handling capabilities, one must use overmoded waveguide. Manipulating rf signals in highly overmoded waveguide is not trivial. The design of even simple components, such as bends, is complicated by the need to ensure the propagation of a single mode without excessive losses due to mode conversion.

Most proposed designs for future linear colliders contain long runs of waveguide. In X-band, room temperature designs, these runs total on the order of 100 km or more for the entire machine. To reduce the length of waveguide involved, we suggested multimoded systems [5]. In these systems, waveguide is utilized multiple times, often carrying different modes simultaneously. One might expect that this would lead to great complications in the design of most rf components. However, in an overmoded component, making an additional propagating mode an operating mode is helped by the fact that its coupling to the other operating mode(s) must be avoided anyway. The impact on the mechanical design can also be modest; most of these components are compact, usually less than four wavelengths in size.

Manipulation of multiple modes in a single component can be easier in rectangular waveguides than in circular waveguides. The philosophy behind many of our designs is to do all the manipulation in two dimensions, i.e., planar designs. This leaves the height of the waveguide as a free parameter to reduce the fields. To transport the rf power, however, we need to use low-loss circular waveguide. To take advantage of this, we present a rf taper which maps modes in circular waveguide into modes in rectangular guide. The modes used are  $TE_{10}$  and  $TE_{20}$  in rectangular waveguide and  $TE_{11}$  and  $TE_{01}$  in circular waveguide. We present the design methodologies for these components, many of which handle two modes simultaneously. The designs feature smooth transitions and avoid sharp edges

and narrow slots to minimize field enhancement. At the same time, losses are kept very low.

We also present a pulse compression system based on our components. The energy is stored in dual-moded resonant delay lines [6]. Since the output pulse width is directly proportional to the delay line length, dual moding of the delay lines allows them to be about half as long as they otherwise would need to be. Dual moding of the transport line, on the other hand, allows power to be directed through a pulse compression path or a bypass path. Power is supplied to the pulse compressor by four 50 MW X-band klystrons run off a common 400 kV solid-state induction modulator. With this system, we have produced a peak rf signal of about 580 MW. In a reliability test, it ran for hundreds of hours at a power level of about 500 MW with pulse energy of more than 200 J. The repetition rate varied from 30 to 60 Hz. This fully demonstrates the power-handling capability of these designs. It also exceeds the previous state of the art [3] by increasing the pulse energy by more than a factor of 3 and the pulse power by more than 25%.

The main part of this paper is organized in three sections. In Sec. II, we briefly describe the design of individual components; i.e., the basic building blocks. Then, in Sec. III, we describe how these components are put together to form the major components of the pulse compressor and the layout of the pulse compression system. Section IV is devoted to the high-power experiment. We describe first our experience with rf processing the system and then the reliability test, an approximately 300 h continuous run to determine its operational robustness.

## II. DUAL-MODE COMPONENTS

In this section we describe the design of the dual-mode components. These are divided into two categories: planar components and circular waveguide components. By the term planar, we mean translationally invariant in the direction of the electric field. The height of the waveguide is irrelevant to the rf characteristics because the fields do not vary in the height direction (i.e., the wave number component normal to two parallel walls of the waveguide is zero).

The two types of components are connected with dual-mode circular-to-rectangular tapers.

### A. Planar combiner/splitter

By reflecting two  $H$ -plane mitered bends across the horizontal waveguide wall one gets a basic dual-mode device, a planar dual-mode combiner/splitter. This is shown in Fig. 1. The width of the oversized port 3 is chosen so that it can support only two modes,  $TE_{10}$  and  $TE_{20}$ . If the device is excited from both port 1 and port 2 with equal amplitudes and identical phases, the  $TE_{10}$  mode is excited in the oversized port 3. If the phase of the excitation in port 1 and port 2 differs by  $180^\circ$ , the  $TE_{20}$  mode is excited in port 3.

These two excitation conditions satisfy, respectively, the condition of a perfect magnetic wall and a perfect electric wall along the symmetry plane bisecting port 3. For the latter case, the junction was matched by adjusting the miter angle and dimensions; for the former case, the junction was matched by the insertion of the oblong post at the symmetry plane, which is designed not to affect the other match. Because any excitation of ports 1 and 2 is a superposition of these two patterns, the two ports are independently matched. By the unitarity of the  $4 \times 4$  scattering matrix representing this device, it follows that the two modes of port 3 are also matched and that port 1 and port 2 are perfectly isolated.

### B. Planar hybrid

If one reflects the dual-mode combiner/splitter around port 3, one realizes a four-port directional coupler with an arbitrary coupling ratio. The coupling ratio depends on the drift length of the common dual-moded waveguide section. One can simplify this device by eliminating the matching obstacles and arranging for the two junctions to cancel each other's mismatch. We presented this type of 3-dB hybrid for the first time in [7], to which the reader is referred for more details. We also reported on a related 8-port superhybrid in [8]. A directional coupler based on the variable coupling idea was also suggested by Kazakov [9].

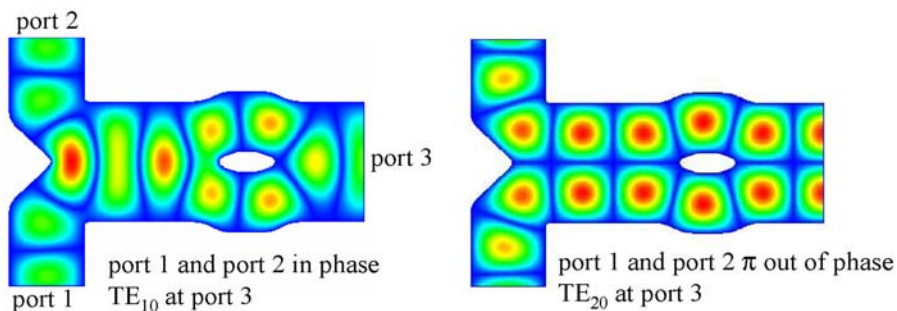


FIG. 1. (Color) HFSS [21] simulation of the basic planar dual-mode combiner/splitter.

### C. Manipulation of modes with planar structures

In the previous section, we presented a method of exciting the  $TE_{10}$  and  $TE_{20}$  modes in planar waveguide. In this section, we present a set of components which couple these modes together. As shown below, these will prove to be indispensable in our system design.

#### 1. Jog mode mixer

An  $H$ -plane bend in a waveguide which propagates both the  $TE_{10}$  and  $TE_{20}$  modes will result in coupling between these two modes. Figure 2 shows a structure composed of two bends back to back. Given our standard dual-mode width, there are only two design parameters, the radius and angle of the bends, which are mirror symmetric. If the bends are gradual enough, backward scattering can be neglected. At an infinite radius we have a straight waveguide with no coupling between the modes. The radius is decreased until a pure incident mode in port 1 results in an equal power mixture of the two modes in port 2.

#### 2. Jog mode converter

We presented this device in [5]. Basically, the radius of the two back-to-back bends of the mode mixer geometry, described in the previous subsection, is chosen such that the powers in each mode are exactly equal at the end of the first bend when a single mode is incident. By the addition of a short straight section between the two bends, the phase between the modes is then adjusted so that full conversion from one mode to the other is accomplished through the device.

#### 3. Bend mode converter

The bend mode converter is very similar to the jog mode converter and is intended for use where mode conversion and a  $90^\circ$  bend are needed. For this device, the direction of the second  $45^\circ$  bend is reversed. The resulting change in coupling sign is then taken into account by lengthening the straight section to readjust the relative phase of the modes.

### D. Dual-mode circular-to-rectangular taper

In this section we present a smooth transition from rectangular to circular waveguide that preserves reflection symmetries. The S-matrix of the total transition connects modes of the same symmetry. This transition is sufficiently

adiabatic that it preserves the TE (or TM) character of modes, introduces negligible reflections, and, in the absence of degeneracy, makes one-to-one, order-preserving modal connections. This transition enables us to carry out basic rf manipulations with the more easily handled over-moded rectangular waveguide modes. We were able to perfect this one taper for both operating modes.

#### 1. Smooth taper geometry

When tapering from one cross section shape, e.g., a circle, to another shape, e.g., a rectangle, the length of the transition  $l$  and the loci of connecting points between the two shapes uniquely define a taper. In cylindrical coordinates a shape  $i$  centered on and lying in a plane perpendicular to the  $z$  axis can be described by a relation  $r_i(\phi)$ , which gives the radius as a function of the angle. The linear taper between two shapes  $r_1(\phi)$  and  $r_2(\phi)$  is then given by  $r(\phi, z) = r_1(\phi) + [r_2(\phi) - r_1(\phi)]z/l$ .

Such a taper is compatible with the process of wire electron discharge machining when the two heads of the machine are moving synchronously with the same angular speed. More complicated tapers are constructed from a set of cascaded subtapers, each having this linear form.

#### 2. Dual-mode taper

Adiabaticity is sufficient to map the  $TE_{11}$  mode in circular guide to the  $TE_{10}$  mode in square waveguide, as is well known. However, it is desirable to convert the  $TE_{01}$  mode in the circular guide to a single polarization of  $TE_{02}/TE_{20}$  in the square guide. Modifying the square waveguide to a rectangular waveguide to break the degeneracy between the  $TE_{02}$  and  $TE_{20}$  modes would permit this. However, in this case, the length of the taper required to achieve an adiabatic transition to a single mode in the rectangular guide is somewhat excessive. If one chooses the dimensions of the square waveguide and the circular waveguide such that the  $TE_{02}$  circular mode and the  $TE_{22}$  and  $TM_{22}$  rectangular modes are cut off, the length of such a taper would be approximately 18 cm at 11.424 GHz, or about seven wavelengths.

Instead, we construct a mode transformer from three sections as shown in Fig. 3. The middle section is a cylinder with the shape  $r_2(\phi) = r_0(1 + \delta \cos 2\phi)$ , where  $r_0$  is the radius of the circular guide and  $\delta$  is a design

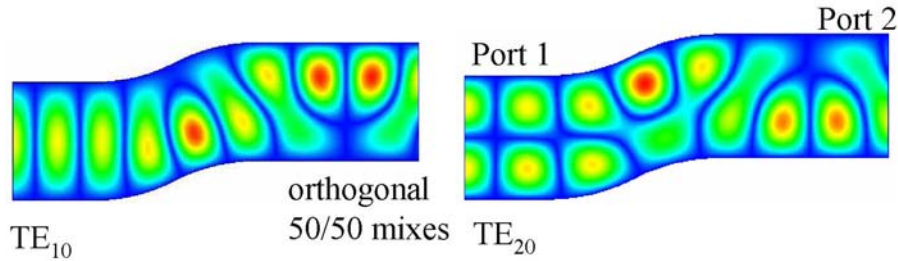


FIG. 2. (Color) HFSS simulation of the jog mode converter. The colors represent a time snapshot of the electric field amplitude.

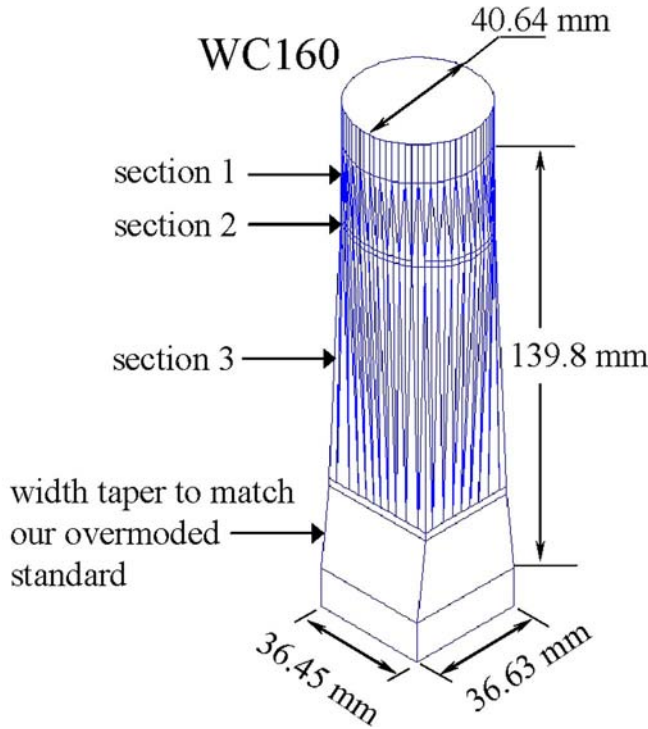


FIG. 3. (Color) Dual-mode taper geometry.

parameter that determines the deviation of the this shape from a circle.

