

Lawrence Berkeley National Laboratory

Lawrence Berkeley National Laboratory

Title

FERMI@Elettra FEL Design Technical Optimization Final Report

Permalink

<https://escholarship.org/uc/item/92x9t83q>

Authors

Fawley, William
Penn, Gregory
Allaria, Enrico
et al.

Publication Date

2006-07-31



ST/F-TN-06/16
LBNL-61333

FERMI@Elettra FEL Design Technical Optimization Final Report

W. Fawley
G. Penn

Lawrence Berkeley National Laboratory, Berkeley, CA 94720 USA

E. Allaria
G. De Ninno

Sincrotrone Trieste, Italy

W. Graves

Massachusetts Institute of Technology, Cambridge, MA 02139 USA

August 2006

SINCROTRONE TRIESTE — Società Consortile per Azioni

S.S. 14 km 163,5 in Area Science Park – 34012 Basovizza – Trieste, Italy – Tel. 040.37581 – Telefax 040.9380902 – P. IVA e Cod. Fisc. IT00697920320
Capitale sociale € 45.022.478,60 – Iscritta al n. 9534 del Registro delle Imprese di Trieste

Società di interesse nazionale ai sensi della Legge 370 del 19/10/1999

Table of Contents

1	<i>Introduction.....</i>	<i>1</i>
1.1	The Harmonic Cascade Approach to Short Wavelength Generation.....	1
1.2	Basic FEL Output Requirements	3
1.3	Comparison between the seeded harmonic cascade and pure SASE approach.....	5
1.4	Important Phenomena which affect the FEL Performance.....	6
2	<i>Undulator & Transport Lattice Design.....</i>	<i>7</i>
2.1	Choice of Undulator Type and Wavelength	7
2.2	Undulator Segmentation and Focusing issues	8
2.3	Undulator Error Tolerance Calculations.....	9
2.4	Post-modulator dispersive section design and issues.....	10
2.5	Delay section needs and issues for FEL-2 FB approach	10
2.6	Vacuum Chamber Longitudinal Wakefield Considerations.....	11
3	<i>FEL-1 Design & Performance Calculations</i>	<i>12</i>
3.1	Basic design	12
3.2	Time-independent Simulation Results	13
3.2.1	<i>Single Parameter Sensitivity Studies at 100 nm.....</i>	<i>13</i>
3.2.2	<i>Single Parameter Sensitivity Studies at 40-nm</i>	<i>14</i>
3.2.3	<i>Time-independent Calculations for Predicted Jitter</i>	<i>15</i>
3.2.4	<i>Undulator Taper Strategies for Reducing Sensitivity to Beam Energy.....</i>	<i>16</i>
3.3	Input Transverse Tilt and Offset Sensitivity Studies for FEL-1	17
3.3.1	<i>Introduction.....</i>	<i>17</i>
3.3.2	<i>Offset and Tilt Sensitivity Studies Using GENESIS.....</i>	<i>18</i>
3.3.3	<i>Tolerance Studies with Offset/Tilt: Simultaneous Variation of Multiple Input Parameters.....</i>	<i>20</i>
3.4	Time-dependent Results for FEL-1	21
3.4.1	<i>Summary of Time-Dependent Simulation Method.....</i>	<i>21</i>
3.4.2	<i>Time-dependent results for the “M2” and “M6” Medium Pulse Distributions</i>	<i>22</i>
3.4.3	<i>Wakefield Calculations for the M2 & M6 Medium Pulse Distributions</i>	<i>23</i>
3.5	Time-Dependent Jitter calculations.....	24
3.6	Diagnostic needs for FEL-1	26
4	<i>FEL-2 Design & Performance Calculations</i>	<i>27</i>
4.1	Basic Design for the Fresh-Bunch Approach	27
4.2	Fresh-Bunch Time-Independent Simulation Results	27
4.3	Fresh-bunch Offset and Tilt Sensitivity Studies Using GENESIS	28
4.3.1	<i>Introduction.....</i>	<i>28</i>
4.3.2	<i>10- and 20-nm Fresh-Bunch Results with Offsets and Tilts.....</i>	<i>29</i>
4.4	Fresh-Bunch Time-dependent Simulation Results.....	29
4.5	Basic Design for the Whole-Bunch Approach.....	30
4.6	Whole-Bunch Time-independent Simulation Results.....	30
4.7	Whole-Bunch Offset and Tilt Sensitivity Studies Using GENESIS.....	31
4.8	Whole-Bunch Time-Independent Optimization and Jitter	31
4.8.1	<i>Optimization for Maximum Output Power</i>	<i>31</i>
4.8.2	<i>Optimization for Reduced Output Jitter.....</i>	<i>32</i>
4.8.3	<i>Optimization for Reduced Bandwidth.....</i>	<i>33</i>
4.8.4	<i>Whole-Bunch Time-Independent Jitter Sensitivity Studies</i>	<i>34</i>

4.9	Whole-Bunch Time-Dependent Simulation Results	34
4.9.1	<i>Introduction</i>	34
4.9.2	<i>Long Pulse “L2” and “L4” Simulation Results</i>	34
4.9.3	<i>Medium Pulse “M2”, “M4”, and “M6” Simulation Results</i>	35
4.10	Wakefield Calculations for the L2 & L4 Long Pulse Distributions.....	36
4.11	Comparison of Known Advantages and Disadvantages of FB vs. WB	36
4.12	Diagnostic Needs Unique to FEL-2	37

1 Introduction

1.1 *The Harmonic Cascade Approach to Short Wavelength Generation*

The FERMI@ELETTRA project is based on the principle of harmonic upshifting of an initial “seed” signal in a single pass, FEL amplifier employing multiple undulators. There are a number of FEL physics principles which underlie this approach to obtaining short wavelength output: (1) the energy modulation of the electron beam via the resonant interaction with an external laser seed (2) the use of a chromatic dispersive section to then develop a strong density modulation with large harmonic overtones (3) the production of coherent radiation by the microbunched beam in a downstream radiator. Within the context of the FERMI project, we discuss each of these elements in turn.

An external laser provides an initial, wavelength-tunable seed signal. This signal, in conjunction with a magnetic undulator (the “modulator”), produces a relatively strong energy modulation $\Delta\gamma$ of the beam electrons via resonant interaction. The modulation has a sinusoidal variation in time identical to that of the seed’s angular frequency, ω_0 ($=2\pi c/\lambda_0$ where λ_0 is the seed wavelength). When the modulator’s length is comparable to or shorter than the exponential gain length for FEL radiation power *and* the number of undulator periods obeys the relation $2N_u (\Delta\gamma/\gamma_0) < 1$, there is very little accompanying density modulation (*i.e.*, microbunching) produced in the modulator. Both these requirements are well satisfied by FERMI’s FEL-1 and FEL-2 modulators.

Following its exit from the modulator, the electron beam then passes through a chromatic dispersion section in which a density modulation develops from path length differences associated with the energy modulation. Typically, $R_{56} (\Delta\gamma/\gamma_0) \sim \lambda_0/4$ which, so long as $\Delta\gamma \gg \sigma_\gamma$, the initial “incoherent” energy spread (slice, not projected), leads to a strong periodic density modulation at wavelength λ_0 containing a large higher harmonic component (up to harmonic number $m \sim \Delta\gamma/\sigma_\gamma$). With the exception of the so-called “whole-bunch” approach to FEL-2, the microbunching fraction at the end of the chicane is quite large, typically in the range of 0.2 – 0.5. Note that at a given longitudinal position in the electron beam, the relative spread (*i.e.*, at different transverse positions) in the induced $\Delta\gamma$ must also be quite small or else the density modulation will be degraded, especially at higher harmonics. Consequently, the transverse extent of the seed laser (presuming a well behaved, Gaussian-like profile) should be significantly greater than that of the electron beam.

At this point the electron beam enters a new undulator (the “radiator”) whose wavelength and magnetic strength are tuned such FEL resonance occurs at an integral harmonic “ m ” of the original seed laser wavelength:

$$\lambda_R = \lambda_0/m = \lambda_w \times (1 + a_w^2) / 2\gamma^2$$

where a_w is the normalized RMS undulator magnetic strength and λ_w is the undulator period. For FERMI, we have typically chosen m to be between 3 and 6 for the first radiator. If (as in FEL-1) this radiator is the final undulator, it generally is made sufficiently long for the FEL radiation to grow to saturation (or even stronger via tapering if greater output power is sought). At long wavelengths (*i.e.*, smaller m) where the microbunching at the wanted harmonic often exceeds 0.2 at entrance to the radiator, the required length

might then be two gain lengths or less. At shorter wavelengths (*i.e.*, larger m) where the initial microbunching might be smaller than 0.1, four or more gain lengths might be necessary for saturation.

For a multistage harmonic cascade such as FEL-2, the first stage radiator is generally much shorter than that necessary for power saturation. In the so-called “fresh-bunch” approach in which radiation from the first radiator is used to energy-modulate part of the electron beam in a subsequent modulator, the first radiator is only made long enough that the radiation is sufficient to produce adequate downstream energy modulation. For reasonable currents, the necessary length can be less than two exponential gain lengths and the radiation is effectively coherent spontaneous emission whose power scales as $I^2 z^2$ (ignoring diffraction and debunching effects). Following the first radiator is an e-beam temporal delay section (essentially a chicane) in order to make the output radiation temporally coincident with a “fresh” section of the electron beam closer to the beam head. This fresh section has not had its incoherent energy spread increased via FEL interaction in the first stage modulator and radiator. Thus, it can be far more easily energy- and density-modulated in the second stage undulators than the “old” e-beam section that interacted in the first modulator and radiator.

The second stage for the fresh-bunch approach consists of a modulator, a final radiator, and, in general, an intervening dispersive section. The modulator uses the radiation from the first stage radiator and must therefore have its magnetic strength be tuned to be resonant at the same wavelength. At present, the first radiator and second modulator in the FERMI design also have the same undulator wavelength although this is not a physics requirement. Since the radiation is freely diffracting once it departs the first radiator, some care must be taken both not to make the temporal delay section too long and that the necessary second modulator length does not exceed ~ 1 Rayleigh length. Otherwise, the coupling between the radiation and e-beam might become too poor for sufficient energy modulation to be developed. Presuming there is good coupling, the second stage modulator, radiator, and intervening dispersive section are quite similar in concept to the first stage. In general, the harmonic upshift number between the second stage modulator and radiator is 4 or less (compared with up to 6 in the first stage). Moreover, the amount of microbunching at the new harmonic in the second radiator is also generally less than half that produced by the first because both the undulator parameter a_w and the radiation intensity are smaller. This generally leads to a smaller energy modulation at the end of the second modulator. The second stage radiator is usually much longer than that of the first stage both because normally the initial bunching is smaller and because the FEL is run to saturation (which requires more distance because the corresponding exponential gain lengths are longer due to the smaller a_w). The physics of this radiator is in the “high gain”, exponential growth regime and is thus much more similar to classic HGHG scheme of Yu [*Phys. Rev.* **A44**, 5178 (1991)] as compared with the first stage radiator (or the single stage FEL-1 configuration) which normally operates in the “low gain” regime with little or no increase in microbunching fraction over the length of the undulator.

In the “whole-bunch” approach to an FEL harmonic cascade, the *entire* electron beam pulse is energy-modulated by the external laser seed and, following the first radiator, there is neither a temporal delay section nor a second modulator. Instead, the electron beam immediately enters a weak dispersive section followed by a second radiator whose FEL resonant wavelength is tuned to an integral harmonic of the first radiator. Due to the relatively small harmonic microbunching at this new wavelength, this second radiator must operate deep in the exponential gain regime. Thus, to keep the exponential gain length and power saturation lengths acceptably small, the energy modulation produced by the first (and only) modulator

must be relatively small compared to $\rho_2 \gamma$ where ρ_2 is the FEL parameter for the second radiator (generally $\sim 1 \times 10^{-3}$). This small energy modulation means that at entrance to the first radiator the e-beam will have a smaller microbunching level at entrance relative to that of the fresh-bunch scheme. Consequently, the whole-bunch approach can essentially fail (in terms of the needed second radiator undulator length for saturation) if the initial energy spread becomes too large. Moreover, because the microbunching level is small at the start of both the first and second radiator, the relative strength of the shot noise microbunching is much higher and the SASE (Self-Amplified Spontaneous Emission) output strength can be 2 or more orders of magnitude greater in the whole-bunch approach than the fresh-bunch approach. Consequently, the output signal-to-noise ratio is likely to be much smaller in the whole-bunch approach. On the other hand, the whole-bunch approach is much less sensitive to shot-to-shot fluctuations of the relative timing between the e-beam and external seed laser.

1.2 Basic FEL Output Requirements

Table 1-1 - Basic FERMI Output FEL Parameters

Parameter	FEL-1	FEL-2
Photon energy range	12 – 31 eV (100 – 40 nm)	31 – 124 eV (40 – 10 nm)
Pulse length	100 fs	≥ 400 fs
Bandwidth	20 meV	~ 10 meV
Repetition rate	10 – 50 Hz	10 – 50 Hz
Harmonic peak power	$\sim 2\%$ of fundamental	$\sim 0.5\%$ fundamental
Photons per pulse	2×10^{13} (@40 nm)	$\sim 10^{13}$ (@ 10 nm)
Peak power	$\sim 1 - 5$ GW	$\sim 0.3 - 1$ GW
Output transverse position stability	50 microns	50 microns
Pointing stability	5 microradians	5 microradians
Power fluctuation	$\leq 25\%$	-----
FEL ρ parameter	2.9×10^{-3} (40 nm)	1.2×10^{-3} (10 nm)

The basic FEL output requirements for FEL-1 and FEL-2 are summarized in Table 1-1; we give some additional details as follows. FEL-1 should have a continuously tunable 40-100 nm operational range in output wavelength; the upper limit is somewhat soft (although the exact value does affect the choice for the undulator wavelength and minimum gap for the radiator). FEL-2 will have a operating range of 10-40 nm. Here too there is some softness in the parameters in that the upper limit of 40-nm can be reduced (and the lower wavelength limit of FEL-1 similarly reduced) if this proves convenient for some reason. The lower limit of 10-nm is also somewhat soft and could be reduced (*e.g.*, to 7-nm) if needed by a particular experiment (albeit with likely a much reduced power output because of the resulting smaller undulator a_w and less gain due to emittance degradation effects). At all wavelengths, both FEL-1 and FEL-2 are to have continuously tunable output polarizations ranging from linear-horizontal to circular to linear-vertical. Consequently, the FEL-1 radiator and FEL-2 final radiator are to have APPLE configurations. Both FEL's will operate at 10-50 Hz; this specification is constrained by the accelerator and not the FEL subsystems; operation at 1-kHz would not significantly increase the demands and cost of the input seed laser system.