The taper from the circle to the intermediate shape scatters the  $TE_{01}$  mode into two modes,  $M_1$  and  $M_2$ , in the intermediate section. The taper between the rectangular waveguide to the intermediate shape scatters the rectangular mode  $TE_{02}$  into the same  $M_1$  and  $M_2$  modes. The lengths of both tapers are adjusted such that the coefficients of the scattered modes  $M_1$  and  $M_2$  are of the same amplitude from both sides. Since  $M_1$  and  $M_2$  have different propagation constants in the intermediate section, the length of that section can be adjusted so that the circular  $TE_{01}$  mode gets completely converted into the rectangular  $TE_{02}$  mode through the composite cascaded transition.

The idea of this design was first reported by our group in [10], and later implemented to split the output of a  $TE_{01}$  gyrokystron [11]. However, in both cases, no care was taken to ensure the adiabatic propagation of the  $TE_{11}$  mode. Because of the odd shape in the middle of this taper, the taper tended to couple the rectangular  $TE_{01}$  to the circular  $TE_{31}$  mode. This is to be expected, as that shape essentially contains a second-order azimuthal deformation, which couples modes that differ in azimuthal index by 2. Conversely, the circular  $TE_{11}$  mode is also coupled to  $TE_{12}$  in the rectangular guide.

The design of Fig. 3 was the result of several iterations on the design, increasing the length until a perfect conversion between one polarization of the  $TE_{11}$  circular mode and the  $TE_{01}$  (not  $TE_{10}$ ) rectangular mode was achieved. Each time the length was increased, the full design was

repeated to ensure perfect conversion between the  $TE_{02}$  rectangular mode and the  $TE_{01}$  circular mode. Then the  $TE_{11}$  circular to  $TE_{01}$  rectangular conversion was checked. Experimental measurements of this taper, for both modes, are shown in Fig. 4.

### E. Dual-mode directional coupler

With the use of overmoded circular waveguide carrying both the  $TE_{11}$  mode and the  $TE_{01}$  mode, we need to monitor the power in each mode separately. A directional coupler for each mode is needed. The coupling not only has to discriminate between forward and backward waves for a particular mode, but also has to discriminate strongly against all other forward and backward modes that propagate in the waveguide, whether excited intentionally or not. In this section, we detail the design of a combined directional coupler that monitors forward and backward waves of both the  $TE_{01}$  and  $TE_{11}$  modes with high directivity and mode selectivity.

The theoretical ideas behind this design were presented for the first time in [12]. The methodology normally used in designing mode selective directional couplers for high-power, overmoded guides discriminates against other modes by proper choice of distance between coupling slots [13]. As the number of modes inside the original guide is increased, the number of holes and the length of the coupler are increased. This makes the coupler hard to manufacture and very sensitive to the accuracy of slots positions and sizes. When this problem was studied for communication systems, it was possible to make the wavelength of the dominant mode in the sidearm equal to that of the main guide mode. In this case, it was realized that natural mode discrimination occurs independent of slot

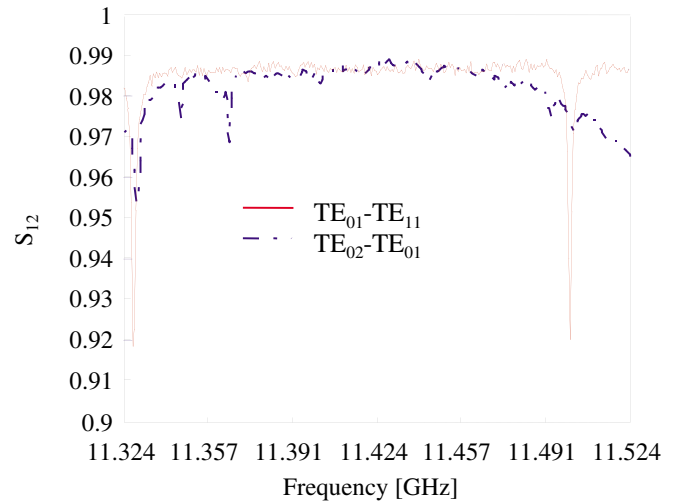


FIG. 4. (Color) Measured  $S_{12}$  between the rectangular  $TE_{02}$  mode and the circular  $TE_{01}$  mode, and rectangular  $TE_{01}$  mode and the circular  $TE_{11}$  mode. These measurements include the response of mode transducers necessary to launch the mode at each end of the taper.

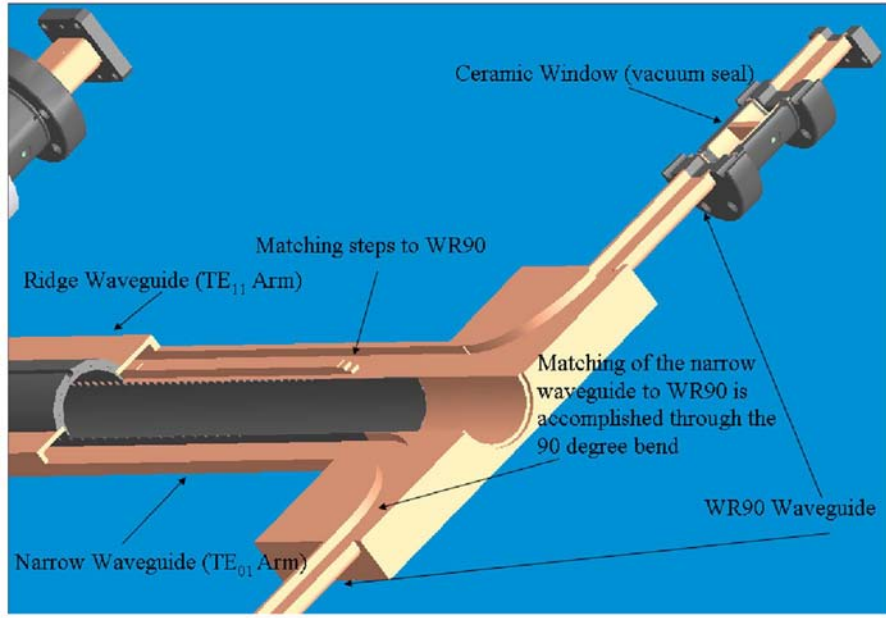


FIG. 5. (Color) Inner structure of the dual-mode directional coupler for system diagnostics.

positions [14]. However, before [12] there was no theory developed for designing couplers in a manner independent of slots positions.

The design of the dual-mode directional coupler is shown in Fig. 5. Two waveguides are coupled through sets of holes to the main circular waveguide. To couple the  $TE_{01}$  circular mode, the width of one rectangular waveguide is adjusted to match the phase velocities. Because the circular waveguide is overmoded and the fundamental  $TE_{11}$  mode has a phase velocity close to the speed of light, one cannot match its velocity in a single-moded rectangular waveguide. To match the velocities, we had to use a ridged waveguide.

To make the coupler directional and to discriminate against coupling of unwanted modes, the coupling hole pattern was chosen according to the above-mentioned theory. Figure 6 shows the calculated response of the sidearms for all modes that can be excited in the circular guide.

Finally, one has to bend the rectangular and ridged waveguides so that one can connect vacuum windows and diagnostic ports. These bends double as tapers to standard size rectangular waveguide, WR90 in our case. The proper matching of these bends and the attached vacuum window is crucial for maintaining good directivity.

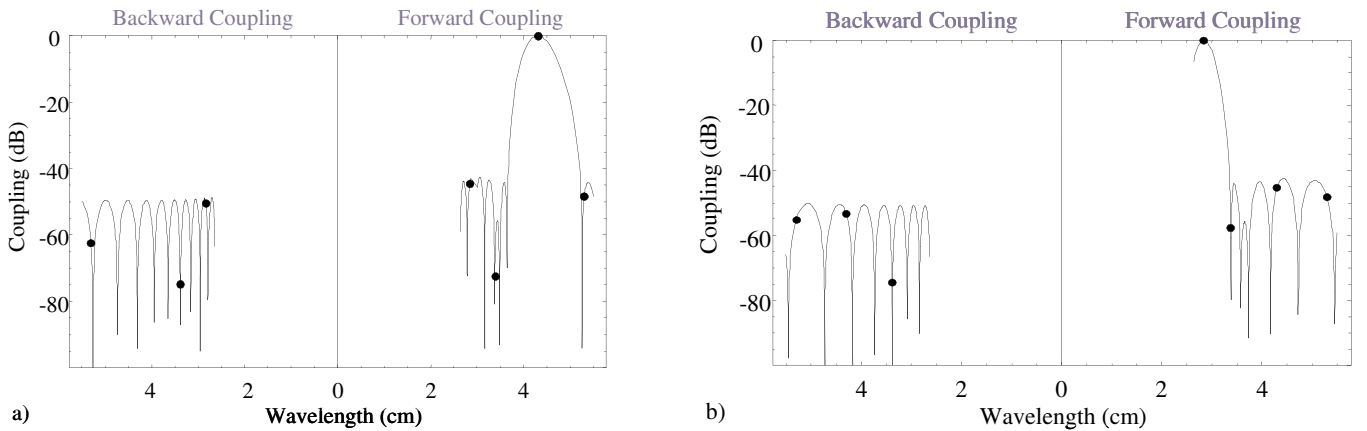


FIG. 6. (Color) Coupling versus guided wavelength. (a)  $TE_{01}$  directional coupler response, (b)  $TE_{11}$  directional coupler response; both normalized to forward coupling of the  $TE_{01}$  mode and  $TE_{11}$  mode, respectively.  $TE_{01}$  coupler length is 50.6 cm. Number of coupling holes is 44.  $TE_{11}$  coupler length is 38.3 cm. Number of coupling holes is 51. Points represent  $TE_{11}$ ,  $TE_{21}$ ,  $TE_{01}$ ,  $TE_{31}$ . The  $TE_{01}$  mode and the  $TE_{11}$  mode are not shown in the backward coupling because they are theoretically zero because of the choice of coupling holes spacing.

### III. DUAL-MODED SLED-II PULSE COMPRESSION SYSTEM

In this section, we start by describing a set of composite supercomponents based on the elements described in Sec. II. From this set of supercomponents, we assemble the main body of our pulse compression system. The pulse compressor is of the SLED-II type [15], a passive waveguide circuit which stores rf energy in a pair of resonantly tuned delay lines for several roundtrips, until a phase reversal in the input pulse causes it to be expelled in a compressed pulse of increased power. The lines are iris coupled at one end to a power directing hybrid and shorted at the other. We then describe the dual-mode storage delay lines. Finally, we describe the overall system integration and the system cold tests.

#### A. Combiner/mode launcher

The combiner/mode launcher is shown in Fig. 7. It comprises the splitter/combiner described in Sec. II A, the circular-to-rectangular taper described in Sec. II D, and rectangular height tapers.

To test the functionality of this device, we connected it to two different mode converters at the circular port (port 3 in Fig. 7). The  $TE_{01}$  mode was launched through the circular port of the taper using a wrap-around mode converter [3]. The output of the taper then launches the  $TE_{20}$  mode in the symmetry port of the combiner. The signal is split equally out the remaining two ports. The results of these measurements are shown in Figs. 8(a)–8(d). From Fig. 8(a), one can estimate the total loss of the combiner

with brazed port tapers to be about 1.2%. Because of the asymmetry of the  $TE_{20}$  mode, one should observe a  $180^\circ$  phase difference between the two ports. This is confirmed, within the measurement and fabrication tolerances, in the experimental data shown in Fig. 8(c). A short, commercial mode transducer was used for the  $TE_{11}$  measurements. Both modes were matched better than  $-30$  dB.

#### B. Dual-mode superhybrid

The main waveguide junction of our high-power system, the dual-mode superhybrid, shown in Fig. 9, was designed using our overmoded rectangular waveguide planar design motif. The planar body of this super component is a four-port device whose ports support the  $TE_{10}$  and the  $TE_{20}$  modes (actually  $TE_{01}$  and  $TE_{02}$  in this device, if one considers that the height is greater than the width). The circular-to-rectangular taper described above allows these ports to be connected to low-loss circular waveguide with one-to-one mode mapping. The planar design allows us to increase power-handling capability of the superhybrid by building it overheight. Additional propagating modes that this introduces need not be taken into account in our manipulations, as there is no mechanism for coupling to them.