At present it is believed that the major application for FEL-1 will involve time-domain experiments such as pump-probe interactions and possibly nonlinear phenomena. Consequently, the requirements for FEL-1 are more related to total photon number per pulse (*i.e.*, $0.4 - 2 \times 10^{14}$) and pulse duration (20-100 fs) than they are to spectral bandwidth. A critical parameter affecting the needed electron beam duration is the timing jitter in the beam relative to that of the seed laser. In order that there be reasonably overlap between the seed and the electrons, the duration of the electron beam must satisfy $\tau_b \geq \tau_L + 2 \Delta t_{jitter}$. If the expected RMS timing jitter from the accelerator is of order 150 fs, an e-beam pulse duration ~ 500 fs or greater will be needed for 100-fs seed pulses. This timing jitter requirement may be one of the more difficult to satisfy in terms of the injector and accelerator subsystems.

Another important parameter associated with FEL-1 time domain experiments is shot-to-shot repeatability. Ideally, for nonlinear phenomena experiments, shot-to-shot RMS jitter in normalized photon number should be 5% or less. As will be seen in the discussion of predicted results (see Sec. 3.5), we do not at present believe that such a low value is obtainable with the expected accelerator and injector parameters. Many FEL-1 experiments will be able to deal (via recording individual shot photon number for later post-processing) with values as high as 25% or greater. Many other FEL-1 output requirement parameters are related to energy jitter parameters --- pointing, virtual waist location and angular divergence jitter, shot-to-shot transverse profile changes. Although none of these is likely *on an individual basis* to prevent FERMI from successfully reaching the 5% goal of (spatially) local intensity fluctuations at the experimental sample, taken together they will likely produce jitter that will exceed this goal even in the absence of photon number jitter. On the other hand, some nonlinear experiments (*e.g.*, those using gaseous samples) may be insensitive to pointing or profile changes.

In contrast to FEL-1 where timing and photon number jitter are critical parameters, many if not most FEL-2 users are interested in frequency domain experiments where longitudinal coherence and narrow bandwidth are most important. Perhaps the most important output goal for FEL-2 is $\geq 10^{12}$ photons/pulse/meV. Consequently, FEL-2 requirements favor long output pulses ($\tau \geq 1$ ps) whose spectral properties ($\Delta E_\omega < 10$ meV) are as close as is reasonably possible to the transform limit. Although the total photon jitter is not critical for most experiments in the frequency domain, shot-to-shot central wavelength jitter is of concern. Consequently, in order to not increase the effective time-averaged, output bandwidth as seen by the user, the wavelength jitter needs to be less than the individual shot bandwidth. At 10-nm output wavelength (124 eV photon energy), ± 5 meV corresponds to $\pm 4 \times 10^{-5}$ normalized wavelength jitter.

Apart from the RMS bandwidth consideration is that of "spectral resolution". For some experiments (such as RIXS where one is examining a small inelastic scattering cross-section in the presence of a much, much larger elastic scattering cross-section), a needed spectral resolution of 10^5 requires that the integrated noise photon level (at the detector) be less than 1 part in 10^5 of the wanted signal. For spectrally unfiltered situations, this could be a much more severe requirement than that of RMS bandwidth. For example, if the integrated noise power is 1 part in 10^4 but has a bandwidth 100 times greater than the main signal, the total (signal + noise) RMS bandwidth increases by only $\sim 40\%$ from that of the signal, but the unfiltered spectral resolution would still miss the 10^5 criterion by a factor of ten.

1.3 Comparison between the seeded harmonic cascade and pure SASE approach

There are a number of advantages and disadvantages when comparing the seeded, harmonic cascade approach to short output wavelength chosen in the FERMI project compared to that of the SASE approach being followed in the LCLS project at SLAC and the FLASH facility at DESY. In the seeded approach, both the duration and (to a large extent) the longitudinal coherence of the output FEL radiation is determined by the characteristics of the input seed from the external laser. In a plain “vanilla” SASE configuration (*i.e.*, one long undulator, no special e-beam manipulations nor undulator interruption by a monochromator), the output radiation duration is comparable to that of the electron beam and the longitudinal coherence is quite limited --- typically of order a few thousand output wavelengths. For the seeded approach, output coherence lengths exceeding 10^4 wavelengths at 10 nm are believed possible. Moreover, shot-to-shot energy jitter in a SASE device directly maps into a jitter in the central wavelength, whereas in the seeded approach the output wavelength is generally not sensitive to jitter in the mean electron beam energy (it will, however, be sensitive to amplitude jitter in the e-beam energy chirp). Another advantage of the seeded approach is that the natural timing synchronism with the input radiation seed permits relatively straightforward pump probe experiments (especially for one stage devices such as FEL-1). On the other hand, as pointed out above, the timing jitter between the radiation seed and e-beam must be small enough that there be reasonable overlap between the seed and the e-beam. Another advantage of the seeded approach is that the required undulator lengths are generally much shorter than those needed for SASE since the effective input power is just one or two orders of magnitude below that corresponding to saturation whereas for SASE the effective input noise power is ~ 6 -8 orders of magnitude below saturation.

On the other hand, the SASE approach is generally much simpler than the harmonic cascade approach, particularly if for the latter multiple stages are needed as is true for FEL-2. In SASE, wavelength tuning (over a limited range) is generally done by simply changing the e-beam energy; no change is needed in the undulators (although it may be necessary to change the quadrupole focusing, if present) nor are phase shifters necessary between the undulator sections. Moreover, the ability to operate at a more-or-less fixed a_w implies that one can optimize this value and always operate the undulator at a relatively small gap. For the seeded approach, one must have a tunable input seed whose power does not change too drastically over the needed tuning range. For FERMI@ELETTRA, the accelerator group has concluded it is too difficult to continuously change the e-beam energy. Instead, the resonant wavelength will be varied by changing the undulator magnetic strength a_w (*i.e.*, via changing the gap). This tuning method also leads to a requirement of longitudinal phase shifters between undulator sections. The two stages in FEL-2 will, in general, also need to change a_w by different fractional amounts (since the output wavelength scales as the quantity $[1 + a_w^2]$). While in principle such simultaneous changes of the input seed wavelength, the undulator gap, the quad focusing strength, and the phase shifters can be automated, FERMI would be the first facility of its kind to do so. Also note that a fixed electron beam energy and variable a_w approach implies that the peak a_w will be used at the longer wavelengths (where the gain is relatively high) and the lowest a_w 's at the shortest wavelengths (where the gain is low), the exact opposite of what one would prefer for minimizing undulator length and maximizing power at shorter wavelengths.

1.4 Important Phenomena which affect the FEL Performance

Harmonic cascade FEL's can be sensitive to many effects and input parameters, some of which are common to essentially all FEL schemes and some of which are more unique to the cascade approach to short wavelengths. First, a poor quality electron beam, be it due to a large emittance or a large incoherent energy spread, suffers greatly from debunching in the radiator and will normally produce a poor level of output radiation. In the harmonic cascade approach, the sensitivity to energy spread is significantly heightened as this approach relies upon nonlinear growth of harmonic microbunching. In a multistage cascade, there can be a sharp cliff in σ_E beyond which there will be very little output radiation from the final stage. There can also be sensitivity to input current although in a multistage cascade one can tune the interaction of the two stages such that the design current can produce a local maximum in output power. For currents much lower than the design current, there can also be a sharp cliff in a multistage device since the output power from the first stage scales quadratically with current. Similarly, if one is relying upon strong exponential gain in the final radiator, the output power will be very sensitive to output current.

Perhaps the most sensitive parameter we have found in the FERMI TOS studies presented in the following sections is that of initial electron beam energy. For FEL-1, a fractional change of 0.1% can lead to a large change in output power and this has serious consequences for shot-to-shot jitter requirements. While we have found some undulator tuning strategies that partially mitigate the energy sensitivity, nonetheless it is still of utmost concern, especially for those experiments in the time-domain in which nonlinear (in intensity) phenomena are studied. For FEL-2, shot-to-shot repeatability in photon number currently appears to be less of a concern than is the final spectral bandwidth. Here, however, the reliance upon chromatic dispersion sections to produce strong microbunching before each radiator also produces a strong output wavelength sensitivity to initial energy chirps on the electron beam. In particular, a temporally broad quadratic chirp in energy leads to a linear chirp in output wavelength and can make it very difficult to reach the transform limit for the output spectral bandwidth. However, in this case, the effect is essentially deterministic and the radiation longitudinal phase space is not truly diluted on a microscopic level (i.e., the radiation emittance does not increase). Temporally narrow fluctuations in electron beam energy (such as might develop from the microbunching instability) also will broaden the spectral bandpass and, if severe enough, can truly dilute the microscopic phase space of the radiation.

Both undulator errors and initial tilt or offset errors on the electron beam (or input laser seed) can also degrade performance. However, at present, our studies suggest that the necessary tolerance specifications appear reasonable compared to what has already been achieved for operational insertion devices at ELETTRA and other FEL facilities.

2 Undulator & Transport Lattice Design

2.1 Choice of Undulator Type and Wavelength

With the requirement that the output radiation polarization be continuously tunable from linear to circular polarizations, the FERMI project quickly chose the APPLE configuration for the final radiator undulators (FEL-1 radiator and FEL-2 second stage radiator). For the initial modulator, a simple, linearly-polarized configuration is optimal both due to its simplicity and because the input radiation seed can be linearly polarized. The essentially zero-gain and short length of the first modulator implies there is little cost or space savings by going with a circular polarization for which the e-beam/radiation coupling is somewhat better. For the case of the first stage radiator and second stage modulator (for the fresh-bunch approach) in FEL-2, we have also decided to use simple, linearly-polarized undulators, mainly because of a perceived simplicity and likely cost savings. However, if for whatever reasons (*e.g.*, suppression of higher harmonic emission) circular polarization is required, such a design change could be made without leading to a requirement of greater for undulator length.

As mentioned previously, wavelength tuning in the undulators will be done by changing the gap (and thus a_w) and not by changing the electron beam energy. Hence, the maximum wavelength reachable for a given fixed beam energy is set by the magnetic field at the undulator pole tips, the minimum gap obtainable, and the undulator wavelength. These considerations actually strongly constrain the available parameter space, especially for FEL-1.

For the first modulator which must satisfy FEL resonance over a nominal wavelength range of 240 to 360 nm, we adopted an undulator wavelength of ~ 16 cm (chosen to agree with that of the LBNL ALS W16 undulator). There should be no difficulty producing the necessary a_w for $\lambda_w \geq 10$ cm. The exact value for FERMI should be driven more by engineering design and cost considerations than becoming fixated on a particular value from a theory point of view. For the second stage modulator of FEL-2, we nominally chose 65 mm to agree with that chosen for the first stage radiator (see the discussion in the next paragraph). To reach the final operating range of 10-40 nm, the upper wavelength need be only 80 nm that in theory could permit a somewhat shorter period if desired for other reasons.

The choice of undulator wavelength for the radiators is driven mainly by two requirements: (1) FEL resonance be physically possible at the longest desired output wavelength (*i.e.*, 100 nm for FEL-1 and 40 nm for FEL-2) for e-beam energies of 1.2 GeV; (2) there be reasonably strong gain (*i.e.*, $a_w \geq 1$) at the shortest desired output wavelength. Requirement #1 drives one to small gaps (for large a_w) and longer wavelengths. Requirement #2 pushes one to shorter wavelengths (but as λ_w decreases, the maximum possible a_w for a given gap opening g begins to decrease exponentially and requirement (1) becomes difficult to meet). At present, we believe the minimum gap opening must be ~ 10 mm; this is subdivided into an 8-mm e-beam “stay clear”, a 1-mm pipe thickness, and 1-mm clearance. With all this in mind, a 65-mm wavelength was chosen for the FEL-1 radiator and first FEL-2 radiator and a 50-mm wavelength for the final FEL-2 radiator. If the upper wavelength of FEL-2 were to be reduced to 30 nm (thus permitting a shorter undulator wavelength), a_w at 10 nm would increase by $\sim 20\%$ from 1.1 to 1.33 and at 7-nm output wavelength by $\sim 30\%$ from 0.74 to 0.96.

Note that if the maximum electron beam energy were to drop to 1-GeV or less, it might be necessary to reconsider the 50-mm wavelength for the second FEL-2 radiator. However, reducing λ_w to 40-mm or less might preclude reaching 40-nm as the maximum operating wavelength for FEL-2 (because of the smaller maximum a_w possible). Even at 1-GeV energy, the FEL-1 radiator a_w exceeds 1.0 at $\lambda = 40$ nm.