The purpose of the superhybrid is to pass the  $TE_{20}$  mode from the input port to the two resonant delay lines of a SLED-II pulse compressor [3,15,16], combining identical reflections from these ports out the fourth port, while also allowing the rectangular  $TE_{10}$  mode to be transmitted from input to output port by a different path, bypassing the pulse

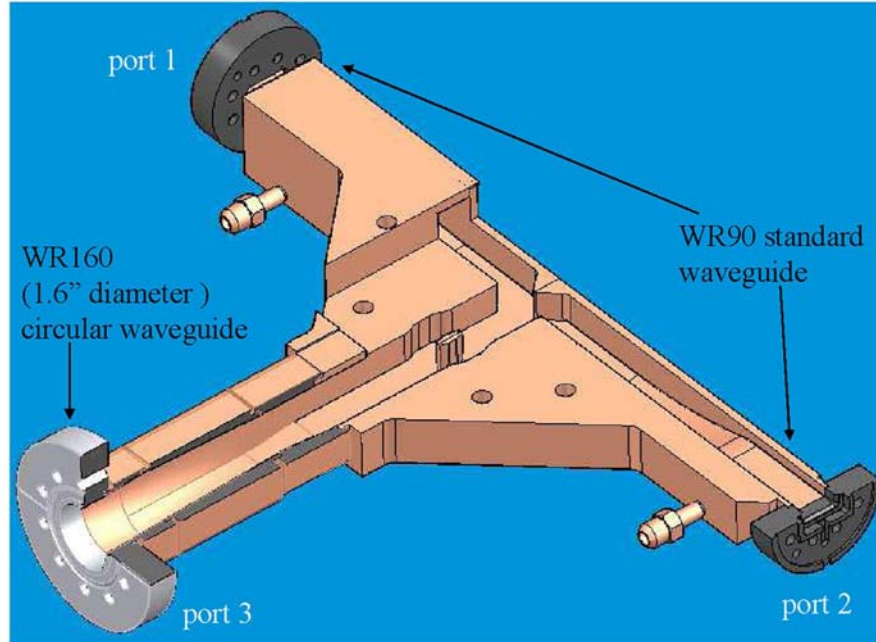


FIG. 7. (Color) Cutaway view of the mode launcher/combiner. Depending on the phases and amplitudes of the waves incident in the rectangular ports either or both of the  $TE_{01}$  and  $TE_{11}$  modes can be excited in the circular port.

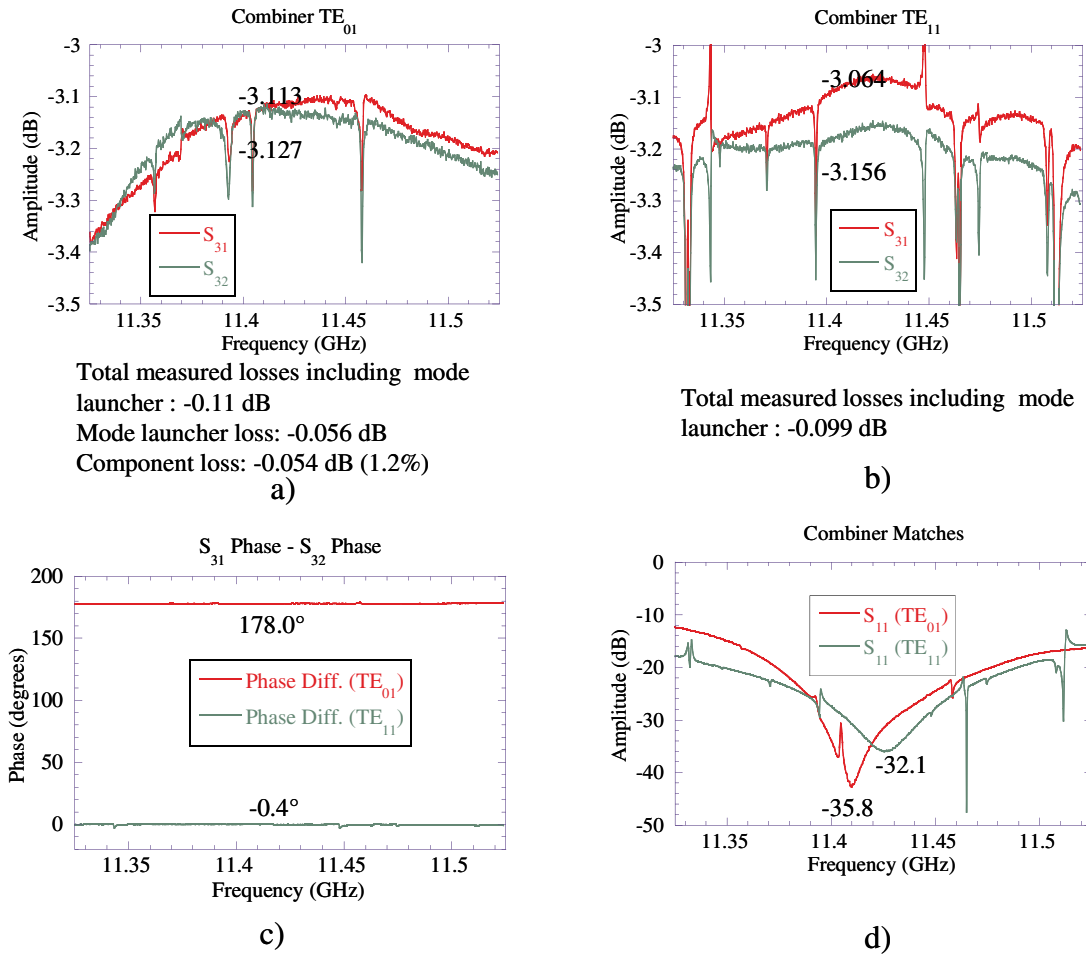


FIG. 8. (Color) Cold test results of the dual-mode combiner. The port indexing (scattering parameter subscripts) is in reference to Fig. 7. All measurements include the response of a cold test circular mode launcher.

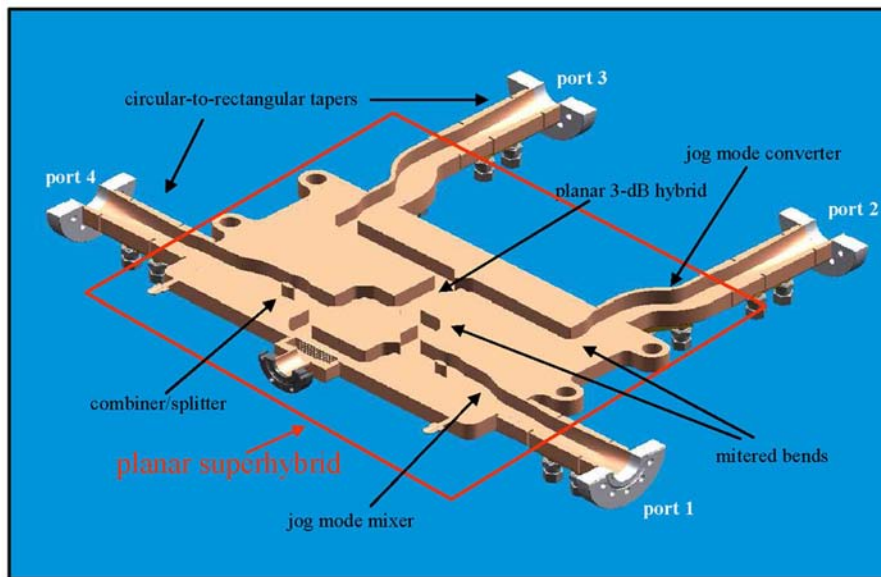


FIG. 9. (Color) The design of the dual-mode superhybrid (a split view).

compressor. While the planar design allowed the superhybrid to be machined out of a single block of copper, like a circuit on a substrate, it is actually composed of several subcomponents, designed individually.

The input port leads directly into a mode mixer, a slight jog or dogleg described above, that couples the two modes, converting a pure input of either  $TE_{10}$  or  $TE_{20}$  into an equal mixture of the two. This is followed by a dual-mode splitter/combiner T junction, which divides power from either pure mode equally between two single-mode (in the horizontal plane) ports. Proper spacing between mode mixer and splitter allows the fields from each of the orthogonal mode combinations to cancel in one splitter arm and add in the other. In combination, these two subcomponents serve to direct the power of the two possible input modes along different paths.

A pair of mitered bends connects the power input as  $TE_{10}$  into an arm of a mirror image of the splitter [see Fig. 10(b)]. The mixed output at its overmoded port is then converted back to pure  $TE_{10}$  at the output port of the superhybrid by a mirror image of the mode mixer. Perforation holes in the outer wall of the short waveguide section between mitered bends allows vacuum pumping in the heart of the device.

Power input as  $TE_{20}$  follows a similar path on the opposite side of the device [see Fig. 10(a)]. Here, the section between mitered bends is occupied by a planar hybrid of the magic- $H$  type. Its coupled ports lead to mitered bends, which turn perpendicular to the power flow of the superhybrid at a spacing designed to accommodate parallel, highly overmoded circular waveguide delay lines (see Fig. 9). Short width tapers (optimized with blended arcs) return to overmoded waveguide.

These are followed by jog converters, which completely convert  $TE_{10}$  power from the magic- $H$  ports to  $TE_{20}$  at the delay line ports of the device. This allows the rectangular-to-circular tapers to couple these ports to low-loss circular  $TE_{01}$  mode delay lines.

Cold test data for the dual-mode superhybrid are shown in Fig. 11. These measurements include the response of the mode transducers used to launch the appropriate modes at each port. They were performed at the circular ports of the fully brazed device. For the circular  $TE_{01}$  mode power to be properly directed forward through the hybrid, the reflections from the delay line ports (ports 2 and 3 in Fig. 9) seen at the hybrid must be identical in amplitude and phase. A phase length difference in the arms between the hybrid and the reflector can degrade the performance. Figure 11(a) shows an initial  $TE_{01}$  reflection measurement, with the delay line ports shorted, and a final measurement made after correcting such a difference with customized flanges. Figure 11(b) shows the circular  $TE_{01}$  mode transmission. To verify that the circular  $TE_{11}$  mode is decoupled from the delay line ports, its transmission and reflection were measured for three cases: both delay line ports shorted, one physically open, and both open. The results are shown in Figs. 11(c) and 11(d). Additional resonances in our setup are apparent in the shorted case, but the overall response is essentially unchanged.

### C. Dual-mode load system

To absorb the output power, one needs a load that accepts both the  $TE_{01}$  mode and the  $TE_{11}$  mode. This is accomplished by using a splitter similar to the combiner described in Sec. II A. First, the power in either mode is split into two arms, as shown in Fig. 12. Irrespective of the

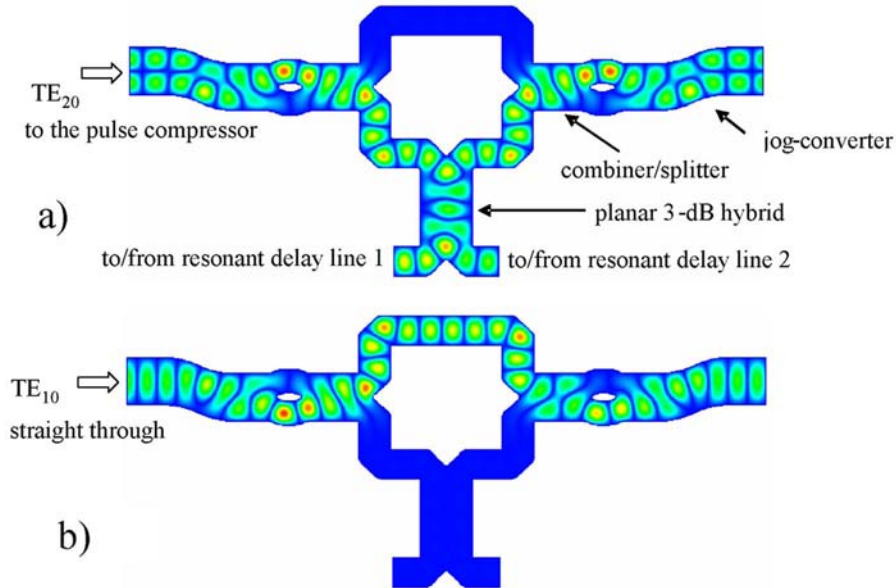


FIG. 10. (Color) The geometry and simulated electric field snapshots of the planar superhybrid with delay line ports shorted for (a)  $TE_{20}$  input and (b)  $TE_{10}$  input.

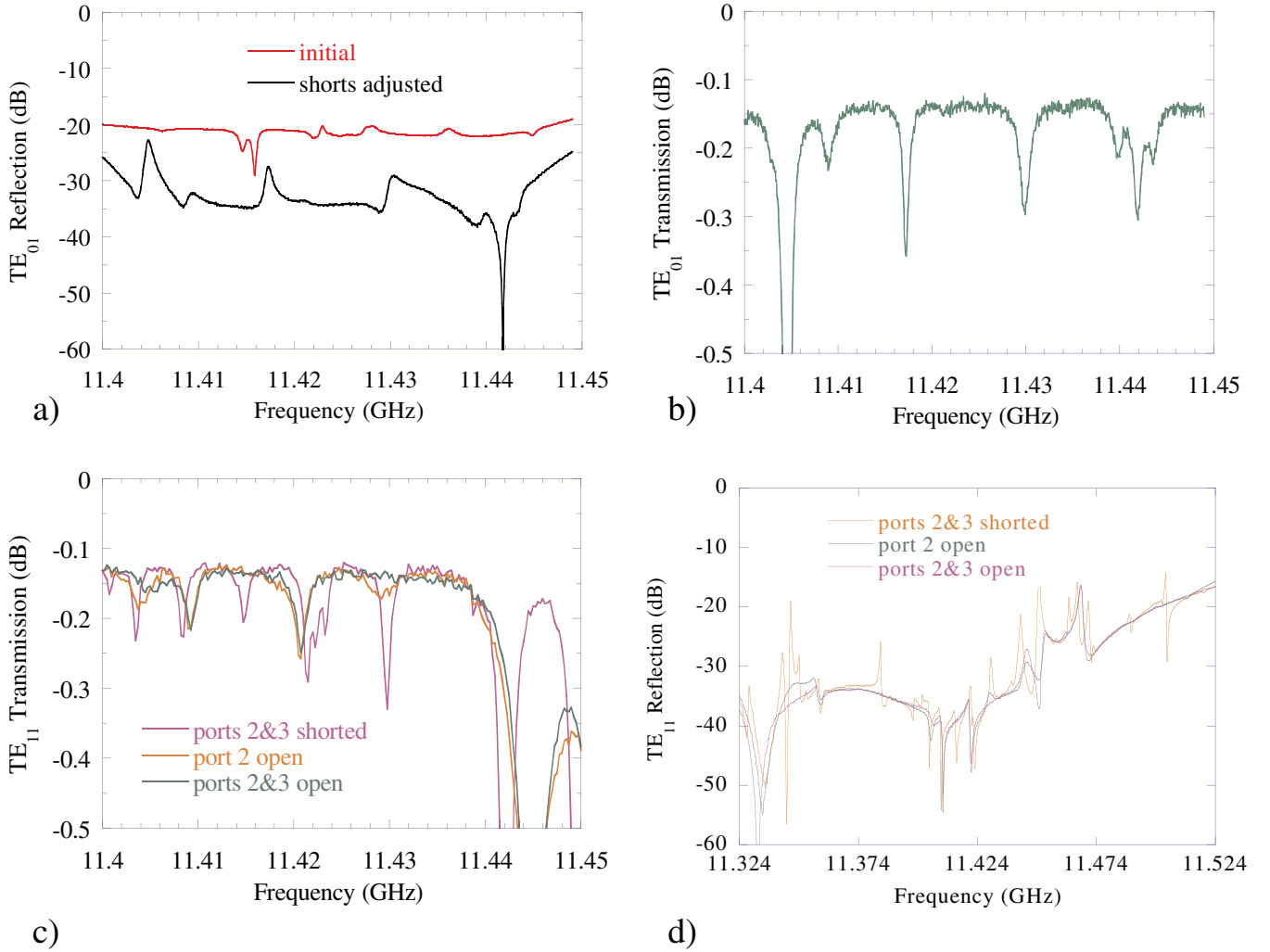


FIG. 11. (Color) Superhybrid cold tests. (a) Reflection measurements ( $S_{11}$ ) for the circular  $TE_{01}$  mode from port 1 while shorting both port 2 and port 3 and matching port 4. A 0.022 in. adjustment of the relative short positions significantly improved the match. (b) Circular  $TE_{01}$  mode transmission between port 1 and port 4 ( $S_{41}$ ) while shorting both port 2 and port 3 (shifted up 0.09 dB to account for loss of mode launchers). (c) Circular  $TE_{11}$  mode transmission between port 1 and port 4 ( $S_{41}$ ) for different conditions at port 2 and port 3. (d) Circular  $TE_{11}$  mode reflection ( $S_{11}$ ) for different conditions at port 2 and port 3.

incident mode, the two outputs of the splitter carry the  $TE_{10}$  rectangular mode. Mode converter bends are then used to convert that mode into the  $TE_{20}$  mode (see Sec. II C 3). The power in each arm is further split between four different waveguides. The total power is thus distributed between eight standard waveguides. These waveguides are terminated with all-metal, fundamental-mode loads, which are water cooled [17].

The load system is shown in Fig. 13. It splits the power eight ways, irrespective of the incident mode. Cold test results of the system for the two incident modes are presented in Fig. 14.

#### D Pulse compression delay lines

The theory for the dual-moded delay lines is detailed in [6]. The theory for pulse compression using resonant delay

lines can be found in [3,15,16,18]. Here we briefly describe these concepts. We then consider the practicalities involved in building these lines. We also describe the cold tests and tuning procedures used during the installation phase of these lines.

##### 1. Dual-moded delay lines

The pulse compressor consists of two highly over-moded, resonant delay lines, iris coupled to two ports of the superhybrid (see Sec. III B). These delay lines, roughly 30 m long in 17 cm diameter circular waveguide, are also dual moded. Here dual moding refers to the fact that they are designed to operate in both the circular  $TE_{01}$  mode and the circular  $TE_{02}$  mode simultaneously. A mode converting reflector at the end of each line transfers power between these two modes.  $TE_{02}$  is cut off at the input taper of each

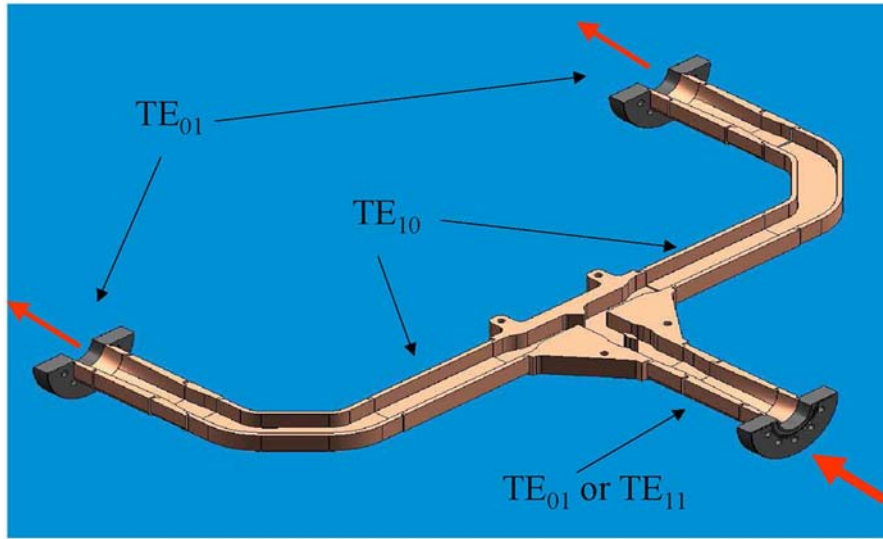


FIG. 12. (Color) Dual-mode splitter. Irrespective of the incident mode, the power is split between two outputs, each of which launches the  $TE_{01}$  mode.

line, which is carefully designed to perfectly reflect this mode without mixing. Because of reciprocity, the reflected  $TE_{02}$  wave gets converted back to  $TE_{01}$  at the far end of the line. Finally, this mode travels back and exits the line. Figure 15 illustrates this process. As a result, it takes two round-trips between each time the wave can impinge on the coupling iris. This effectively doubles the delay time for a line of given length, allowing the system to be considerably

more compact than standard delay lines would be for the desired compressed pulse width. For resonant tuning of the lines, the reflectors are mounted on accurately centered, stepping-motor driven vacuum feedthroughs. With a 400 ns double-bounce delay, we get a compression ratio of 4 from our  $1.6 \mu s$  input pulse and a gain of approximately 3.

Transitioning between the moderately overmoded diameter (4.06 cm) of the superhybrid ports and the highly overmoded delay line waveguide diameter (17 cm) and between the latter and the intermediate  $TE_{01}$ – $TE_{02}$  mode converter diameter (8.255 cm) requires carefully designed

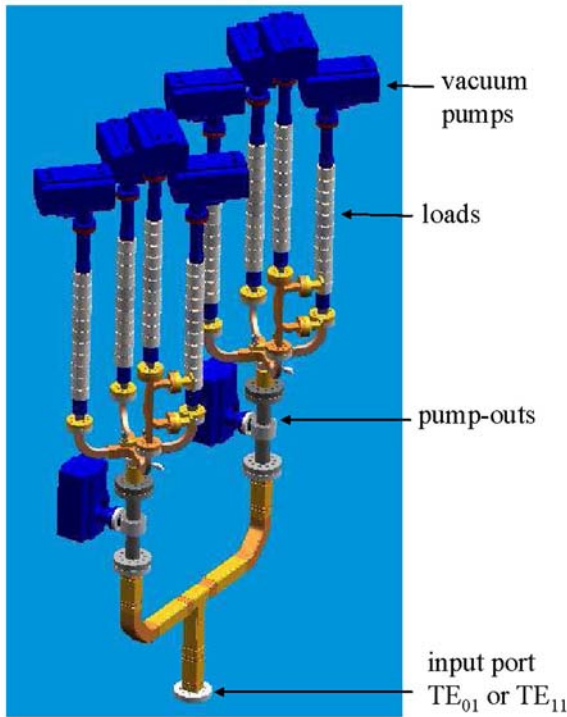


FIG. 13. (Color) The “load tree”: a dual-mode high-power load system.

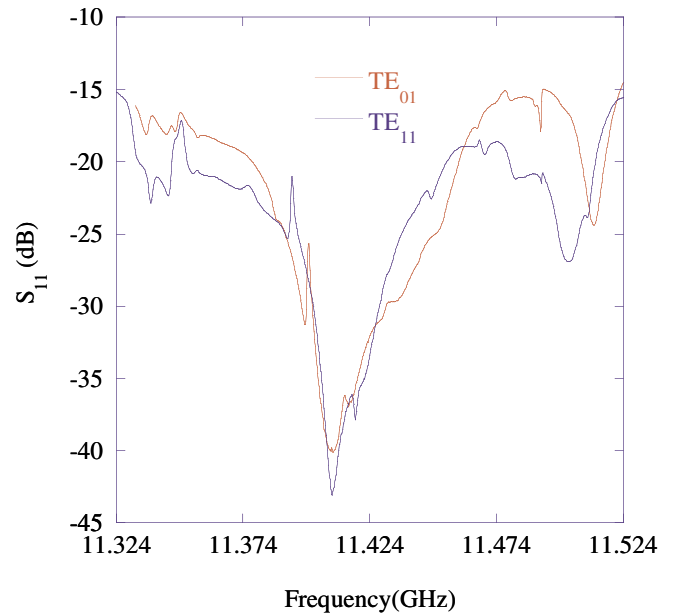


FIG. 14. (Color) Load tree response for different incident modes.