2.2 Undulator Segmentation and Focusing issues

In order to produce high powers, the necessary active radiator lengths for FEL-1 and FEL-2 are ~ 15 -30 m. For many reasons (*e.g.*, magnetic forces, alignment, external focusing and diagnostic needs, possible tapering needs), this is far too long to construct as one continuous magnetic structure. Consequently, the radiators will be subdivided into a number of modules, each consisting of an active undulator segment and a break segment with the latter containing a number of items such as quadrupoles, a longitudinal phase shifter, beam position monitors and dipole correctors, and diagnostics. Our preliminary discussions with the undulator and diagnostics teams indicate that ~ 1.0 m will be needed for the break segments; the exact distance will eventually be determined by more detailed engineering considerations. In order to keep the fraction of space occupied by the active magnetic segment a reasonable value (*i.e.*, ≥ 0.5), we have tentatively chosen lengths of 2.34 m (= 36 periods) for the FEL-1 and 1st FEL-2 radiators, and 2.40 m (= 48 periods) for the 2nd FEL-2 radiator. These are the “active”, full strength lengths and do not include the 2-3 poles at the beginning and end of each undulator segment needed for adiabatic matching. While longer undulator segment lengths are permitted in terms of magnetic force considerations, the desire for a Twiss beta function ~ 7 -10 m precludes total module lengths much longer than 3.5 m.

The coupling between the radiation and the electron beam can depend strongly upon the latter’s radius. The FEL radiation emissivity scales directly with the e-beam density and thus there is a premium for minimizing the beam radius. However, as the electron beam radius is reduced, transverse velocity spread and diffraction effects both increase which can reduce the e-beam-radiation coupling despite the larger current density. Consequently, for a given emittance and radiation wavelength, there can be an optimum e-beam radius for maximizing energy extraction. For FERMI, this is generally in the vicinity of 100 microns ($r_{RMS} = (\langle x^2 \rangle + \langle y^2 \rangle)^{1/2}$). For normalized e-beam emittances ~ 1.5 mm-mrad, this radius is much smaller than that obtainable with the “natural” focusing of the undulator, especially for FEL-2 where a_w is generally much smaller than is true for FEL-1. Consequently, the FERMI design includes external quadrupole focusing to produce an average value of 10 m for the Twiss beta function in each plane.

Another function of the external quadrupoles is to compensate for changes in natural undulator focusing both as a_w is changed in conjunction with changing the desired output wavelength, and, at a fixed wavelength, as the undulator polarization is changed (*e.g.*, from vertical to circular). The APPLE-type undulators have the “interesting” feature that the focusing for circular and vertical polarization can be negative (*i.e.*, defocusing) in the x -plane; for long undulators such as the final radiator in FEL-2, this must be compensated for by quadrupoles (or other forms of external focusing). Thus, the operational control system will need to actively modify the quadrupole strengths as non-negligible changes are made in the undulator gaps and/or polarizations.

Once the decision has been made to shift output wavelengths by changing the undulator gap opening (*i.e.*, a_w values) as opposed to changing the electron beam energy (as is true for fixed gaps devices such as the

DESY FLASH and the SLAC LCLS facilities), one must actively control the longitudinal phase slip between the electron beam and radiation in the breaks between undulator sections. The phase advance in a simple drift section of length L_B , $\Delta\theta = k_w L_B / (1 + a_w^2)$, is not necessarily a multiple of 2π as the FEL wavelength λ_s and thus a_w are varied. Consequently, a very weak magnetic chicane (strength $R_{56} \sim 2\lambda_s$) is needed to act as a “phase shifter” in the break section. Based upon discussions with the undulator group, we believe ~ 10 cm of longitudinal space will be needed for this element.

2.3 Undulator Error Tolerance Calculations

Apart from electron beam errors such as offset, tilt, and mismatch, there are also errors which are possible within the undulator. Some of these include: 1) tilt and offsets of entire undulator segments 2) “global” segment mistuning errors such that the average a_w is offset by a constant amount within each segment (e.g., due to an incorrect gap setting 3) “local” undulator errors due to individual pole strength errors which can lead both to longitudinal phase errors between the electron beam and the FEL radiation and to the electron beam wandering offset the undulator and radiation central axis.

Regarding issue (1), tilt and offset in the undulator to lowest order are equivalent to equal and opposite values in the initial electron beam position and tilt. Sensitivity to such is discussed in Sec. 3.3 and we do not discuss it further other than to note that in a multi-segment undulator setup that, depending upon the exact method of undulator alignment, the effect of these errors might in a statistical sense grow as $\sqrt{N}$ where N is the number of segments. Hence, if a criterion for beam tilt and offset is a value Y , then the equivalent RMS criterion for the individual segments might need to be reduced to $Y/\sqrt{N}$. On the other hand, there will be active dipole correctors between segments so this estimate may be unduly pessimistic and severe.

“Global” segment mistuning will lead to a longitudinal phase error growing with z . If this error becomes comparable to $\sim \pi/2$ radians, there can be significant loss of FEL radiation gain. On the other hand, smaller scale errors in some circumstances can lead to increases of radiation power over what is possible from a constant K undulator; this appears due to a serendipitous tapering which extracts more power. In order to obtain a rough criterion for how accurate a_w (equivalently the gap opening) needs to be set in an RMS sense for the FEL-1 radiator and the final radiator of FEL-2, we did a series of GINGER runs in which for each run a random mistuning with a given RMS expectation value was applied to each individual radiator segment. For FEL-1 at 40-nm output wavelength, the results (see Fig. 2.3-1) show that in an average sense, the RMS segment mistuning error in a_w must exceed 0.002 before the output power begins to drop more than a few percent. We believe that this constraint should be relatively easy to meet for the FERMI undulators. For FEL-2, the results (see Fig. 2.3-2) are quite similar with RMS errors below 0.002 showing essentially no effect on the average output power at 10-nm and the standard deviation remains less than 10%.

To examine the effects of “local” errors, we used a simple numerical code (XWIGERR) to generate random sets of pole strength errors whose distributions followed simple Gaussians. Within each undulator period, the two individual pole errors of one undulator period are decomposed into an “even” component which leads to no net transverse kick but does induces a phase error through a net change in K , and an “odd” component that does produce a net transverse kick on the beam. This odd component causes the electron beam both to wander off-axis and also suffer a net phase error (because there is an accompanying decrease in $\langle v_z \rangle$ associated with the increase in $\langle v_\perp^2 \rangle$). In the absence of any correction, the size of the beam wander tends to increase steadily with z and as does the longitudinal phase error.

However, in the actual physical situation, for a given sorting of the individual poles within a given undulator segment, the strength errors are then “frozen” (i.e., do not vary on either a shot-by-shot basis nor a day-by-day basis) and may therefore be corrected in some average sense in z by the insertion of dipole shims. After discussion with the FERMI@ELETTRA undulator group, we have implemented “virtual” shims in the XWIGERR code by using the following algorithm:

- 1) First calculate a set of random pole errors for each undulator segment, decompose those errors period by period into even and odd components, and then calculate the uncorrected longitudinal phase and transverse trajectory with z .
- 2) Within a predetermined number of evenly spaced intervals for each segment, calculate the necessary shim corrections to the first and second integrals of $B_{\perp}(z)$ to remove the tilt and offset at the end of each interval.
- 3) With the presence of these shims, recalculate the longitudinal phase as a function of z in each sub-segment interval and then the necessary change in a_w within that sub-segment by the appropriate amount to bring the net change of phase over the sub-segment to zero.
- 4) Write out a shimmed undulator lattice file giving the net kick and Δa_w as functions of z to be used by the GINGER FEL simulation code.

In Figs. 2.3-3 through 2.3-8, we show time-independent GINGER calculations for output power sensitivity to undulator pole strength errors for FEL-1 at 40-nm and FEL-2 (fresh-bunch approach) at 10-nm; two virtual shims were used per undulator segment. Each point on these plots corresponds to a different random pole strength error set. The results are plotted versus RMS residual (after shimming) phase error, tilt, and offsets. Each residual is a function of z and is calculated only within the individual segments. One sees that there is a relatively tight correlation between the residual phase error and the output power. For FEL-1, here is virtually no loss of power for residual phase errors below 0.2 radians; for FEL-2 at 10-nm the equivalent value is 0.15 radians. The correlations are less for tilt and offset residuals. For FEL-1, it appears that there is little loss of power for 20-micron offsets or less and 50-microradian tilts; for FEL-2 the equivalent numbers are approximately 10 microns and 25 microradians.

2.4 Post-modulator dispersive section design and issues

Following each modulator is a break section that contains a magnetic chicane whose chromatic dispersion is used to develop a strong coherent microbunching from the energy modulation impressed upon the electron beam by the FEL interaction in the modulator. For reasonably large input seed powers (e.g., ≥ 10 MW) and short wavelengths ($\lambda_0 \leq 300$ nm), the necessary R56 dispersion parameter is ~ 100 μm or less in the first stage modulator. For the fresh-bunch approach, the second stage dispersive element is typical ~ 5 times smaller. Some preliminary design work for the dispersive sections has been done by the FERMI undulator group and no significant engineering or space issues have been reported. To date, there has not been any detailed analysis upon the size limits for higher order optics terms, such as would be induced by fringe fields and/or geometric aberrations. However, we believe there is unlikely to be any practical problem given the allowed longitudinal space of ~ 30 -cm or more and the relatively large seed wavelength. Depending upon the actual design of the dipole elements of this chicane, there may be a transverse focusing effect that eventually should be properly modeled by the FEL simulation codes.

2.5 Delay section needs and issues for FEL-2 FB approach

In the fresh-bunch approach to FEL-2, it is necessary to delay the e-beam by $\sim 500 - 1000$ fs relative to the FEL radiation in order that a “fresh” section of the e-beam be energy-modulated in the second stage

modulator. A 1-ps delay is equivalent to an R_{56} of 600 μm . Since there is no concern of maintaining any microbunching between harmonic stages inasmuch as the radiation field carries the “imprinting” signal, we believe the allowed longitudinal space of ~ 1.8 m is more than adequate to contain the needed chicane. The delay section must also contain various diagnostics and at least a quadrupole singlet (and possibly doublet) for rematching the e-beam to the second stage optics. For purposes of FEL simulation, it is important to include diffraction effects as radiation transits this section.

2.6 Vacuum Chamber Longitudinal Wakefield Considerations

As the electron beam travels through the metallic vacuum chamber within the undulator and break sections, it will feel the effects of time-dependent wakefields whose strength normally linearly scales with the beam charge. Of particular concern for FEL interaction are the effects of longitudinal wakes that in general will decelerate the e-beam but, depending upon the characteristics of the wake function and the temporal current distribution, may actually accelerate some temporal regions of the pulse. The wake function among other things includes a resistive term arising from finite “DC” and “AC” wall conductivity, a surface roughness term arising from imperfections on the inner wall, and a geometric term arising from any cross-section interruptions such as pumping bellows. Since the FEL output is sensitive to energy mistuning of $\sim 0.05\%$ (*i.e.*, 600 keV on a 1.2 GeV beam) and a typical total undulator lattice length is ≤ 40 m, a temporal variation on the wake of ~ 15 keV/m could pose problems. Sample calculations are presented in Secs. 3.4.3 and 4.10 for FEL-1 and FEL-2, respectively.

3 FEL-1 Design & Performance Calculations

Table 3-1 - Basic Beam and FEL-1 Undulator Parameters

Input Seed Laser	
Power [MW]	100
Wavelength [nm]	240 - 300
Waist size [microns]	300
Input Electron Beam	
Energy [GeV]	1.2
Current [A]	800
RMS energy spread [keV]	200
RMS emittance [mm-mrad]	1.5
Modulator Undulator	
Period [m]	0.16
Length [m]	3.04
Number of periods	19
Radiator Undulator	
Period [m]	0.065
Section length [m]	2.34
Break length [m]	1.04
Number of sections	6
Total length [m]	20.28
FEL parameter ρ	2.9×10^{-3}

3.1 Basic design

In Table 3-1, we present the base seed laser, electron beam, and undulator characteristics adopted as a design point for FEL-1. Because at present it is believed that most users of FEL-1 are interested in pump-probe and other experiments in the time domain, we have chosen electron beam parameters characteristic of the so-called “medium pulse” option where the main body current is ~ 800 A, the flat-top duration is ~ 700 fs (which makes allowance for timing jitter), and the total charge is ~ 0.7 -nC. We chose a (hopefully) conservative value of 200 keV for the incoherent energy spread. For FEL-1, the output is relatively insensitive to the actual value in this region so long as the input laser power and modulator length are sufficiently long to produce a coherent energy modulation ΔE an order of magnitude greater than σ_E . We adopted a nominal laser power of 100 MW; lower values could be used although the modulator length would need to be increased over the ~ 3 -m used here. The laser comes to a focus halfway in the modulator with a waist size of 300 μm . This value is significantly greater than the e-beam transverse size and was chosen to minimize the amount of induced incoherent energy spread.

Figure 3.1-1 illustrates the beam line lattice for FEL-1 which includes a modulator, dispersive chicane, and the radiator sections consisting of active undulators and breaks. The electron beam and seed laser enter from the left. The drift length between the modulator exit and radiator entrance, and between the individual radiator sections, is 1.04 m, which must include the space associated with the partial strength poles at entrance and exit. By the end of the modulator, the peak-to-peak energy modulation $2 \Delta E$ is about 4 MV and the RMS energy spread is about 1 MV. A simple scaling argument ($R_{56} \times \Delta E/E = \lambda_{\text{MOD}} / 4$)

suggests that the necessary R_{56} is about 35 microns, close to what was actually adopted (see Section 3.2). For purposes of FEL gain calculation, each radiator undulator section in the simulations is 2.34-m long and is composed 16 *full-strength* 6.5 cm periods. We assumed that 2 periods or fewer are required for adiabatic transition to and from each break, resulting in a physically-usable drift section length of 0.84 m. For calculation purposes, each drift section was presumed to include a “perfect” phase shifter that ensures the longitudinal phase slippage in the drift is an exact multiple of 2π . Since photon number is a critical parameter for FEL-1, the nominal layout includes sufficient sections (6) in the final radiator to ensure power saturation at the shortest design wavelength (40 nm) although as few as three sections are necessary for saturation at 100-nm wavelength.

Our general procedure for optimizing the undulator parameters for each wavelength was as follows. The normalized modulator strength a_w ($= K/\sqrt{2}$ for a linearly-polarized undulator) was set to the nominal FEL resonance value. Then we optimized the radiator performance relative to values of the dispersion parameter R_{56} and a_w . Figure 3.1-2 displays the growth of power and coherent microbunching at 40-, 60-, and 100-nm wavelengths as predicted by the GENESIS and GINGER codes for the parameters of Table 3-1. At the longer wavelengths power and bunching saturation was reached well before the end of the sixth section. This figure also shows there is good basic agreement between the GENESIS and GINGER predictions. One does expect some differences to arise from different numerical integration schemes for the wiggler-period-averaged FEL equations and possibly also from different radial grid resolutions. Nonetheless, for both these FERMI studies and others associated with the LCLS there do not appear to be any serious systematic differences in predictions by the two codes. The peak power of 2.5 GW at 40 nm corresponds to an extraction efficiency of 0.26%, quite close to the 3D FEL parameter of 2.9×10^{-3} (which likely is to be an overestimate because the Ming Xie fitting formula neglects the physics appropriate to the 1-m break sections).

3.2 *Time-independent Simulation Results*

In order to obtain a rough feeling for the output power sensitivity to electron beam and laser parameters, we performed an extensive set of GENESIS and GINGER simulations varying *single* parameters at a time. The studies described in this section were limited to only axisymmetric results for both the electron beam and input laser; Section 3.3 discusses sensitivity to non-axisymmetric effects such as an input transverse offset or tilt of the entering electron beam. Each individual beam or seed parameter was varied around a central values for the “medium pulse” case as given in Table 3-1. The studies here were done in the “time-independent” or “time-steady” limit in which any and all properties of a time-varying electron and laser pulse are replaced by a single, representative value. This approximation allows one to model performance with just a single longitudinal “slice” and is thus computationally much faster than a full time-resolved simulation. For calculation purposes, each drift section was presumed to include a “perfect” phase shifter that ensures the longitudinal phase slippage in the drift is a multiple of 2π .

3.2.1 **Single Parameter Sensitivity Studies at 100 nm**

To obtain output at 100-nm wavelength, we presumed an input seed laser wavelength of 300 nm and a 3:1 harmonic upshift from the modulator to the radiator undulator. The individual focusing quadrupoles (assumed to have 20-cm effective length) in the radiator were set to the values 100, -100, 100, -100, 100 G/cm to give an average beta function of 10 m in both x and y within each undulator section (horizontal linear polarization was presumed for this study). The corresponding initial β_x and β_y at modulator entrance were 8.75 and 12.6 m, respectively. Due to the very high electron beam quality (*e.g.*, normalized emittance and energy spread), FEL-1 performs extremely well at 100 nm.

Before beginning the input parameter sensitivity scans, we first optimized the radiator performance relative to values of the dispersion parameter R_{56} and normalized undulator strength a_w . As shown in Figure 3.2-1, there is a relatively broad peak in power versus R_{56} . For variations in a_w , there is a relatively much narrower peak centered at 3.98. This narrowness can be directly translated to the acceptable range in electron beam energy centroid; *i.e.*, $\Delta\gamma/\gamma \sim (a_w^2 / [1 + a_w^2]) \Delta a_w / a_w$. Based upon these results, for the following parameter sensitivity studies we adopted $R_{56} = 40 \mu\text{m}$ and $a_w = 3.980$.

Only 3 radiator undulator sections are needed for radiation power saturation at a level of somewhat greater than 4 GW for the optimized parameters. The microbunching fraction (see the 100-nm curve in the right plot of Fig. 3.1-2) jumps above 30% following the dispersive section and then grows to a peak of nearly 60% by the middle of the third undulator section. This level quantitatively is about as large as one ever achieves in normal FEL interaction.

Due to the small emittance and the fact that we are not optimizing the focusing strength as the emittance varies, there is only a weak dependence (see Fig. 3.2-2) of output power upon this parameter. There is a much stronger dependence upon energy spread, especially well before saturation. At saturation and beyond, the dependence becomes quite weak. With respect to beam current, one sees a nearly I^2 dependence; this is expected in the strong bunching limit because the FEL emission is essentially coherent spontaneous emission which scales as $I^2 b^2$.

Figure 3.2-3 shows the equivalent plots for variations in electron beam energy and input laser power. Again because of the high beam quality and large FEL parameter ρ , there is only a weak dependence upon beam energy with the power dropping by only 25% for a fractional change in energy of 0.25%. As the input laser power is varied, there is a very weak dependence at small z that essentially disappears when saturation is reached at the end of the third undulator.

Figure 3.2-4 shows far field radiation diagnostics extracted from GINGER results. The mode quality M^2 never gets much better than a value of 2.7 indicating that the output is significantly contaminated by higher order modes. We believe the reason for this is the small Rayleigh range associated with the electron beam size so the self-similar mode is very non-Gaussian. If a better quality mode is desired, it is possible that by decreasing the quadrupole focusing (and thus increasing β and the electron beam size) would lead to a better output mode profile and smaller M^2 . The virtual emission point is quite insensitive to various e-beam parameters but the waist size seems monotonically dependent upon both e-beam energy and current.

3.2.2 Single Parameter Sensitivity Studies at 40-nm

To produce 40-nm output, the input seed wavelength was set to 240 nm and the radiator FEL resonance tuned to the sixth harmonic; an alternative set would be 200 nm and the fifth harmonic. As with the 100-nm studies, we determined the optimal R_{56} and a_w parameter for peak performance (Fig. 3.2-5). Here we found that five undulator sections are necessary to reach first saturation of the 40-nm radiation power (see Fig. 3.1-2), approximately two more than was true at 100-nm. The saturated output power is ~ 2.7 GW, $\sim 70\%$ of 100-nm saturation value. The parameter sensitivity scans (Figs. 3.2-6 and 3.2-7) show nothing particularly surprising, with the HWHM of energy tolerance being $\sim 0.2\%$, much narrower than was true at longer wavelengths. Regarding other electron beam parameters, until saturation the power appears quite sensitive to the initial energy spread and current (Fig. 3.2-7).

Figure 3.2-8 shows far field radiation diagnostics extracted from both GINGER and GENESIS results. In

general, the agreement between the two codes is excellent; some differences are to be expected because of the non-identical algorithms used to extract the virtual waist size and location. As shown in the upper left plot, the mode quality M^2 together with the output radiation power becomes very degraded as the beam energy differs by more than 0.2% relative to the value corresponding to peak output power. However, at saturation, both the output power and mode quality factor M^2 are relatively invariant with respect to σ_E and current. Note that the best value of M^2 is about 1.6 at power saturation, significantly better than the values found at longer wavelengths. This improvement is likely to be due to the greater Rayleigh range (*i.e.*, less diffraction) at 40-nm wavelength. The effective emission point monotonically increases by ~ 5 m as σ_E is increased from 100 keV to 300 keV.

3.2.3 Time-independent Calculations for Predicted Jitter

In order to develop a rough estimate for the expected shot-to-shot jitter in the output power and photon number from FEL-1 at 40 nm (*i.e.*, the shortest FEL-1 wavelength because we expect the sensitivity to be greatest there), we did some time-independent calculations in which the input laser seed power and various electron beam quantities were allowed to vary *independently* around their individual design values following a tolerance budget summarized in Table 3-2. The modulator segment was tuned to $a_w = 3.914$ for resonance at 240-nm wavelength, and the laser seeding power was set to 100 MW with a waist of 300 μm in the center of the modulator. The following dispersion section was set to R_{56} of 19 μm and the quadrupole strengths adjusted in order to have an average beta function in the radiator of about 10 m. The electron and seed beam parameters were those given in Table 3-1 except that the input incoherent energy spread was reduced to 100 keV from 200 keV; the change resulted in ~ 2.8 GW of FEL radiation at saturation.

Table 3-2 - Adopted Shot-to-Shot Variation Budget

Parameter	Normalized shot-to-shot variation
Emittance	10%
Peak current	8%
Mean energy	0.10%
Energy spread	10%
Seed power	5%

Two particular sets of calculations were done. First we consider fluctuations on a single parameter only. For each of the different electron parameter (*e.g.*, energy, current, *etc.*), we created a Gaussian distribution of 50 parameter values with the corresponding standard deviation (Table 3-3). We then used each of these values in different FEL-1 simulation runs with the GINGER code to initialize the electron beam (or input seed laser). These runs produced the data for the red curves of Figs. 3.2-9,11,13,15,17 and the $\langle P(z) \rangle$ and $\sigma_{P(z)}$ curves shown in Figs. 3.2-10,12,14,16,18. A second set of calculations were done with simultaneous multiparameter jitter in which a set of 400 parameter values were created where each and all beam parameters were randomly varied following the appropriate Gaussian distribution.

We first consider the effect of a jitter only in the mean electron energy. According to the FERMI linac group, a Gaussian distribution with a normalized rms of 0.1% is the design goal for beam energy at the end of the linac. The single parameter sensitivity scans discussed above in Sec. 3.2.1 showed that γ plays a crucial role in the FEL performance of FERMI. The multiparameter results (see Figure 3.2-9), although

they show scatter due to the other parameters being varied, remain very well correlated to the electron energy variation. Examination of Fig. 3.2-10 clearly shows that the normalized output power fluctuations increase along the radiator and, importantly, the fluctuation level does not significantly decrease upon entering the saturation regime.

For the case of current jitter (see Figs. 3.2-11 and 3.2-12), the output power monotonically climbs with increasing current and the normalized output power fluctuations become significantly reduced when approaching saturation; this property should be considered when trying to minimize sensitivity of the FEL output to input parameter jitter. As was true for the electron beam energy, there is a clear correlation between multiparameter jitter output power and the input electron beam current.

For the adopted central design point, FEL-1 shows a very small sensitivity to the electron beam incoherent energy spread. As is apparent from Figs. 3.2-13 and 3.2-14, although the FEL gain before saturation can depend strongly on the value of the energy spread, the value of the saturation power is much less sensitive. In the case of multiparameter jitter the contribution of the energy spread to the total power fluctuations is very small and there is no correlation between output power and energy spread. There is similarly very small sensitivity at saturation with respect to transverse emittance. The power fluctuations due to emittance jitter are significant and almost constant in the first part of the radiator but drastically decrease when approaching power saturation (Figs. 3.2-15 and 3.2-16).

Table 3-3: FEL-1 40-nm Power Fluctuations or single and multi- parameter jitter

Parameter	Shot-to-shot variation	Output power fluctuations
Mean energy	0.10%	13%
Peak current	8%	9.4%
Emittance	10%	1.7%
Energy spread	10%	0.3%
Seed power	5%	1.3%
Multiparameter		22%

FEL-1 also shows little sensitivity at saturation to the seed laser power. While the normalized fluctuations are high at the beginning of the radiator (because the bunching immediately following the dispersion section is tightly correlated to the input seed power), their level (see Figs. 3.2-17 and 3.2-18) monotonically decreases with a minimum value occurring close to the saturation point in z . Also, there is little correlation between the multiparameter results and the input seeding power.

Table 3-3 and Fig. 3.2-19 summarize the output power jitter results from these time-independent studies. It is clear for the adopted distributions given in Table 3-2 that the e-beam energy and current are the main parameters responsible for output power fluctuations. Consequently, for a significant reduction of the total amount of jitter for FEL-1 it is necessary to reduce the sensitivity to these parameters.

3.2.4 Undulator Taper Strategies for Reducing Sensitivity to Beam Energy

One possible way for reducing FEL-1 output power sensitivity to the mean electron energy is to enlarge the effective energy bandwidth of the radiator. This may be possible by using different a_w values in different radiator undulator sections. With this approach, electron bunches with a mean energies slightly different from the nominal value will encounter some section whose a_w is close to FEL resonance.

The simplest variation strategy is to monotonically decrease or increase the undulator strength section by section. In particular, in order to increase the gain similarly to what is done by tapering in and beyond the saturation region in an FEL, one would monotonically decrease a_w along the radiator. The red dotted curve of Fig. 3.2-20 shows one of the best undulator profiles we have found for this approach in terms of reducing the output power sensitivity to electron beam energy. Figure 3.2-21 plots output power plotted versus beam energy and shows a clear increase in the energy bandwidth; the power fluctuations due to energy jitter are reduced from 13% for a constant undulator strength down to only few %. However the results become worse if we consider multiparameter jitter; in that case the output power fluctuations are in both cases larger than 20%. The reason for this is the monotonic taper strongly increases sensitivity to the electron current.

In order to reduce sensitivity to the mean electron energy without simultaneously increasing sensitivity to electron beam current, we investigated a second simple tapering configuration. Section by section, we alternatively set a_w to higher and lower values (see Fig. 3.2-20, dashed blue line) relative to a constant tapering with z . To lowest order, this type of variation neither favors higher nor lower beam energies (relative to the nominal value). To the contrary, the previous strategy favored higher energies and higher currents that reach saturation earlier in z . Consequently, we hoped that this new strategy would reduce the effect sensitivity to current also. The alternating strength strategy appears to work well; plots of the output power versus the electron beam energy displayed in Fig. 3.2-21 clearly show the reduction of sensitivity to beam energy relative to that shown by the constant strength radiator. For this new configuration the normalized standard deviation of the output power for the adopted energy jitter is less than 5%, and, more importantly, the power fluctuations are reduced also for the multiparameter jitter case (see, *e.g.*, Figs. 3.2-11, 3.2-13, and 3.2-19).

3.3 *Input Transverse Tilt and Offset Sensitivity Studies for FEL-1*

3.3.1 Introduction

Displacements and tilts of the electron beam are an important consideration for studying the performance of the FERMI FEL. Electron offsets can occur due to upstream pointing errors, undulator misalignments, or to internal structure in the electron bunch arising from time-dependent linac wakefields, in which case they will also be sensitive to timing jitter. FEL performance in the presence of such offsets was modeled with the GENESIS code because a fully three-dimensional field solver is necessary to capture all non-axisymmetric effects. Our simulation studies included initial offsets for the electron beam only; the laser seed and undulators were assumed to lie along a common axis. Time-independent results for FEL-1 are presented at different output wavelengths, focusing on examples for 100-nm wavelength both with and without undulator tapering, and for 40-nm wavelength without tapering. The most prominent effect of electron beam offsets is a large drop in output power when the transverse overlap in the first undulator between the electron beam and input radiation seed drops significantly, although in addition the FEL output develops offsets comparable to those of the electron beam. Information on predicted output phase is also included; a strong phase variation with longitudinal position can have a large impact on the output radiation spectral width.