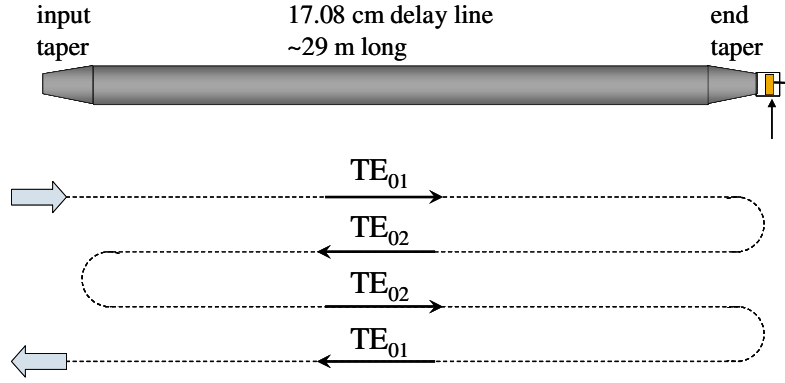


FIG. 15. (Color) Dual-mode delay lines.

circular tapers. This is due to the fact that there are six circular  $TE_{0n}$  modes supported by the delay lines. The approach of using an optimized set of many discrete steps, rather than long, adiabatic, smooth tapers, allows our designs to be quite compact ( $\sim 30$  cm). This compactness arises from the large number of free parameters inherent in this approach. The input taper must purely transmit  $TE_{01}$  and must purely reflect at the larger port  $TE_{02}$ , while the end taper must purely transmit both modes. The 32-step designs were produced using a mode-matching code and automated optimization. After the end taper, a movable mode converter is inserted. For the details of the taper designs and the end mode converter design, the reader is referred to [6]

Other delay line components include the coupling irises, which are machined from double-sided blank flanges, and vacuum pump outs at 17 cm diameter.

## 2. Delay line cold testing and tuning

We cold tested the individual dual-mode SLED-II delay lines with a network analyzer in conjunction with a time-domain program for producing compressed pulse waveforms [3,6]. We observed small but noticeable steps

in the middle of the double-bounce time bins, indicating mode impurities. Figure 16 shows time-domain waveforms simulated from frequency domain measurement data for the two delay lines. Gains from such measurements ran in the range of 2.9 to 3.2, out of an ideal (lossless) 3.44. Both lines together, with the superhybrid and directional coupler in the circuit, still produced gains of approximately 3.1.

The mode impurity evidenced by mid-time-bin discontinuities suggests imperfection in the juggling of power between  $TE_{01}$  and  $TE_{02}$ . To diagnose the cause of these small mid-time-bin steps in our SLED-II waveforms, we removed the iris, shortened our time pulse to 150 ns, and measured the amplitude and phase of the reflection seen after the first bounce, when no power should emerge. A typical result from such measurements is shown in Fig. 17. We see a small initial reflection after a time equal to one round-trip along the line and then the full reflection after two round-trips. Measurements were taken while the end mode converter was moved in 1 mm steps over a range of 24 mm. The data is plotted in Fig. 18.

The top delay line was found to have an impurity on the order of  $-26$  dB, while the bottom line was significantly

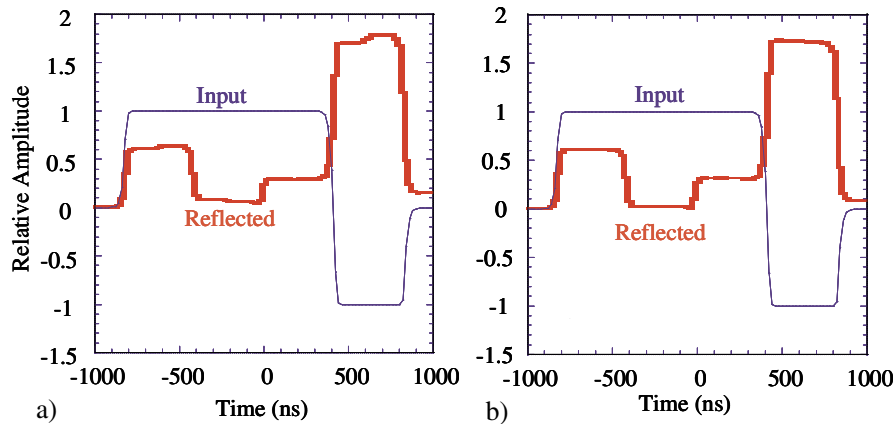


FIG. 16. (Color) Individual amplitude responses for (a) the upper resonant delay line and (b) the lower resonant delay line, measured separately (see Fig. 19). Small mid-time-bin steps indicate mode impurity.

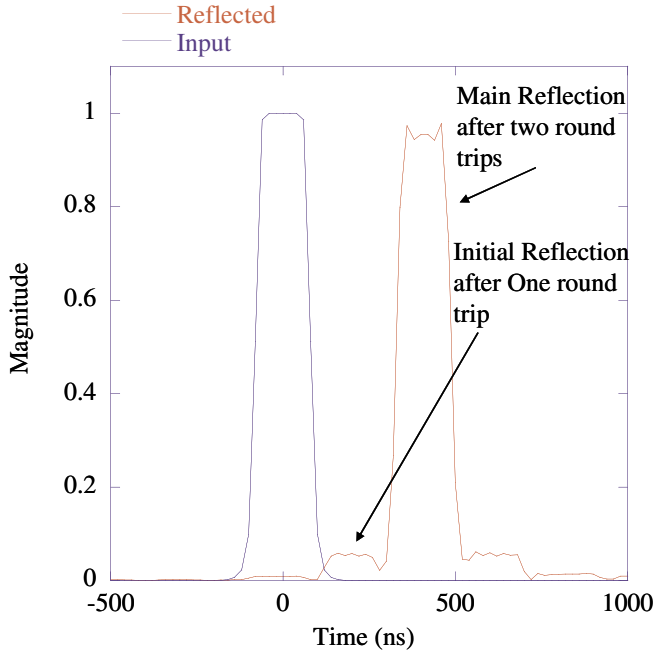


FIG. 17. (Color) A typical top delay line response without an input iris.

better, with an impurity of  $-34$  dB, as shown in Fig. 18(a). The small amplitude oscillation is attributable to interference with a cold test mode launcher mismatch of about  $-48$  dB. It is more prominent in the better-behaved bottom line, where this mismatch has greater effect on the combined phasor.

A mode impurity can come from two sources: imperfect conversion in the end mode converter, or mode scattering due to imperfections in the tapers and delay line. In the former case, as the end mode converter is moved a distance  $\Delta L$ , the phase of the power seen after the first bounce

would change as  $2\beta_{01}\Delta L$ ; the signal travels down as  $TE_{01}$  and remains unconverted on the return trip. For the latter case, the phase would change as  $(\beta_{01} + \beta_{02})\Delta L$ ; both the portion converted on the way down and the portion converted on the way back travel as each mode either before or after the end mode converter. Here the  $\beta$ 's are the mode wave numbers in the 8.255 cm diameter guide in which the end mode converter moves. The effective periods for these two cases are 14.2 and 16.1 mm, respectively. Figure 18(b) shows plots of the phase of the first bounce reflections for each line. Also plotted is the phase of the second bounce for the top line divided by two, which has the same rate of change expected for the first bin with taper/delay line errors. Finally, the phase change expected for end mode converter error is plotted. All phase plots are adjusted to start at  $180^\circ$  and then freely shifted by  $\pm 360^\circ$  to fit the scale. The measured phases follow very closely the  $(\beta_{01} + \beta_{02})\Delta L$  behavior, indicating unwanted conversions in the tapers and/or delay lines. Small excursions by the bottom line data are again attributable to the effect of mode launcher mismatch interference.

To solve this mode purity problem, the spacings of the irises from the delay line tapers were adjusted to keep the spurious first round-trip reflection as out of phase as possible during the charging time of the lines. Further, since the delay lines and the tapers all have different imperfections, the best combination of delay lines and tapers was found in order to minimize this undesirable side effect of multimoding. The problem was further circumvented by choosing the best resonance positions of the shorting plungers, so that parasitic couplings do not add in phase. Each plunger has a travel range encompassing three such positions. After these heuristic modifications we arrived at a system giving an output pulse with a reasonably flat top and a gain of about 3.1. To make the output pulse absolutely flat in high-power operation, we used a feedback

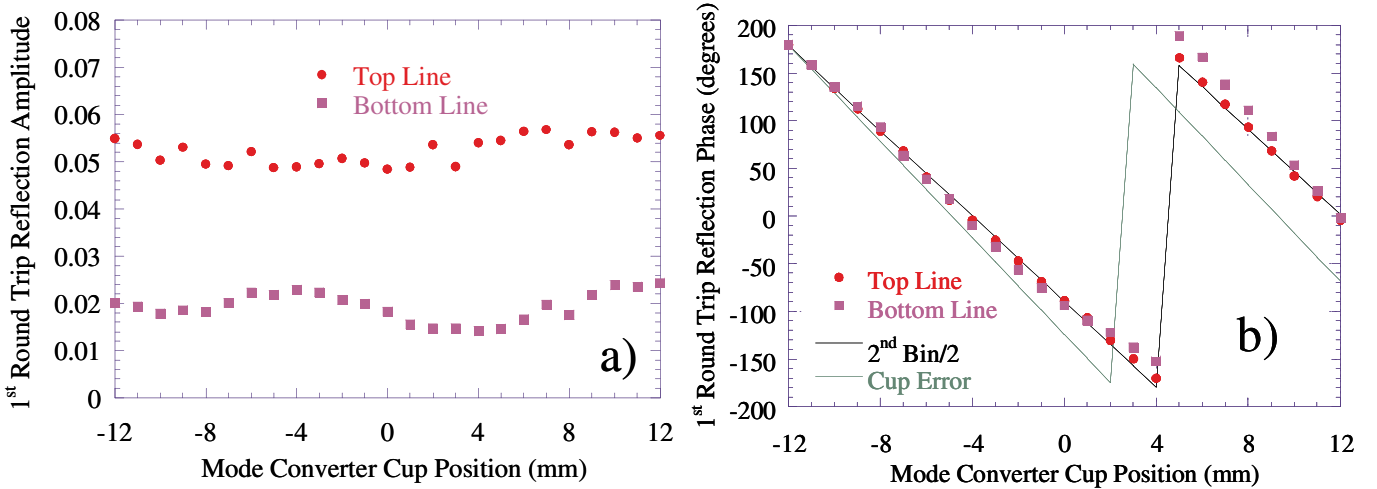


FIG. 18. (Color) Variations in (a) amplitude and (b) phase of spurious reflection (initial reflection after one round-trip) as a function of the end mode converter position.

system on the input to the klystrons. This feedback is discussed in Sec. IV B below.

### E. System cold tests

Here we describe the integrated system and report on the cold testing procedure as well as the tuning procedures performed on the whole system.