“Global” sensitivity studies simulating various types of jitter simultaneously but in this section also including jitter arising from initial tilt or offset, were performed for the untapered configuration at various wavelengths. We find that again, as was true for the jitter studies without offsets or tilts (see Sec. 3.2.3), the most significant source of predicted shot-to-shot fluctuations is jitter in the electron beam energy. However, when using expected values for fluctuations in the electron beam and laser seed power, the combined effect of the jitter in the other FEL parameters is predicted in some cases to be comparable to

the effect of energy jitter. For untapered undulators, the predicted output power fluctuations (normalized standard deviations) are 13% at 100 nm, 24% at 60 nm, and 28% at 40 nm; tapered undulators show less output sensitivity to electron beam energy jitter. These sensitivity studies also reveal correlations between beam parameter errors and output power. In addition to some apparent mistuning of a_w so that the nominal energy does not quite yield the optimal power, which would be easily corrected in practice, the only significant correlation is between the beam transverse emittance and the output power. Rather than reaching a local maximum at the nominal electron emittance, the performance improves significantly as emittance is decreased.

The base electron beam and radiation parameters used for the offset and tilt studies were essentially the same as used in Sec. 3.2: 1.2-GeV energy, 800-A peak current, 200-keV rms incoherent energy spread, and 1.5-micron rms normalized emittance. The beam optics focus the electron beam to an average Twiss beta function of 10 m. Up to 6 undulator sections are used in the FEL-1 radiator, fewer for the untapered cases; for example, the 100-nm wavelength case peaks in power after 3 sections. For the tapered 100-nm wavelength cases, each undulator section has a constant a_w but the value decreases from section to section. Specifically, a_w drops by a fixed amount in each of the first three undulator sections, then by a fixed but larger amount in the final three undulators. Tapering effectively extends the undulator length that remains resonant with the electron beam as it loses energy to the radiation. The input seed power was 100 MW with a Gaussian waist of 210 μm (in intensity) halfway through the undulator.

First, we review the nominal performance in terms of radiation power and bunching for a variety of cases. The results are shown in Fig. 3.3-1. For the 100-nm case, the maximum power is 4 GW, and the peak bunching is close to 0.6; the initial bunching after the dispersive section is 0.33. Tapering the magnetic strength in three additional 2.34-m undulator sections increase the output power by a factor of roughly 2.5. Figure 3.3-2 plots the evolution of the radiation spot size with z . Here, we can see the effect of diffraction in the breaks between undulator sections, as well as gain-guiding effects that partly overcome diffraction and refocus the radiation beam. This effect is most noticeable at 40-nm wavelength where the peak power is 2.5 GW and the peak bunching grows to 0.5 from a post-dispersion section value of 0.2. At 40-nm wavelength, tapering the undulator and using 6 segments yields 4.8 GW of power.

3.3.2 Offset and Tilt Sensitivity Studies Using GENESIS

These sensitivity studies focused upon analyzing the final output radiation as functions of electron beam offset in position and angle. Plots are given of the output power and phase, the maximum power that can be contained within a pure Gaussian mode, the transverse mode quality (given in terms of M^2), the divergence angle, the location of the virtual waist and the radius of the waist. The latter properties are determined by analytically refocusing the output radiation to a waist. The minimum spot size that can be achieved determines the optical quality parameter M^2 , which has a value of unity for a pure TEM00 Gaussian mode. The power contained within a single, pure Gaussian mode is calculated independently by determining the mode properties that would have maximum overlap with the GENESIS-calculated output radial profile. However, the corresponding Rayleigh length and waist location of this mode does not in general correspond to the calculated virtual waist described above except in the limit where M^2 is close to 1. Note that all green curves in the following graphs are plotted against the right axis. It is clear from the following graphs that the fraction of power in a pure Gaussian mode is not determined by the M^2 parameter.

The most appropriate figure of merit depends upon the optical and experimental configuration downstream. There are two additional factors to consider in this analysis that do not apply to an ideal radiation beam. When the electron beam has an offset, the radiation field need not be centered on axis

either; consequently, we include a separate calculation of the mean center and orientation (*i.e.*, tilt) of the output radiation. Other quantities are then calculated with this mean offset subtracted out. The one exception to this rule is the calculation of the power in a pure Gaussian mode. Here we still assume that the radiation is exactly aligned with the undulator axis, an assumption that can seriously affect the calculated results. Thus, beams that enter with a large tilt may exhibit an artificially low value for the fraction of power in a pure Gaussian mode while at the same time have a calculated M^2 that is (properly) close to 1. The fractional power diagnostic, however, may be pertinent to experiments that are sensitive to shot-to-shot transverse jitter in the FEL output.

Figures 3.3-3 and 3.3-4 show the resulting analysis of the mode properties for the baseline 100-nm case as a function of a pure initial offset or tilt, respectively, in the electron beam at entrance to the modulating undulator. The actual displacement of the output radiation is shown in Figs. 3.3-5 and 3.3-6. Note that the effect of an initial tilt inside the modulator on the output radiation properties is extremely small, despite the fact that the peak tilt value considered, 40 microradians, is comparable to a 400-micron displacement for a 10-m beta function. The displacement of the electron beam within the modulating undulator is, in contrast, extremely important due to the change in the spatial overlap fraction between the laser seed and the electron beam. Electron beam displacements within the modulator comparable to or greater than the seed radiation spot size cause significant reductions in both the energy modulation and the emitted power from the downstream radiator. Nonetheless, the actual misalignment of the output radiation is, as expected, comparable to or smaller than the maximum offset or tilt in the electron beam itself. Note that the effect of offsets on output phase is very small, with at most a 1.5-radian shift for the largest offsets considered. For all of the configurations considered below, the phase shift remains near or below π radians, and is primarily due to path-length changes from geometric effects. This suggests that similar tolerances will apply when considering offsets between different longitudinal slices of the electron beam.

In Figs. 3.3-7 through 3.3-10, we show equivalent results for the tapered undulator configuration at 100-nm wavelength. There are two main differences brought about by the combination of tapering and the increased length of undulator. First, the sensitivity to electron beam offsets is greatly reduced because a smaller initial energy modulation resulting from electron beam offsets still has enough nonlinear growth possible in the longer undulator to “catch up” with the nominal aligned case, whose undulator and dispersion section parameters were tuned for a larger energy modulation. Additionally, the relationship between the initial offset and tilt of the electron beam to that of the output radiation is altered because the electron beam undergoes more of a betatron oscillation in the longer radiator.

Figures 3.3-11 through 3.3-14 display the output radiation and mode properties at 40-nm wavelength as functions of input offset and tilt. The reduction in power due to electron beam offsets is similar to that of the untapered case at 100 nm. For a 300-micron initial transverse offset, the power drops by more than a factor of two in both cases, as compared with a 15% drop for the tapered, 100-nm wavelength example.

Table 3-4 : Nominal design value and presumed RMS jitter for FEL-1 parameters

Parameter	Nominal value	Assumed jitter
Laser seed power	100 MW	5%
Beam energy	1.2 GeV	0.1%
Peak current	800 A	8%
Energy spread (rms)	200 keV	10%
Normalized emittance	1.2 micron	10%
Average beta function	10 m	0
Transverse size	60 – 80 micron	5% (from emittance)
Angular spread (rms)	6 – 8 microradian	5% (from emittance)
Position offset	0	100 micron
Angle offset	0	10 microradian

3.3.3 Tolerance Studies with Offset/Tilt: Simultaneous Variation of Multiple Input Parameters

This section discusses GENESIS calculations at both 100- and 40-nm output wavelength detailing the impact on FEL-1 output power of the simultaneous variation of the input parameters, including beam transverse offsets and tilts. The nominal working point and the standard deviation for each electron beam and input laser parameters are given in Table 3-4. In all following jitter studies, the input Twiss parameters were kept fixed at the nominal values and the undulator strength was untapered --- see Sec. 3.2.4 for details on how tapering can reduce the sensitivity to electron energy jitter

For 100-nm output, we first consider the case where we vary *only* the energy and hold all the other parameters constant at the nominal parameters; this case is equivalent to the energy scan presented in Sec. 3.2.1. The results are shown in Fig. 3.3-15 where the average power over multiple simulations is 2.82 GW and the standard deviation is 9.7%. Other FEL output characteristics vary by different fractional amounts; for example, the radiation spot size only varies by 6.5%, while the on-axis intensity varies by 18%.

Next, we consider a “global” jitter study with simultaneous variation of all the following parameters: initial energy, current, energy spread, emittance, laser seed power, transverse offset and angle. The results are shown in Figs. 3.3-16 through 3.3-19 where the output power is plotted against initial electron beam energy, emittance, transverse offset and tilt. The average power over all the individual jitter runs is 2.80 GW and the fractional standard deviation is 13%. While the variation of beam energy still has the largest effect on the properties of the output radiation, jitter in other quantities does lead to significant additional variability, including the total output power. In addition, the jitter in transverse position and tilt of the laser output are comparable to the input electron beam jitter in these quantities with proper allowance being made for the phase advance in the undulator and any changes in the Twiss beta function between the modulator entrance and radiator exit.

Figures 3.3-20 through 3.3-23, present 40-nm wavelength results for “global” parameter jitter (e.g., multiple parameters being jittered simultaneously). The average power for these runs is 1.94 GW, and the standard deviation is 28%. The variation of beam energy is clearly the dominant effect, although there is again a noticeable increase in output power as the emittance is decreased.

FEL-1 jitter study results including tilt and offset may be summarized as follows. Performance at the shortest wavelengths shows the greatest sensitivity to jitter in input parameters. This is particularly true for electron beam energy jitter. It is noteworthy that there is also a moderately strong linear correlation between emittance and output power in all examples. This effect is more noticeable than the $\sim$ quadratic dependence of output power upon beam current. This suggests that further improvements could be made by redesigning the FEL to work optimally at the actual beam emittance rather than some nominal value. In the current configuration, there is a significant improvement in performance as the beam emittance is reduced. At 40-nm wavelength, there is also a slight improvement of performance for electron beam offsets of about 100 microns. This suggests that some additional retuning of a_w might increase the predicted output power.

Overall, these tolerance studies show that jitter in electron beam transverse position and tilt has a significant effect on FEL performance. However, for the standard deviations chosen (*i.e.*, 100 microns in position and 10 microradians in tilt), these contributions are not as critical as 0.1% jitter in the electron beam energy. In practice, these different jitters may in fact be correlated, either due to chromatic effects in the linac or to a timing jitter underlying both effects. Our time-dependent jitter study in discussed in Sec. 3.5 does include such correlations.

3.4 Time-dependent Results for FEL-1

3.4.1 Summary of Time-Dependent Simulation Method

While time-steady (*i.e.* monochromatic) simulations can give good indication of the performance sensitivity of an FEL, it is now generally recognized that more-or-less complete “start-to-end” (S2E) simulations that begin at the emitting cathode and end at the undulator exit (for the electron beam) and/or the experimental sample (for the photons) give far more accurate estimates. For the FERMI TOS, the injector and linac groups did extensive modeling and produced a number of so-called “golden” macroparticle files for use as input by the FEL simulation group. Initially, it was hoped a “short pulse” case with a temporal FWHM of $\sim$ 300 fs would be appropriate for use in pump-probe experiments. However, due to the relatively large temporal jitter ($\sim$ 150 fs) expected at the linac exit, we concluded that such a short pulse would lead to unacceptable shot-to-shot output jitter as many seed pulses would fall temporally outside the electron beam pulse. Work then began on a “medium” bunch case whose FWHM is $\sim$ 650 fs and which would thus be wide enough to deal with the temporal jitter.

Both the GINGER and GENESIS FEL simulation codes were used to predict the full time-dependent radiation output corresponding to these macroparticle “golden” files. In general, the FEL codes use a total number of macroparticles per time interval that is greater than that generally available from most ELEGANT output files. In the case of GENESIS, this problem is solved by use of a special macroparticle creation algorithm that creates as needed locally in 6D phase space new macroparticles in the “empty” spaces between the ELEGANT macroparticles. GINGER uses a much different algorithm in which to fully populate a given time slice, ELEGANT macroparticles from adjacent temporal regions are used with their 5D coordinates (x, x', y, y', γ) carefully interpolated in order to maintain their individual deviation from a coarse-grained average in time. In principle, both algorithms should maintain the local time-dependence

of various higher order correlations (*e.g.*, $\langle xy' \rangle$, $\langle \gamma x' \rangle$, *etc.*).

For both codes, it was also necessary to rematch the 4D phase space (x, x', y, y') to the FEL-1 undulator lattice; in operational practice this will be done by a series of dipoles and quadrupoles upstream of the modulating undulator. The rematching was accomplished by determining the Twiss α and β in the central temporal regions of the ELEGANT files, computing the needed transformation matrix to give the correct match, and then applying this matrix to ALL the macroparticles. We used only the temporally-central, “well-behaved” portion because there are often current spikes the head and/or tail regions with “abnormal” phase space properties.

Because it was felt that the 40-nm wavelength performance was the most critical area to model, nearly all our FEL-1 time-dependent simulations have been done at this wavelength. For the input radiation seed, we either used a Gaussian temporal profile pulse with a FWHM of 100 fs (appropriate for pump-probe experiments) or a constant intensity flat-top pulse in which the laser fully covered the e-beam (appropriate for experiments for which maximum photon number but not timing synchronization is needed). In all cases the seed beam had a Gaussian transverse profile with a 210-micron waist (in intensity) occurring at the mid-point of the modulator. The simulations normally adopted temporal slice spacing of either 0.8 fs (*i.e.*, 240 nm) or 1.6 fs. After each modulator run, the particles are written out to disk and then read in by the subsequent radiator run with the longitudinal phases (measured relative to a plane wave) multiplied by the harmonic upshift number, in this case 6 (=240 nm/40 nm). However, the temporal spacing and resolution in the radiator runs remains the same (*i.e.*, 0.8 fs); in other words, the macroparticles are not reorganized into independent 40-nm slices. Rather, physics quantities such as current and microbunching fraction at 40-nm wavelength are effectively averaged over a 240-nm interval. So long as the normalized output spectral bandpass is small compared with $\lambda_p / c\Delta t = 40 \text{ nm} / 240 \text{ nm} = 1/6$, this temporal resolution is more than adequate.