#### 1. System design

The components described above compose the heart of our dual-moded pulse compression system. The layout of

our system is shown in Fig. 19. The four XL4 klystrons [19] incorporated in this project, powered from a shared solid-state modulator [20] and driven in pairs through two traveling wave tube (TWT) amplifiers, offer a nominal 200 MW input at a pulse width of  $1.6 \mu\text{s}$ . The outputs of these klystrons are combined in pairs through planar hybrids [8]. The fourth port of each of these brings mis-phased/mismatched power into a high-power load. WR90 waveguide carries the combined rf from each of these klystron pairs to a port of the dual-mode combiner (Sec. III A). This combiner is a three-port device, whose third port launches power from the single-moded inputs

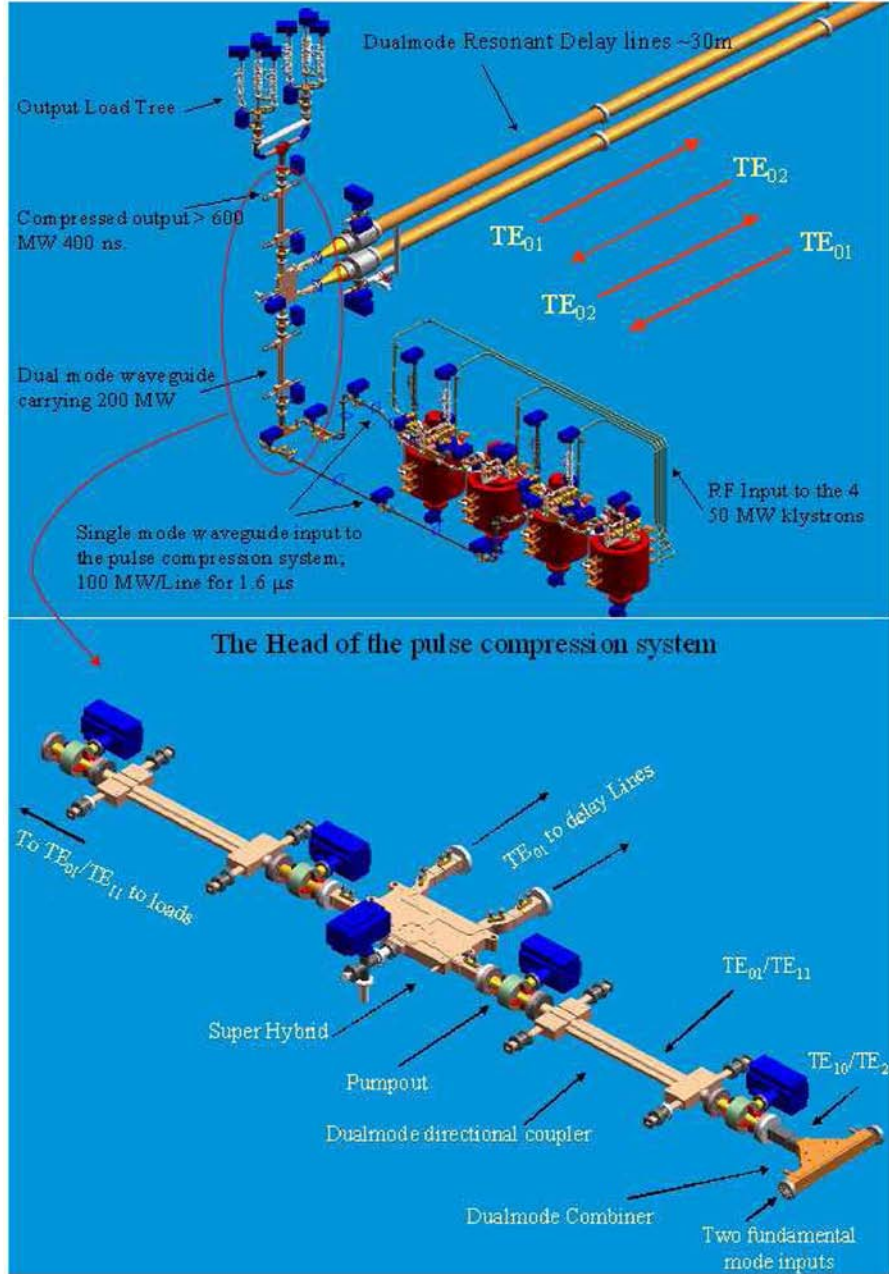


FIG. 19. (Color) Isometric layout of the dual-moded rf system with dual-moded SLED-II pulse compressor.

into overmoded circular waveguide in either the  $TE_{11}$  or  $TE_{01}$  mode, depending on the relative phase of the inputs.

The total combined power is fed into the superhybrid described in Sec. IIIB, with the klystron pairs phased to launch  $TE_{01}$  into the system, power is directed through the resonant-delay-line (also called SLED-II) pulse compressor [15]. The first original publication of the scheme is in [16]. These lines, as described above, are also dual moded (see Sec. IIIC). For a detailed design of the dual-moded delay lines, the reader is referred to [6].

With orthogonal relative phasing of the klystron pairs, power is directed around an alternate path to the same dual-mode output port. Thus, the system can be run in compressed ( $TE_{01}$ ) or uncompressed ( $TE_{11}$ ) mode. Dual-mode directional couplers in circular waveguide for monitoring the power in each operating mode (see Sec. IIE) are installed before and after the superhybrid.

A splitter similar to the input combiner divides the output power from either operating mode between two arms, each of which further divides the power again in four, so that it can be sent into eight high-power loads (see Sec. IIIC).

## 2. System cold testing

Before system assembly, components were individually cold tested with a network analyzer. Results were quite satisfactory (see above sections), with port matches and desired isolations typically somewhat better than  $-30$  dB and component losses on the order of a percent. The system itself was also cold tested at various stages of installation. After the system was fully assembled, we measured its response through the dual-mode directional couplers. A signal was injected into the  $TE_{01}$  arm of the directional coupler just before the superhybrid. The output was measured at the  $TE_{01}$  arm of the directional coupler at the output of the superhybrid, just before the load tree (see Fig. 19), which terminated the system. Initial results are shown in Fig. 20. Both the frequency and time domain

indicated a spurious resonance that occurred at almost exactly 11.424 GHz, our operating frequency. Because the system is highly overmoded and has single-moded waveguide terminals, it can trap several parasitic resonances. These can greatly degrade efficiency and can lead to dangerously high localized energy storage and field levels. To deal with such an eventuality, we prepared several double-sided flanges of various widths, which could be used as spacers to perturb the length of the system at various points. Once we inserted a spacer of appropriate width just after the combiner (see Fig. 19), the resonance mode frequency shifted away from 11.424 GHz, and the system performed adequately. This is documented in Fig. 21.

After completion of installation, the system was pumped down and carefully baked, after which vacuum levels on the order of  $10^{-9}$  torr could be reached. The cold test of the system was then repeated and the same results were obtained, indicating that the baking did not distort any part of the system.

## IV. HIGH-POWER EXPERIMENT

In this section we discuss the high-power tests. We start by describing our processing procedures and observations. Next we describe our low-level rf system that led to an output pulse suitable for multibunch acceleration; i.e., with phase and amplitude controlled to sufficient accuracy. Finally, we describe a reliability test performed on the system for many hours of operation.

### A. Processing

Processing of high-power rf components and high gradient accelerator structures is usually done with a narrow pulse width at the beginning ( $\sim 100$  ns). The power is then increased gradually to the maximum value. When the system runs with an acceptable breakdown rate at a given pulse width and maximum power level, the power is reduced and the pulse width increased. This process is re-

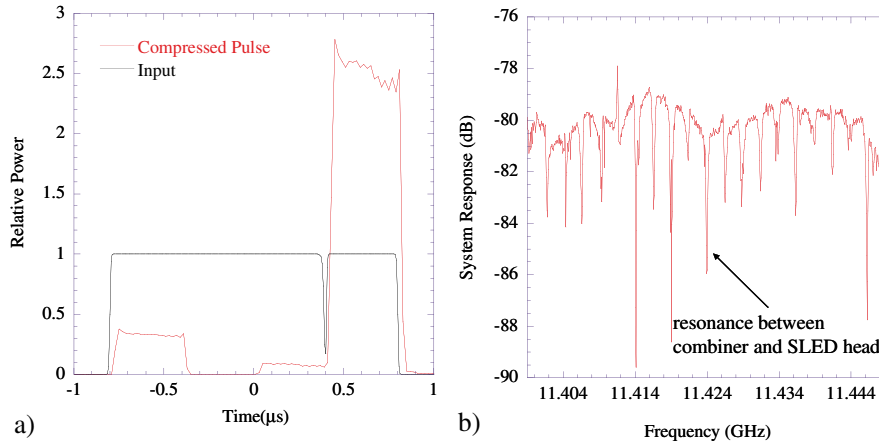


FIG. 20. (Color) Initial overall system response. A spurious resonance occurred at 11.424 GHz.

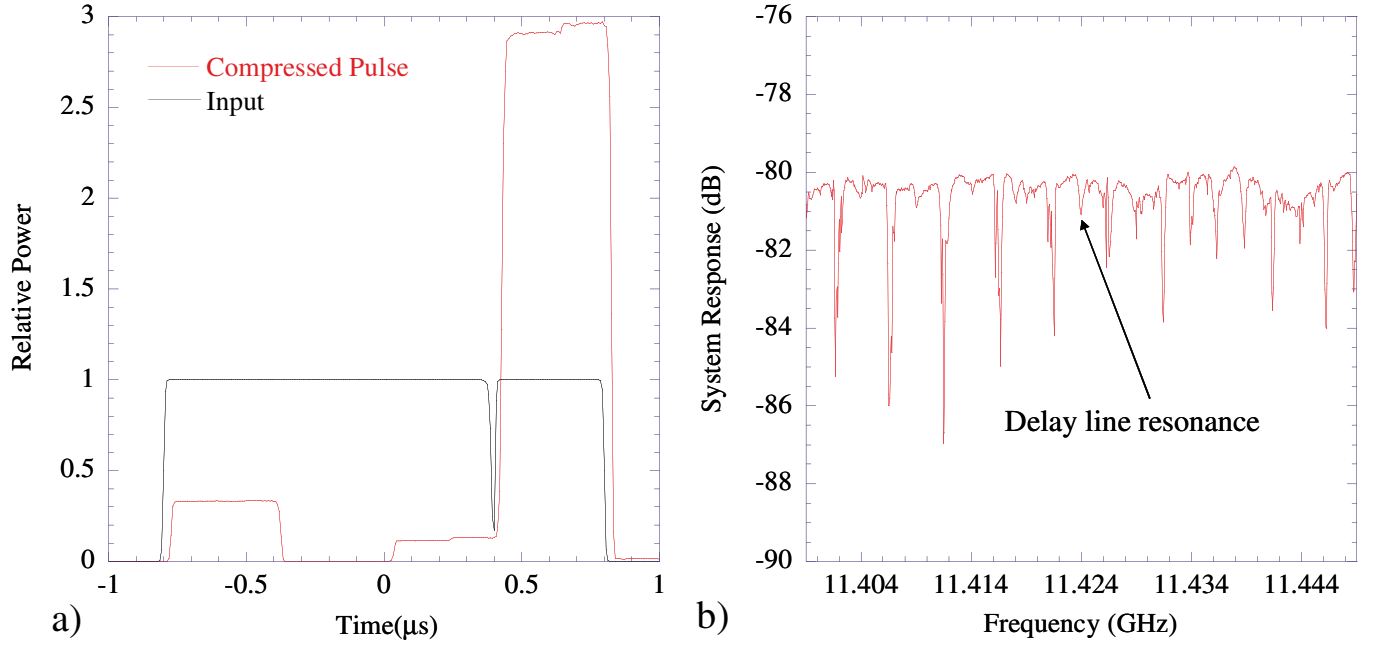


FIG. 21. (Color) Final system response.

peated until the system reaches the design pulse width and power level.

The motivation of this process is to avoid damage to the system due to high energy content per pulse. In our system, the four klystrons provide at their peak 200 MW of power for  $1.6 \mu\text{s}$ , which amounts to 320 J per pulse. This is enough to cause considerable damage if all of this energy is dumped into one breakdown arc. The physics of processing rf copper structures under high electric and magnetic field conditions is not well understood. However, we conjecture that it is due to the high applied electric field, and that if the application time is short, one has a better chance of improving the surface condition without causing

damage to the structure. Hence the pulse widening procedure is adopted.

To provide a short pulse to the system while preserving its functionality as a pulse compression system, the input pulse was divided into four short pulses or time bins. This was done through modulation of the inputs to the klystrons. The first three pulses are in phase, while the last one is shifted by  $180^\circ$ . The response of the system under these input conditions is shown in Fig. 22(a). The processing started with 100 ns time bins, and the system was pushed all the way to produce close to 500 MW in the fourth output time bin, or the compressed pulse. The pulse width was then increased as shown in Fig. 22(b). During this initial

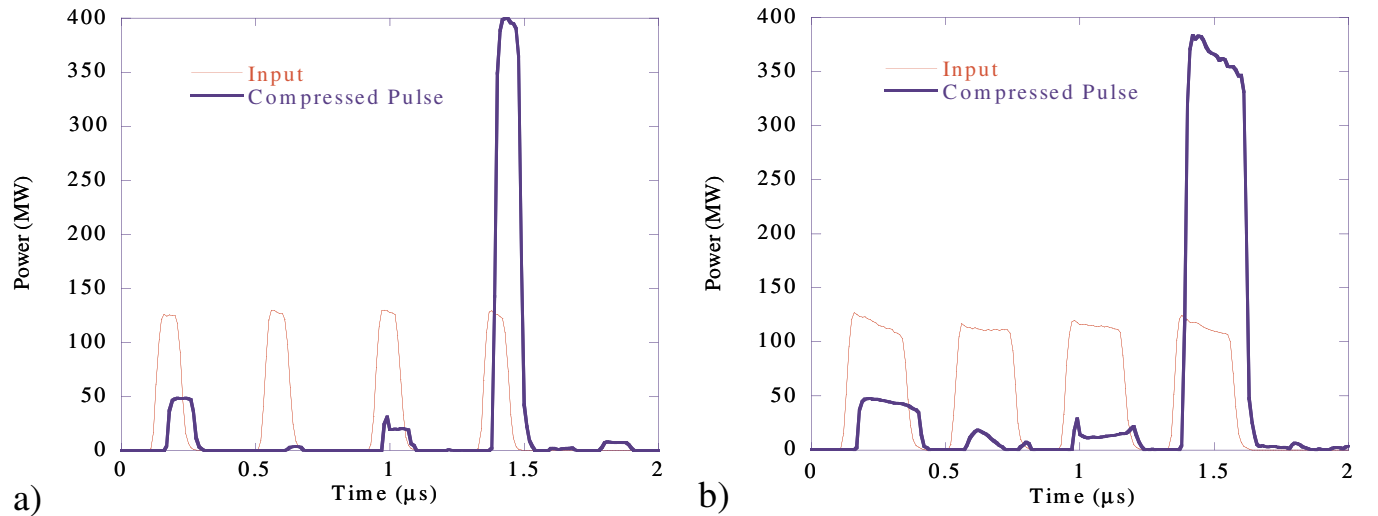


FIG. 22. (Color) Typical processing waveforms.

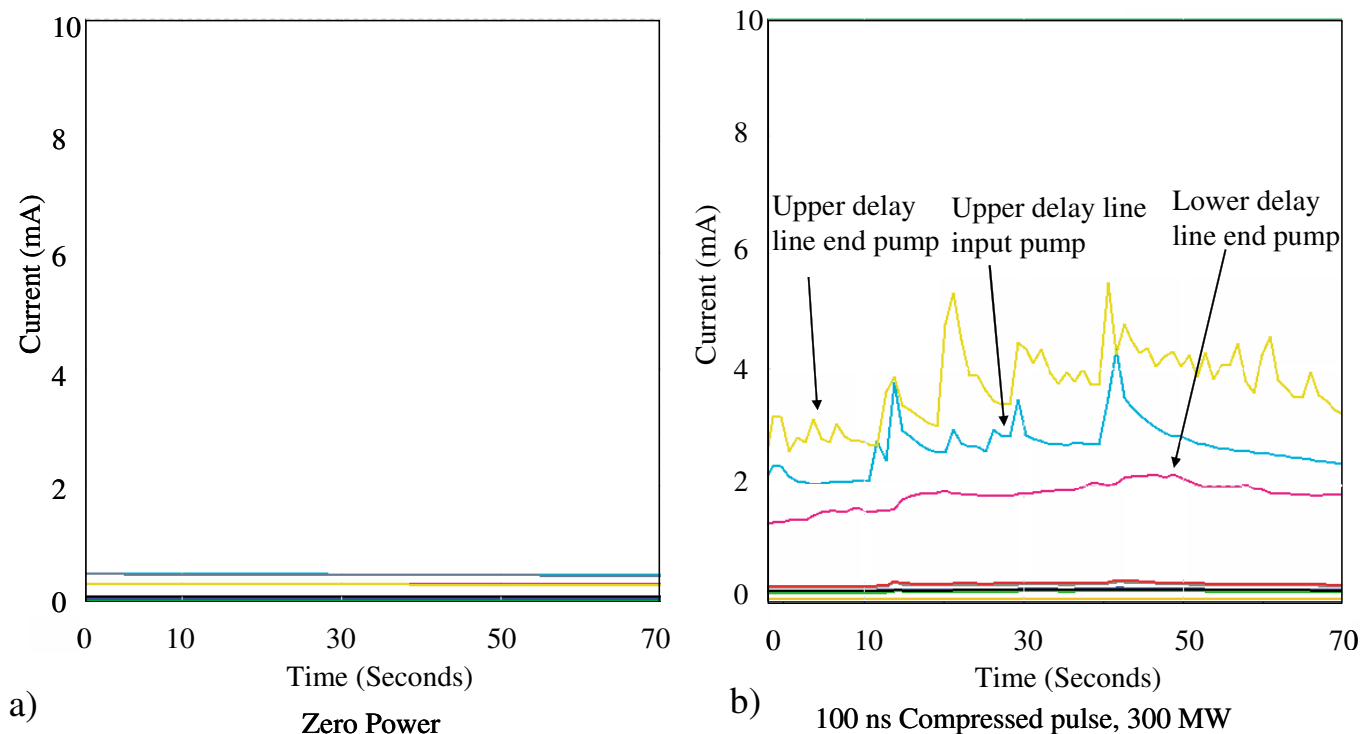


FIG. 23. (Color) Typical ion pump currents during processing (a) without and (b) with power. The figures contain traces for many ion pumps throughout the system. These typically do not show any activity. The delay lines each have an 80 liter/s pump at each end.

process, the delay lines were the cause of most of the out-gassing. Figure 23 shows a typical chart for the currents in the ion pumps at different locations in the system. The current in the ion pumps attached to the beginning and end of the top delay line were noticeably more active than other parts of the system. During cold tests, this was the line that showed more spurious mode conversion between the  $TE_{01}$  and  $TE_{02}$  modes.

The processing procedure continued with a limit on how high the pressures or ion pump currents were allowed to rise. After a period of time, typically 15 min, of running at a given power level, the pressure would start to go down, and we would raise the power. This process was automated. During this process we seldom saw any evidence of breakdown. When a rf breakdown occurred, we observed it through a sudden change in the reflected and transmitted rf power levels monitored throughout the system and detected by a pattern recognition system. Breakdowns during processing were usually preceded by a sharp increase in pressure. A limit on the rate of pressure rise was thus set in the processing algorithm; if the pressure rose too quickly, the power would be dropped to avoid a rf breakdown and slowly raised again. In about 60 h the system reached close to 600 MW of compressed power. This is shown in Fig. 24. During this processing period we were only limited by out-gassing in the delay lines. Occasionally the load tree or the waveguide leading to it would breakdown, but these events were rare.

The power shown in Fig. 24 was measured through the directional couplers using a peak power meter. To verify this reading, measurements of the average calorimetric power in the output loads were performed. A comparison

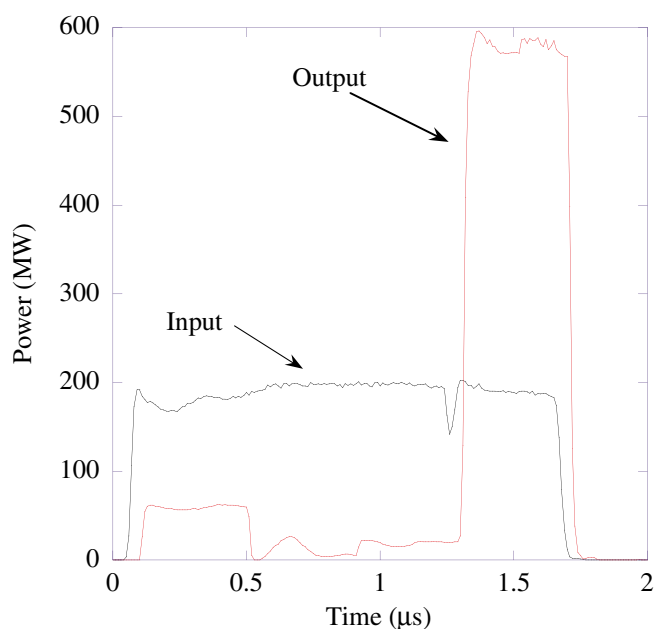


FIG. 24. (Color) System output after processing and without low-level rf corrections (see below).

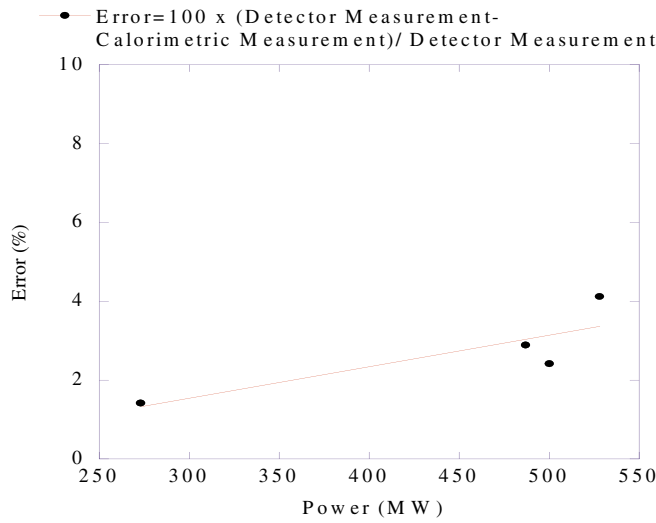


FIG. 25. (Color) Discrepancy between readings of the peak power analyzer connected to the dual-mode directional coupler at the output and calorimetric measurements on the cooling water of the loads. The peak power analyzer signal was integrated to get average power for comparison.

between the calorimetric measurements and the power meter measurements is shown in Fig. 25.

### B. Pulse shape correction and the low-level rf system

Because the output pulse is intended to accelerate a multibunch electron beam, the pulse top needs to be accurately controlled in both phase and amplitude. The pulse shown in Fig. 24 is rough due to the variation of the input signal from the four klystrons in both phase and amplitude.

To correct the input signal, the combined power level injected into the TE<sub>01</sub> mode was measured using the dual-mode directional coupler. The signal from that coupler was down-converted to 476 MHz, and both amplitude and phase were measured using a logarithmic amplifier/detector. The measured signal was digitized and imported into a computer. The computer then calculated the needed corrections for the input of the klystrons, and modified the output of the two 80 MHz arbitrary function generators that controlled an *I/Q* modulator. This modulator controlled the signal sent to the TWTs that drove the klystron. The system is shown in Fig. 26. We assumed a linear model for all components, including the klystrons, although this does not give an accurate approximation of the response. The system thus converged slowly, making a small correction each time around the loop. After about 15 iterations, the system input signal was flat in phase to about 1°, the limit of our crude measurement system. The input pulse amplitude could not be made as flat as desired because of the variation of the klystron input voltage across the pulse. Especially at the beginning of the pulse, the voltage was low enough that the klystron would deeply saturate. Nonetheless, the output pulse, shown in Fig. 27, showed excellent phase and amplitude flatness.