3.4.2 Time-dependent results for the “M2” and “M6” Medium Pulse Distributions

In this section we present results for two particular S2E golden file distributions: “M2” with a current of ~1-kA and “M6” with a ~700-A current (somewhat below the 800-A design current for the steady state calculations). Figures 3.4-1 and 3.4-2 show the time-resolved current and energy profiles of these distributions at entrance to the undulator. One sees that the M2 pulse has a relatively flat current distribution between current spikes at the head and tail but a fairly large quadratic energy chirp in time. The flat current region has an incoherent energy spread of ~100 keV, a value that presently we believe may be artificially low by a factor of 1.5-2.0 because longitudinal space charge instability physics were not included in the corresponding ELEGANT runs. The M6 distribution is much flatter in energy; its core region incoherent energy spread is also ~100 keV.

In Figs. 3.4-3 and 3.4-4 we plot the predicted output power versus time at the end of six radiator segments ($z=20.28 \text{ m}$) for the two distributions for a situation where the input seed covers the entire electron pulse. The final radiation pulse energies, [1.5 mJ – GINGER; 2.0 mJ - GENESIS] and 1.7 mJ for M2 and M6 respectively, are equivalent to slightly more than 3×10^{14} photons/pulse. For GINGER with an axisymmetric field solver, the average output powers of ~1.8 GW are smaller than those calculated in the time-steady case (~2.4 GW). This is to be expected because S2E distributions will have various effects (*e.g.* transverse centroid offsets) ignored by the time-steady calculations. The GENESIS calculation shows a much greater power in the tail region for M2, the distribution with a fairly large quadratic energy chirp.

It is unclear if this occurs because of a correlation between offset and higher energy in the tail region that leads to more particles being in resonance in a fully 3D calculation or some other effect.

The output spectra (Fig. 3.4-5 in which we show only a narrow bandpass region near the central wavelength) are not dissimilar but the M6 result has a much narrower, well-defined line that is possible because of the nearly flat electron beam energy distribution. The M2 output is spread out over a bandwidth nearly four times larger, an effect directly attributable to the interaction between the dispersion section ($R_{56} = 32\text{-}\mu\text{m}$) and the quadratic energy chirp. The time-integrated far field RMS emission angles are 64 and 52 microradians for the M2 and M6 cases, respectively, suggesting that the effective mode size and content are quite similar..

We also performed a set of calculations with a short duration input laser seed (100-fs FWHM Gaussian profile) as might occur in a pump-probe type experiment. In order to roughly approximate the effects of e-beam/laser-seed timing jitter, a series of six runs with the relative e-beam-laser time sequentially increasing 100 fs from run to run were done for the M2 distribution. Particularly noteworthy (see Fig. 3.4-4) is the formation of an intensity spike at the leading edge (*i.e.*, head) of the output pulses up to twice the average level of the pulse main body. The enhancement appears to be a strong function of the e-beam current inasmuch as an 1200-Amp region near the tail produces a spike exceeding 9 GW whereas a 800-Amp region near the head shows a 4 GW spike level or less. Presumably, the spike formation is related to a super-radiant phenomenon of the sort seen by Giannessi in SPARK-X simulations (private communication). The change in output pulse energy as the temporal center of the seed laser pulse moves from the e-beam head to tail suggests that a 350-fs temporal jitter *alone* could introduce a 15-30% output jitter. This would be in addition to any jitter associated with shot-to-shot fluctuations in the mean electron beam energy.

3.4.3 Wakefield Calculations for the M2 & M6 Medium Pulse Distributions

As discussed in Sec. 2.6, there will be wakefields arising from the interaction of electron beam electromagnetic fields with the vacuum chamber walls and geometric breaks occurring from any cross-section interruptions. We have done some very preliminary calculations for the M2 and M6 current distributions using a wakefield code based on a numerical physics package developed by H.-D. Nuhn at SLAC. This code currently includes wakes from vacuum chamber resistivity, surface roughness, geometric breaks, and a “synchronous term”. For the calculation here, we presumed a circular Al vacuum pipe of 10.0-mm inner diameter, a surface roughness of 100-nm amplitude with a longitudinal period of 25 microns. The geometric wake was calculated presuming a 10-cm break occurring every 3.4 meters. An “AC” conductivity model was used for the resistive wake.

With these choices, the resultant longitudinal wakefield of both the M2 and M6 distributions does not appear to seriously threaten FEL-1 output. Apart from large spikes at the head and the tail of ~ 50 kV/m amplitude, the M2 and M6 wakes (see Figs. 3.4-6 and 3.4-7) have temporal variations of only 5 KV/m or less. Over a 24-m vacuum pipe length, a fluctuation of ± 96 kV (or less than 0.01%) is not at all worrisome. The spike region wakes are large enough to cause a reduction in output --- however, there are other phenomena (*e.g.*, higher emittances) which will also suppress emission. The approximately constant wake term over the interval (-0.3 ps,+0.3 ps) of ~ 8 kV/m can be compensated by a very slight undulator taper. As the FERMI undulator and vacuum chamber design becomes more finalized, these wake calculations must be redone to include more realistic roughness numbers and perhaps a non-circular geometry.

3.5 Time-Dependent Jitter calculations

To examine the affects of injector and accelerator jitter upon the shot-to-shot, *fully time-resolved* properties of the output FEL-1 radiation, 100 individual files of ~1M macroparticles were propagated starting from the injector (GPT code) through the linac (Elegant code) for the M2 medium bunch case (including quadratic energy chirp). Each file included the effects of random jitter in the individual injector and accelerator cell voltages and timing. The jitter followed Gaussian distributions with variances set by the budget allowances allocated by the S2E group to satisfy the requirement on the output electron beam that $\Delta E/E \leq 0.1\%$ and $\Delta I/I < 10\%$. After analyzing the mean arrival time of the electron bunch in each file, we discovered that the RMS temporal jitter was 337 fs (see Fig. 3.5-1), a factor two larger than the expected value predicted by **LiTrack** code for the nominal gun and linac parameter fluctuations. This unexpected difference may be due an anomaly in the Elegant program and there are ongoing studies to examine its nature.

Table 3-5 : Time-Averaged Envelope Quantities from 100 Medium Bunch Jitter Files

Quantity	Mean Value	Std. Dev.
Gamma	2306	0.08%
Current (A)	1010	5.6%
Incoherent energy spread	0.208	7.7%
Normalized emittance (m-rad)	1.59e-6	11.3%

Although this amount of temporal jitter precludes highly meaningful short pulse studies, these particle files do contain an elaborate wealth of information, including both the time-resolved fluctuations of individual quantities such as current or emittance AND the full correlations of these quantities with one another. Consequently, we decided to perform “whole-bunch” simulations in which the radiation pulse is constant in time and covers the *entire* electron beam, thus making the timing jitter irrelevant. From these simulations, one may determine the predicted RMS jitter in output photon number and benchmark this value against time-independent simulations of Sec. 3.2.3 where the various electron beam quantities independently vary following prescribed Gaussian distributions.

Figures 3.5-2 through 3.5-5 plot the mean energy, current, energy spread, emittance versus time for all 100 files. For each file, these quantities were determined from the output macroparticle positions in the ELEGANT files using the SDDS program **elegant2genesis** in which we adopted a temporal resolution of 16 fs. This program also recenters the time coordinate at the mean arrival time of all electrons in a given file. We then statistically analyzed these quantities in a time window extending from -200 fs to +200 fs encompassing the pseudo-flat current region of which the great majority of all FEL emission will occur. Table 3-5 reports the mean value and the standard deviation of each quantity inside this central time window

GINGER time-dependent simulations for the FEL-1 lattice at 40-nm output wavelength were performed over a time window of 1.0 ps with 4-fs resolution. For each jitter file, a simulations was done using artificial macroparticles created from the time-dependent envelope quantities previously determined by the **elegant2genesis** code. The results for FEL output power are shown in Fig. 3.5-6. Over a central

time window of $[-350 \text{ fs}, +350 \text{ fs}]$, statistical analysis of these data gives an average power of $\sim 2.5 \text{ GW}$ and a normalized standard deviation of $\sim 20\%$ (see Figs. 3.5-7 and 3.5-8), both of which are in a good agreement with time independent results obtained with the nominal values.

Similar whole-bunch seeding simulations for FEL-1 were performed with GENESIS using directly the (10^6) elegant macro-particles. A simulation time window of about 600 fs was chosen which encompasses the central part of the bunch. Statistical analysis of the results gives an average energy of about 1.24 mJ ($=2.1 \text{ GW}$) with a standard deviation of about 24% (Fig. 3.5-9). This result is in reasonable agreement with Ginger's prediction and with time-independent simulations. Note that GENESIS has a full 3D (x-y-z) field solver using the ELEGANT macroparticles while the GINGER field-solver is 2D (r-z) and the created macroparticles followed a 5D Gaussian distribution (x, y, x', y', γ) . These differences could strongly affect the output power calculations if there are time-dependent offset or tilts within the electron bunch and/or unexpectedly strong correlations between quantities such as x and y' .

In order both estimate the expected jitter in output power for short pulse seeding *in the absence of temporal jitter* and to examine more local dependences of output power with various electron beam parameters, another series of GINGER runs over a time window of 400 fs with 31-fs resolution were done with the time-centered ELEGANT files but now using a seeding pulse with a Gaussian temporal width of 100 fs. The average power and its normalized standard deviation were determined to be 1.5 GW and 28% respectively over a time window $[-200 \text{ fs}, +200 \text{ fs}]$ (Figs. 3.5-10 and 3.5-11). The increased fluctuation level when compared with the flat seeding results above may be due both because the lower average output power in jitter files with lower average current will be partially compensated for in the flat seeding case by the typically longer pulse durations corresponding to these lower current files. In addition, the larger time-window of the flat-seeded case will tend to average out fluctuations in electron beam parameters with time scales $\sim 75\text{-}300 \text{ fs}$ far more strongly than will be true for the 100-fs seed case.

However, the more temporally localized nature of the 100-fs Gaussian profile seed allows us to identify the individual electron beam parameters to which the FEL performance is most sensitive by plotting for each jitter file the average output power versus the average value of a given electron bunch parameter. Moreover, we can confirm trends in previous time-independent results.

After calculating the averages of the electron beam parameters over the time window $[-200 \text{ fs}, +200 \text{ fs}]$, one can see from Fig. 3.5-12 that the time-independent results for output power dependence upon electron energy and current are in fact confirmed. However, the time-dependent results do show a much stronger dependence upon emittance and energy spread. This is probably due to some correlation between electron parameters. In particular, the dependence of output power from the energy spread is correlated to the approximately linear dependence of incoherent energy spread upon the electron beam current. In order to understand the large effect of the emittance on the output power that was not present in time-independent simulations, we examined the correlation between current and energy spread, and Twiss parameters and emittance in the 100 files. As shown in Fig. 3.5-13, there is a clear correlation in these quantities. Consequently, while relatively simple one parameter scans around a central design point are of course useful to illuminate key dependences, one does need to do full start-to-end simulations that include the full correlations of different parameters in realistic electron beams in order to obtain accurate estimates of the full sensitivity of output power.

3.6 Diagnostic needs for FEL-1

From the point of view of FEL physics, there are a number of obvious measurements that could help guide the operational staff in maximizing output performance from FEL-1. Within the different undulator segments and other lattice elements, electron beam position monitors are needed to ensure that the beam orbit stays as close as possible to the magnetic axis. Following the chromatic dispersion element, the electron bunch will have both strong energy modulation and microbunching; the latter can be measured with coherent optical transition radiation (COTR) from an insertable foil. In the radiator, the build up of the coherent harmonic signal can be determined section by section by purposely mistuning the magnetic strength of the downstream sections so as to eliminate any additional emission. COTR in the break sections can also be used to determine the z -dependent evolution of the microbunching. Note that the COTR signal should also be rich in harmonics of the initial seed laser. Ideally for comparison with simulation predictions, measurements of the microbunching and FEL radiation should be done following each radiator undulator section. Finally, it will also be useful to measure the final energy spread of the e-beam upon exit from the final undulator section as this too should be commensurate with the FEL emission from the radiator.

4 FEL-2 Design & Performance Calculations

4.1 Basic Design for the Fresh-Bunch Approach

As mentioned in the Introduction, we presently believe that most applications for FEL-2 output will be in the spectral domain where there is a premium for obtaining narrow bandwidth. Consequently, our work on FEL-2 has concentrated utilizing a relatively long (~ 1.5 ps) electron beam pulse with a moderate (~ 500 A) current. For the fresh-bunch approach, such a long pulse is nearly essential given the practicalities of temporal jitter and the accuracies of a temporal delay section. For the whole-bunch approach, one conceivably could use shorter pulses of the sort suggested for FEL-1. However, this raises the question of how small the actual incoherent energy spread will be for a pulse with ~ 800 - 1200 Amp current; too large an input energy spread leads to larger energy spreads in all downstream stages and consequently reduced gain and longer required radiator lengths in the final stage to reach the same output power.

Our core design for the first stage of fresh-bunch (FB) FEL-2 is extremely similar to that described earlier for FEL-1. Specifically, relatively strong input seed power (~ 100 MW) energy-modulates a 500-A electron beam in a long period (~ 16 -cm wavelength) undulator followed by a dispersion section with an $R_{56} \sim 25$ μm . This produces strong bunching at the fundamental ($b \geq 0.5$) and also at the 2nd-6th harmonic in the following radiator. Where our design for FEL-2 FB layout (see Fig. 4.1-1) begins to differ significantly is that the first stage radiator is relatively short (*e.g.*, 2-3 segments) and only brings the radiation to a sufficient level (~ 200 MW) to provide adequate coherent energy modulation in the following undulator. This contrasts to the FEL-1 radiator that is long enough (*e.g.*, 5-6 segments) to reach FEL power saturation (~ 1 - 4 GW). The basic parameters for the second stage modulator are the same as for the first stage radiator (*e.g.*, 65-mm period; 2.34-m segment length) that, in principle at least, should minimize costs. In general, we believe it best to keep the first stage as short as possible, both for cost reasons and to minimize SASE growth which can increase the incoherent energy spread of the “fresh” portion of the e-beam to be used in the second stage modulator and radiator. The FB delay section is described in Sec. 2.5 and is provisionally presumed to have a 1.8-m length (a value is necessary in order that the numerical simulations properly model radiation diffraction effects across the delay section).

The second stage (and final) radiator has a somewhat shorter period (*i.e.*, 50 mm) and is subdivided into 2-m long active undulator segments separated by 1-m breaks. These breaks contain a quadrupole singlet for focusing, a phase shifter, dipole correctors, and diagnostics. The length of the final radiator is somewhat arbitrary; in general we have presumed for the FB approach one would want sufficient length for power saturation, ~ 6 segments at 10-nm wavelength. However, one could certainly increase the output power by adding more radiator segments whose magnetic strength would then be tapered. The FEL rho parameter for 10-nm output wavelength is $\sim 1.2 \times 10^{-3}$ and the 3D exponential gain length for 100 keV incoherent energy spread is 2.4 m. Ignoring the effects of the undulator module breaks, the predicted saturated power according to the M. Xie formalism is 610 MW.

4.2 Fresh-Bunch Time-Independent Simulation Results

Following similar strategies for optimizing the different components associated with the FB approach as were employed in the FEL-1 time-independent studies in Sec. 3.2, we did a series of simulations to

optimize the output power of the second stage. Given the presumption that most of the first FEL-2 stage would more or less be identical to that of FEL-1 (an assumption well warranted by cost considerations but perhaps not necessarily optimal for final performance), the most sensitive things to vary are the number of radiator sections in the first stage and the strength of the dispersive section following the second stage modulator. We found that a good design point was a three-section first radiator which would put out ~ 250 MW of power at 40-nm wavelength. Then at the end of the second modulator, the peak-to-peak energy modulation is ~ 1.8 MeV (as compared with 2.5 MeV modulation in the first stage) which works best with an $R_{56\text{-RAD2}}$ of 7.0 μm . The final radiator is 6 sections long and produces (see Fig. 4.2-1) a peak power of 0.35 GW at 10-nm for 500-A beam current and 1 GW for 800-A current. As can be seen in Fig. 4.2-2, the peak coherent microbunching fraction exceeds 0.5 at the longer wavelengths and reaches ~ 0.3 or greater at 10-nm. The output mode quality is generally quite good with $M^2 \sim 1.2$ and a far field opening angle of 15 microradians. These particular runs presumed an initial incoherent energy spread of 200 keV.

Similar runs at 500-A current (but with 100-keV incoherent energy spread at the undulator entrance) were made with the GENESIS code. As shown in Fig. 4.2-3, one obtains a maximum power at 10-nm of 0.63 GW, and peak bunching of 0.44; the initial bunching in the final radiator is 0.17. At 20-nm output wavelength, the maximum power is 1.4 GW and the peak bunching is 0.53; the initial bunching in the final radiator is 0.19.

We also did a small number of time-independent, single parameter variation scans to determine the sensitivity of the final output. Not surprisingly given the much larger total number of undulator periods in FEL-2 as compared with FEL-1 plus the additional sensitivity connected with having a second stage of modulation, there is a much greater dependence upon electron beam parameters, especially energy. In Fig. 4.2-4, we plot output power versus initial beam energy. One sees that the FWHM is not much larger than 0.2%; this places rather severe constraints upon accelerator jitter if shot-to-shot reproducibility is needed for a particular FEL-2 experiment. There also is a fair amount of sensitivity upon beam current (Fig. 4.2-4). In previous LUX studies at LBNL, it was found that in a 4-stage cascade one could tune the multiple dispersion sections so that to first order in the current the output power was at a local maximum; however, that may be more difficult in a two stage harmonic cascade. The emittance variation results also show generally more sensitivity as was found in FEL-1. This is due both because of the reduced gain at the shortest wavelengths and because of the greater sensitivity in general of multistage cascades to any parameter variation.

4.3 Fresh-bunch Offset and Tilt Sensitivity Studies Using GENESIS

4.3.1 Introduction

FEL-2 is typically more sensitive to variations in input parameters than is FEL-1. This is largely due to the shorter wavelengths targeted for FEL-2, which leads to similar displacement errors translating into larger longitudinal phase errors, and there is the additional complexity of having two rather than one stage of harmonic conversion. Furthermore, at shorter wavelengths radiation diffraction is less important, which renders FEL-2 more sensitive to deviations in the electron orbit and to misalignments. As in the FEL-1 studies, offsets and tilts were modeled in GENESIS because a fully three-dimensional code is necessary to capture fully these effects. While the simulations include initial offsets for the electron beam, the laser seed and undulators were assumed to lie along a common axis. Time-independent results for FEL-2 using a 240-nm wavelength laser seed are presented at wavelengths of 20 nm and 10 nm, using stage 1 wavelengths of 60 nm and 40 nm respectively. In these examples, there is no undulator tapering.

Critical parameters plotted include output power and phase, the maximum power which can be contained within a pure Gaussian mode, the transverse mode quality (given in terms of M^2), the divergence angle, the location of the virtual waist and the radius of the waist. The diagnostics of the output radiation are performed in the same way as the FEL-1 sensitivity studies. Note that all green curves in the following graphs are plotted against the right axis. The most prominent effect of beam offsets is a large drop in output power if the electron beam and laser seed fail to overlap each other in the first undulator, although in addition the FEL output develops offsets comparable to those of the electron beam

4.3.2 10- and 20-nm Fresh-Bunch Results with Offsets and Tilts

Figures 4.3-1 and 4.3-2 show the resulting analysis of the mode properties for the fresh-bunch 20-nm case as a function of a pure initial offset or tilt respectively in the electron beam at the start of the modulating undulator. The actual displacement of the output radiation is shown in Figs. 4.3-3 and 4.3-4. The effect of beam tilts is fairly significant, although for a tilt of 30 microradians the output power is still more than 50% of the ideal result. The shorter wavelength and the additional stage of harmonic generation lead to more stringent requirements on beam alignment. Beam offsets from the laser seed are even more of a concern, as the power essentially drops to zero for a 200-micron offset, even though the rms spot size of the laser seed is 210 micron. Offsets of 50 microns do not have a severe effect on the output power, although even this offset results in a 0.5 m shift in the apparent waist location and a 5% reduction of the radius of the output radiation field. It is clear that the delay section in the fresh-bunch configuration, where the information from the seed is separated from the electron beam and is contained entirely in the 60-nm radiation pulse, does not present any particular difficulties so long as that section itself is aligned well.

Figures 4.3-5 through 4.3-8 indicate the analogous results for a fresh-bunch configuration with an output radiation wavelength of 10 nm. The sensitivity is more serious in all criteria. In both fresh-bunch and whole-bunch configurations, the power drops to essentially zero for offsets of 150 micron, and a 50 micron initial offset leads to a 30% reduction in power. For the case of an electron beam entering with a non-zero tilt, the power is reduced by more than half for tilts exceeding 22 microradians.

4.4 Fresh-Bunch Time-dependent Simulation Results

A number of full start-to-end simulations were done for the fresh-bunch approach to FEL-2. We again concentrated upon 10-nm output cases as these were likely to be the most sensitive to imperfect electron beam parameters; the first stage radiator output was at 40-nm so the second stage had a 4:1 up-conversion ratio. Here we present results for both the L2 and L4 “golden file” long-pulse distributions that have about 1 nC of charge and differ mainly by the longitudinal variation of the electron energy across the bunch. The “L2” bunch has a large quadratic energy chirp, with a total energy variation of 2 MeV across the bunch (ignoring the large current spikes at either end of the bunch). The “L4” bunch has a predominantly cubic energy chirp, which leaves the core of the bunch essentially flat in energy, but the head and tail of the bunch both have an energy chirp with decreasing energy towards the head of the bunch. There is only a small current enhancement in the tail of the L4 distribution. The beam current and energy profiles are compared in Figs. 4.4-1 and 4.4-2.

Figures 4.4-3 through 4.4-6 present plots of 10-nm output power and far field on-axis eikonal phase versus time and power spectra in the central spectral region for the two cases. For the GINGER runs, the input laser seed was a 470-fs (FWHM) Gaussian pulse with 120 MW peak power centered at $t=+250$ fs for L2 and $t=+750$ fs for L4. The fresh-bunch delay section shifted the field pulse 650 fs toward the head for L2 and 1250 fs for L4 (note: these positions and delays were not at all optimized and it would not be

surprising if one could obtain at 1.5-2X increase in power and spectral brightness). The GENESIS runs used a somewhat longer pulse (600-fs FWHM) at a slightly lower power (100 MW), Both codes used first stage dispersive sections with $R_{56} \sim 20$ microns.

The L4 distribution performs significantly better, both in terms of power and spectral brightness. The total radiation energies in the L2 and L4 pulses were 60 and 200 microjoules respectively. GINGER results showed equivalent spectral widths (FWHM) were 38 and 23 meV, about 8 and 5X larger than the transform limit for this particular seed pulse; GENESIS shows a narrow spike for L4 with a FWHM of roughly 10 meV. The far field opening angle is 19.5 microradians for L2 and 17.5 microradians for L4 and in both cases the near field transverse profile appears to be extremely clean. Other GINGER diagnostic quantities for (L2, L4) include: $M^2 \sim (1.25, 1.20)$, (96.0 96.5%)% of radiation contained in a lowest order Gaussian mode whose apparent emission point was (11.5, 12.0) m upstream of undulator exit with a waist of $\sim (125, 140)$ microns. Equivalent GENESIS results for the (L2, L4) distribution showed (89%, 93%) of the power in a Gaussian mode with a waist (8.1, 6.2) m from the FEL end and a waist size of (150, 170) microns.

Although more output power can be obtained by increasing the seed duration, the spectral brightness will generally not increase for the L2 distribution since the quadratic chirp leads to the spectral width increasing nearly linearly with pulse width. In theory one can compensate for the chirp by placing a comparable reversed chirp on the input seed but the output wavelength will be sensitive to timing jitter between the seed laser and electron beam. If the RMS timing jitter could be reduced to 100 fs or less, use of 750-fs FWHM or longer seed pulses is certainly conceivable and for the L4 pulse could significantly increase both the total number of photons and to some extent the spectral brightness

4.5 Basic Design for the Whole-Bunch Approach

As mentioned in the Introduction, the whole-bunch approach to FEL-2 differs in a number of ways. First, the entire duration of the electron beam will be energy modulated by interaction with the external seed laser whose strength is typically much weaker than that for the fresh-bunch approach --- we found that 45 MW appeared to work well for 10-nm output and values as low as 30 MW for 20-nm output wavelength. One needs to minimize this value in order to keep the induced energy spread small and thus not increase the exponential gain length too greatly for the second stage radiator. The first stage dispersive section has an $R_{56} \sim 30$ μm ; before it was realized that quadratic chirp effects could greatly degrade the output bandwidth, smaller values of input seed laser and larger dispersions (sometimes exceeding 100 μm) were considered. Second, the first stage radiator can be quite short (*e.g.*, 2 or 3 sections) and is followed by a weak dispersive chicane ($R_{56} \sim 4$ μm) and then a long (up to 10 sections), second stage radiator. As was true for the FEL-1 radiator, the average beta function was set to 10 m.

4.6 Whole-Bunch Time-independent Simulation Results

Time-independent GENESIS code results for the whole-bunch approach are shown in Fig. 4.2-3 together with those corresponding to fresh-bunch. The nominal electron beam parameters were 1.2 GeV energy, 500 A current, 150 keV incoherent energy spread, and 1.5 mm-mrad normalized emittance; these numbers correspond to the “long bunch” pulse duration. Some additional studies in following sections were done with 800-Amp current and “medium bunch” pulse duration. One should also note that over the course of the TOS study the adopted value for initial incoherent energy spread varied over a range of 100- to 150

keV and as of June 2006 it is still unclear from the studies of longitudinal instability growth in the linac as to what an appropriate “safely conservative” value is.

For the 10-nm output wavelength, the maximum power over a 10-section final radiator is 0.31 GW, and the peak bunching is 0.28; the initial bunching in the final radiator is 0.04. For the 20-nm whole-bunch case, the maximum power over an 8-section final radiator is 0.93 GW, and the peak bunching is 0.38; the initial bunching in the final radiator is 0.05. We note that for both the fresh and whole-bunch approach, the microbunching value peaks two undulator sections before the output power reaches its peak; this suggests that appropriate tapering might improve the output power.

4.7 Whole-Bunch Offset and Tilt Sensitivity Studies Using GENESIS

In order to examine the output effects of initial tilts and offsets on the electron beam (whose parameters were the same as in the previous section), we performed FEL-2 whole-bunch studies equivalent to those done for the fresh-bunch approach. At 20-nm output wavelength, the whole- and fresh-bunch results are very similar as shown in Figs. 4.7-1 through 4.7-4. The main difference is that the whole-bunch case is less sensitive to initial offsets (but still drops to zero power with a 200 micron offset). In both cases, an initial offset yields a roughly identical offset in the output radiation field, and a tilt that is roughly given by the ratio of the original offset to 20 m. An initial tilt yields an output tilt roughly half in magnitude. The offsets which results from an initial tilt are very different in the two configurations: in the fresh-bunch case, the final offset is roughly 1 m times the initial tilt and peaks at 35 micron; in the whole-bunch case, the final offset is much smaller for small initial tilts, but shoots up to very high values at initial tilts of 30 microradians or greater.