### C. Reliability test

To demonstrate the reliability of the high-power rf system, an extended run of round-the-clock operation at 500 MW was undertaken. This test accumulated 357 h of smooth running at a repetition rate of 30 Hz and 99 h of running at 60 Hz. The higher repetition rate operation had to wait for a chiller to be added to the modulator cooling

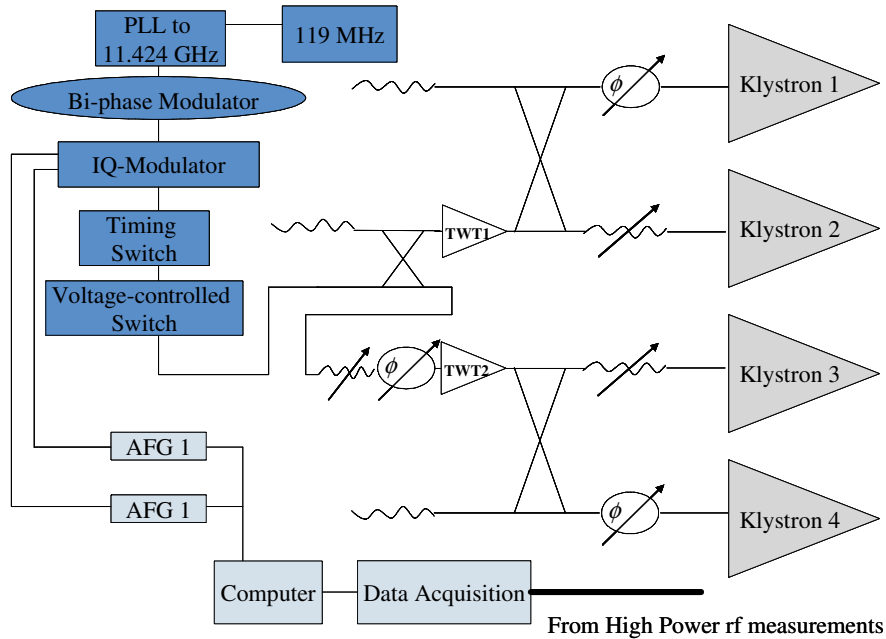
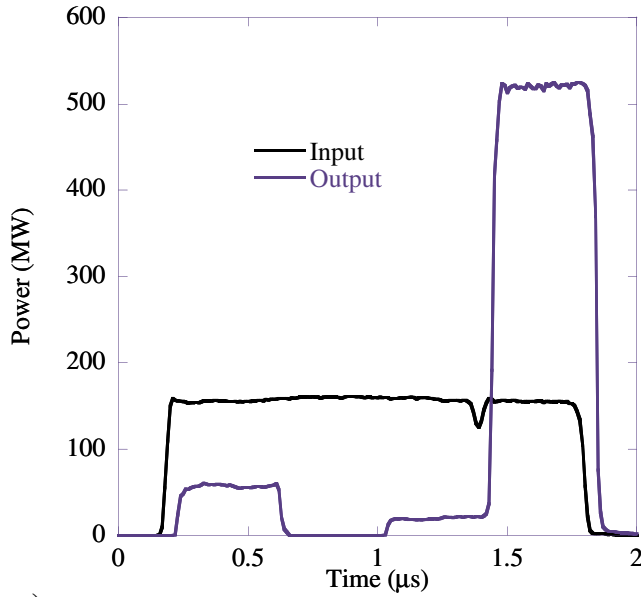
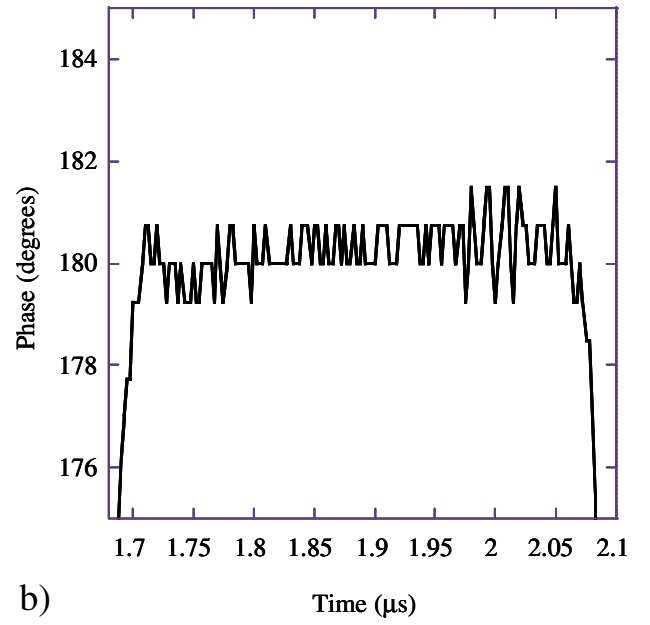


FIG. 26. (Color) Low-level rf system architecture.



a)



b)

FIG. 27. (Color) System response after low-level rf correction is applied. (a) Power levels and (b) close-up view for the phase of the output pulse.

system. The total represents the equivalent of 555 h at 30 Hz. These numbers are integrated on times, excluding automated recovery time which was set to 1 min.

For machine protection, the system was interlocked to trip off on various vacuum and rf power signals. For the last 366 30-Hz-equivalent hours, or  $3.95 \times 10^7$  pulses, a full set of diagnostic signals was saved and archived for each trip, allowing determination of its cause. The most common cause of trips was the phenomenon of pulse tearing, or pulse shortening, in the klystrons, followed in frequency by breakdown in the waveguide between the klystrons and the combiner. By disconnecting the klystrons from the pulse compressor and running them into a load, it was determined that this pulse tearing phenomenon was unrelated to our overmoded waveguide system.

Out of  $39.5 \times 10^6$  pulses, we recorded only 209 trips, the distribution of which is given in Table I. Of these, about 29 appeared to originate in the region from the overmoded

combiner through the SLED-II system. This rate, 0.079 per hour, meets our goal of 0.083 per hour, or 1 in  $1.3 \times 10^6$  pulses. Correcting for trips attributable to operator interference or to delay line mistuning (before implementation of automated tuning software) brings this number comfortably below that limit.

#### D. $TE_{11}$ processing

For the sake of completeness, before performing the reliability test described above, we tested the system with the alternative mode,  $TE_{11}$ , injected, rather than  $TE_{01}$ . As

TABLE I. Results of reliability test trip analysis.

| Source of trip      | Source of trip |
|---------------------|----------------|
| SLED-II or combiner | 29             |
| Klystron 1          | 1              |
| Klystrons 1 & 2     | 15             |
| Klystron 2          | 18             |
| Klystron 3          | 72             |
| Klystrons 3 & 4     | 28             |
| Klystron 4          | 38             |
| Loads               | 1              |
| Undeterminable      | 7              |

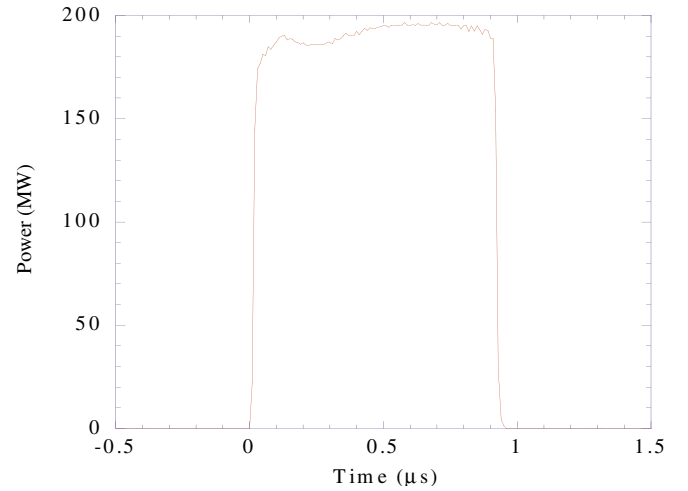


FIG. 28. (Color) System response after low-level rf correction is applied. (a) Power levels and (b) close-up view for the phase of the output pulse.

described earlier, this mode of operation is obtained by a simple change in the relative phase between the klystron pairs. Under  $TE_{11}$  operation, no power is delivered to the delay lines; the power goes directly through the super-hybrid to the load assembly. The field pattern throughout the system is different than in the case of  $TE_{01}$  operation, and the flanges see the axial currents of the  $TE_{11}$  mode. Therefore, some additional processing was required for the system to run at high power in this mode. After approximately 10 h of running, we reached power levels approaching 200 MW with a pulse width close to 1  $\mu$ s. This is shown in the measured peak power meter waveform of Fig. 28. We did not commit time for further processing of the system under these conditions due to the need to use the system to power accelerator structures.

## V. CONCLUSION

We have presented theoretical designs for a set of waveguide components suitable for ultra-high-power rf manipulations. Most of these components are dual moded. Experimental verification of the functionality of many has also been presented. From these various components, we have built and operated at 11.424 GHz a dual-moded pulse compression system. The system produced close to 600 MW of output power with a pulse width of 400 ns and a repetition rate of 60 Hz. We believe this sets a new record in pulsed rf high-power systems with precision control of amplitude and phase.

## ACKNOWLEDGMENTS

We acknowledge years of discussions and input from Professor Norman Kroll who had a great influence on the design directions and techniques used in our study. We thank Ronald Ruth, Perry Wilson, Roger Miller, Gordon Bowden, and David Farkas for many useful discussions. We thank David Burke and David Schultz for their leadership roles during the course of this project. The authors wish to thank Marc Ross, Tonee Smith, and Douglas McCormick for all their help and support during the commissioning phase of this experiment. We also wish to thank Andrei Lunin, Steffen Doebert, Chris Adolphsen, and Toshiyuki Okugi for taking shifts during the processing phase. We also thank the modulator team, Richard Cassel, Jeffrey De Lamare, Marc Laruss, and Minh Nguyen, for their 24-hour support. We thank Kathleen Ratcliffe and the vacuum team for their support. Finally, we thank Elias Andrikopoulos for his continuous support on the low-level rf system and measurements.

---

[1] NLC ZDR Design Group and NLC Physics Working Group, S. Kuhlman *et al.*, in *Proceedings of Snowmass '96*, Snowmass, CO (Reports No. SLAC-R-0485,

No. BNL-52502, No. FERMILAB-PUB-96-112, No. LBL-PUB-5425, No. LBNL-PUB-5425, No. UCRL-ID-124160, 1996), p. 197; see also *Proceedings of the DPF/DPB Summer Study on New Directions for High-Energy Physics (Snowmass '96)*, Snowmass, CO, 1996 (Report No. SLAC-R-732, 1996).

[2] Sami G. Tantawi, in *Proceedings of the 20th International Linac Conference (LINAC 2000)*, Monterey, CA, 2000 (SLAC Report No. SLAC-PUB-8582, 2000).

[3] Sami G. Tantawi *et al.*, IEEE Trans. Microwave Theory Tech. **47**, 2539 (1999).

[4] V.A. Dolgashev, in *Proceedings of the 2003 Particle Accelerator Conference (PAC'03)*, Portland, OR (IEEE, Piscataway, NJ, 2003), pp. 1267–1269.

[5] S.G. Tantawi *et al.*, Phys. Rev. ST Accel. Beams **5**, 032001 (2002).

[6] S.G. Tantawi, Phys. Rev. ST Accel. Beams **7**, 032001 (2004).

[7] C.D. Nantista, W.R. Fowkes, N.M. Kroll, and S.G. Tantawi, in *Proceedings of the 1999 Particle Accelerator Conference (PAC'99)*, New York (IEEE, Piscataway, NJ, 1999).

[8] Christopher D. Nantista and Sami G. Tantawi, IEEE Microwave Guided Wave Lett. **10**, 520 (2000).

[9] Sergey Kazakov, in *Proceedings of the 6th International Study Group*, Tsukuba, Japan, 2000, <http://lcdev.kek.jp/ISG/ISG6.SKaz2.pdf>

[10] S.G. Tantawi, N.M. Kroll and K. Fant, in *Proceedings of the 1999 Particle Accelerator Conference (PAC'99)*, New York (Ref. [7]).

[11] I. Spassovsky, E. S. Gouveia, S.G. Tantawi, B.P. Hogan, W. Lawson, and V.L. Granatstein, IEEE Trans. Plasma Sci. **30**, 787 (2002).

[12] S.G. Tantawi, in *Proceedings of the 1993 Particle Accelerator Conference (PAC'93)*, Washington, DC (IEEE, Piscataway, NJ, 1993).

[13] W.X. Wang *et al.*, Int. J. Electron. **65**, 705 (1988).

[14] S.E. Miller, Bell Syst. Tech. J. **33**, 661 (1954).

[15] P.B. Wilson, Z.D. Farkas, and R.D. Ruth, in *Proceedings of the 1990 Linear Accelerator Conference*, Albuquerque, NM (Reports No. LA-12004-C, No. UC-910, and No. UC-414, 1991).

[16] A. Fiebig and C. Schiebllich, in *Proceedings of the 1988 European Particle Accelerator Conference (EPAC'88)*, Rome (World Scientific, Singapore, 1989), pp. 1075–1077.

[17] Sami G. Tantawi and A.E. Vlieks, in *Proceedings of the 1995 Particle Accelerator Conference (PAC'95)*, Dallas, TX (IEEE, Piscataway, NJ, 1996).

[18] S.G. Tantawi, R.D. Ruth, and A.E. Vlieks, Nucl. Instrum. Methods Phys. Res., Sect. A **370**, 297 (1996).

[19] George Caryotakis, in *Proceedings of the 1997 Particle Accelerator Conference (PAC'97)*, Vancouver, Canada (IEEE, Piscataway, NJ, 1998).

[20] R.L. Cassel, G.C. Pappas, M.N. Nguyen, and J.E. DeLamare, in *Proceedings of the 1999 Particle Accelerator Conference (PAC'99)*, New York (Ref. [7]).

[21] HP High Frequency Structure Simulator (HFSS), Version 5.4, Hewlett-Packard Co., copyright 1996–1999.