At 10-nm output wavelength, the whole-bunch case is much more sensitive to initial tilts, with the output power reduced by half for a 12-microradian tilt (see Figs. 4.7-5 through 4.7-8). This increased sensitivity is probably related to the reduced performance of the whole-bunch configuration at 10 nm for the relatively low current of 500 A. The offset of the output is of the order of the initial offset of the electron beam. An initial tilt leads to an offset of the output given by roughly 7 m times the initial tilt, while the final tilt is roughly half the initial tilt.

4.8 Whole-Bunch Time-Independent Optimization and Jitter

Because the whole-bunch approach uses the same electrons in both FEL stages, electrons enter the final radiator with an increased incoherent energy spread due to FEL interaction in the first stage. Consequently, both the initial (*i.e.*, at entrance to the first undulator) and added energy spread plays a crucial role and can strongly degrade FEL performance. To optimize the output performance of FEL-2, one needs to find a compromise where the middle radiator produces sufficiently large microbunching at the desired harmonic without overly enlarging the energy spread. As a possible optimization strategy we considered a simple setup that neglects any concern for maintaining easy hardware and layout compatibility with the fresh-bunch scheme. The number of undulator segments per radiator is then allowed to vary and the quadrupole strength was set to produce an average beta function of about 10 m in the final radiator.

4.8.1 Optimization for Maximum Output Power

The first optimization study aimed for peak output power and was performed for the long bunch case by

simultaneously varying the power of the input seed laser power and the R_{56} values of the two dispersive sections. The base electron parameter values are given in the second column of Table 4-1. The results appear relatively insensitive to using 2 or 3 sections on the middle radiator; however, 10 sections in the final radiator are always needed. At peak 10-nm power optimization, we found 0.5-GW output by choosing a seed power of 95 MW and $R_{56\text{-RAD1}}=81\mu\text{m}$ and $R_{56\text{-RAD2}}=9.6\mu\text{m}$. However, note these parameters strongly depend on the value of the initial (incoherent) energy spread of the electron bunch.

Table 4-1: Input Parameter Mean and RMS Jitter Values

Parameter	Nominal value	Normalized Parameter Fluctuation Level (%)
Norm. Emittance	1.5 mm mrad	10.0
Peak current	500 A	8.0
Mean energy	1.2 GeV	0.1
Energy spread	150 keV	10.0
Seed power	95 MW	5.0

4.8.2 Optimization for Reduced Output Jitter

An important consideration when optimizing parameters is sensitivity to electron beam parameter fluctuations. We performed a jitter analysis on the setup described in the previous paragraph by employing time-independent simulations using “long bunch” mean parameter values and adopting normalized shot-to-shot fluctuations at a level tabulated in Table 4-1. We found (see Table 4-2 below) very high sensitivity to many of the parameter fluctuations. This sensitivity occurs due to the crucial role that energy spread (and the balance between coherent and incoherent energy spread) plays in the gain of the whole-bunch approach. In the WB approach one must modulate the electrons as much as possible in order to provide good microbunching, but not too much energy modulation in order to enter the second radiator with reasonable energy spread. In the first radiator the gain, and consequently the induced energy spread, will depend on electron parameters like $(\delta\gamma, \epsilon, D)$; these will then affect the final output power.

Table 4-2: Output Power Jitter Sensitivity for Nominal Parameters

Parameter	Average output power (MW)	Normalized fluctuations σ (%)
Emittance	440	18.3
Peak current	443	21.0
Mean energy	334	45.1
Energy spread	372	26.0
Seed power	434	4.4

Table 4-2: Sensitivity to the considered parameter fluctuations of the setup that maximizes the nominal output power. Multi-parameter jitter sensitivity has been estimated by a statistical analysis of the simulation results. The standard deviation of power is normalized to its mean value.

Table 4-3: Output Power Sensitivity to Modified Parameters

Parameter	Average output power (MW)	Normalized fluctuation σ (%)
Emittance	355	7.6
Peak current	353	13.5
Mean energy	258	43.6
Energy spread	344	4.9
Seed power	352	12.9

A possible way of reducing the output power fluctuation level is to find a different configuration with hopefully less output sensitivity to $\delta\gamma$. By studying different combinations of input seed power and R_{56} strengths of the two dispersive sections, we found that a significant reduction of the sensitivity to $\delta\gamma$ can be obtained by using a stronger seeding laser, a weaker strength of the first dispersive section and a slightly weaker second dispersive section. In particular we considered a seeding power of 125 MW, $R_{56\text{-RAD1}}=65\mu\text{m}$ and $R_{56\text{-RAD2}}=9.0\mu\text{m}$; jitter results are reported in Table 4-3. Figures 4.8-1 and 4.8-2 show the predicted growth of power with z for the long- and medium-bunch cases respectively with this particular optimization of the laser power and dispersive sections.

4.8.3 Optimization for Reduced Bandwidth

An additional important consideration is the output spectral bandwidth and its sensitivity to the interaction of a dispersive section with a temporal quadratic energy chirp in the electron bunch. In order to minimize bandwidth increase due to this effect, the FERMI FEL design group has decided to limit $R_{56\text{-RAD1}} \leq 30\mu\text{m}$ and $R_{56\text{-RAD2}} \sim 5\mu\text{m}$. With these constraints, we performed a new optimization for electron beam parameters typical of medium and long bunch cases developed by the S2E group.

For the long bunch case that has a nominal current of 500 A and an initial incoherent energy spread of 100 keV, the optimized setup results in a seeding power of 45 MW, $R_{56\text{-RAD1}}=30\mu\text{m}$, 2 radiator sections for the first stage, $R_{56\text{-RAD2}}=4.0\mu\text{m}$, and, finally, 10 second stage radiator sections. Table 4-4 gives settings for input seed power and the dispersion sections before each radiator. The nominal output power with this setup is about 0.3 GW.

Table 4-4: Bandwidth-optimized Long Bunch Setup for WB

Wavelength (nm)	Seed power (W)	$R_{56}(\mu\text{m})$	$R_{56}(\mu\text{m})$
($\delta\gamma=100$ keV) 240→40→10	45.0	30.0	4.0
($\delta\gamma=100$ keV) 240→60→20	20.5	30.0	5.0
($\delta\gamma=100$ keV) 270→90→30	10.0	30.0	4.0

Table 4-5: Bandwidth-optimized Medium Bunch Setup for WB

Wavelength (nm)	Seed power (W)	$R_{56}(\mu\text{m})$	$R_{56}(\mu\text{m})$
($\delta\gamma=125$ keV) 240→40→10	40.0	30.0	4.2
($\delta\gamma=125$ keV) 240→60→20	12.2	30.0	5.0
($\delta\gamma=125$ keV) 270→90→30	14.0	11.8	2.5

For the medium bunch case in which the nominal current is 800 A and an initial incoherent energy spread of 125 keV, the optimized setup at 10-nm wavelength has a seeding power of 40 MW, $R_{56\text{-RAD1}}=30\mu\text{m}$, 2

radiator sections for the first stage, $R_{56\text{-RAD2}}=4.2\mu\text{m}$. Here a maximal output power of 1.1GW is reached after 9 second-stage radiator sections. Table 4-5 gives input laser seed and R_{56} settings for 10-, 20- and 30-nm output wavelengths. Figures 4.8-3 and 4.8-4 plot the radiation power, and electron beam microbunching and energy spread versus z in the second radiator.

4.8.4 Whole-Bunch Time-Independent Jitter Sensitivity Studies

We have done some parameter jitter studies for the whole-bunch configuration described in the previous section concerning reduced bandwidth. For a case corresponding to “long bunch”, the e-beam current was 500 A, incoherent energy spread 100 keV, input laser power 45 MW, and $R_{56\text{-RAD1}}=34\mu\text{m}$ and $R_{56\text{-RAD2}}=3.5\mu\text{m}$. The adopted RMS fluctuation levels for the different e-beam parameters were the same as listed in Table 3-4 for FEL-1. For simultaneous variation of all parameters, the resultant output power level dropped to 135 MW with a standardized fluctuation level of 81%. Figures 4.8-5 and 4.8-6 show the shot-to-shot scatter plotted against different variables. One sees that variations in mean beam energy continue to dominate but there are also strong variations of output power with current and emittance. This effect is to be expected because the net exponential gain in the final WB FEL-2 radiator is sensitive to both current and emittance.

The medium bunch jitter scans used an electron beam with 800-A current and 100-keV energy spread. The bandwidth-optimized setup had an input laser of 40 MW with $R_{56\text{-RAD1}}=30\mu\text{m}$ and $R_{56\text{-RAD2}}=3.1\mu\text{m}$ (note that this last value is slightly smaller than that used in the previous section). The jitter scans (see Figs. 4.8-7 and 4.88) showed an average output power of 655 MW with a standardized fluctuation level of 49%. The beam energy variations again dominate the output fluctuations but, in comparison to the long pulse case, there is less sensitivity to emittance and current, presumably because the FEL ρ parameter is larger.

4.9 Whole-Bunch Time-Dependent Simulation Results

4.9.1 Introduction

Relatively numerous time-dependent simulations have been performed using electron beam distributions produced by the S2E group, which mostly fall into two categories, a medium-duration (0.7 ps) with 0.8 nC of charge, or a long-duration bunch (~ 1.5 ps) with 1 nC of charge. Either case is suitable for the whole-bunch FEL-2 configuration which comprises: the initial modulator is 3.04 m long, then a 1.04-m break which includes a chicane having $R_{56} \sim 30$ microns, followed by two 2.34-m undulator sections separated by a 1.04-m break; these are followed by a 1-m break with a weaker chicane ($R_{56} \sim 3\text{-}4\mu\text{m}$) that is tuned to the range of 3 to 4 microns, and a final radiator consisting of multiple sections, each 2.4 m in length and separated by 1-m breaks. Quadrupole focusing keeps the average beta function to around 10 m. The seed laser has a peak power of 40 MW in a 1-ps pulse (FWHM), and is focused to a 200-micron spot size. In the modulating undulator, the laser seed overlaps the electron bunch roughly in the center of the bunch. The **GENESIS** simulations presented here are tuned to produce 10-nm output radiation from a 240-nm seed laser.

4.9.2 Long Pulse “L2” and “L4” Simulation Results

Simulations of long-duration bunches include the “L2” and “L4” distributions. Briefly, the L2 distribution bunch has a large quadratic energy chirp, with a total energy variation of 2 MeV across the bunch while the “L4” bunch has a predominantly cubic energy chirp, which leaves the core of the bunch essentially flat in energy, but the head and tail of the bunch both have an energy chirp with decreasing energy towards the head of the bunch. See Sec. 4.4 and Figs. 4.4-1 and 4.4-2 for more detailed information. In the long-duration bunch examples, the peak current is in the range 500 – 600 A, the slice (i.e., incoherent)

energy spread is about 70 keV, and FEL power saturation is not reached until after 10 sections of final undulator. In the medium duration bunch examples, the peak current ranges between 750 and 1000 A, and saturation occurs after 8 undulator sections.

The predicted output pulses are slightly different in the two cases. The “L4” distribution yields more peak power and a longer pulse, with a peak power of roughly 450 MW and a FWHM of 470 fs. The total energy in the pulse is 0.25 mJ, of which 93% is in a Gaussian mode with $Z_R=7.7$ m and waist location -7.2 m from the end of the FEL. The bunching has a double-peaked distribution, which implies that the center of the pulse has reached saturation, but the edges may still be growing. The photon energy spectrum has a peak corresponding to 8.4×10^{11} photons per meV, and a FWHM of 8 meV. There is a secondary peak at higher energy, but it has less than half the magnitude of the primary peak. This is sufficient to raise the spectral rms width to 29 meV, however.

The FEL output for the “L2” distribution has a peak power of 350 MW and a FWHM of 450 fs, and the power profile has more fluctuations. The total energy in the pulse is 0.13 mJ, of which 93% is in a Gaussian mode with $Z_R=7.5$ m and waist location -7.8 m from the end of the FEL. The bunching has a flat-top profile with a peak value of 0.3. The photon energy spectrum has three prominent spikes with a peak corresponding to 1.5×10^{11} photons per meV, a FWHM of 40 meV and an rms width of 32 meV. Figures 4.9-1 through 4.9-5 show the power, spectrum and phase profiles of the output radiation.

4.9.3 Medium Pulse “M2”, “M4”, and “M6” Simulation Results

Simulations of medium-duration bunches include the “M2”, “M4”, and “M6” distributions. The oldest medium-duration bunch example shown, the “M2” distribution, is similar to the “L2” distribution in that there is a noticeable quadratic variation in the electron energy along the bunch, as can be seen in Fig. 4.9-7. The total bunch duration, however, is only 780 fs. There is a slight current ramp, from 900 A in the head to 1100 A in the tail, and the energy spread is roughly 80 keV. The phase space distribution and FEL output are shown below. Although the final bunching is highest in the tail of the electrons, going up to 0.4, the output power is relatively flat. The typical power is 1.5 GW, although some spikes in power go significantly higher, and the FWHM is 500 fs. Note that there is also a “notch” near $t=200$ fs where the beam does not bunch or radiate effectively at all, probably due to a disruption in the beam profile. It is unclear if the small-length scale features of the distribution are realistic or not. The total energy in the pulse is 0.66 mJ, of which 89% is in a Gaussian mode with $Z_R=9.0$ m and waist location -8.9 m from the end of the FEL. The photon energy spectrum has one main spike with two significant side bands, one of which has roughly 2/3rds of the peak intensity. The peak corresponds to 4.4×10^{11} photons per meV with a FWHM of 10 meV; however, the overall rms width is 104 meV due to broad sidebands. Figures 4.9-8 through 4.9-10 show the power and phase profiles, and the spectrum of the output radiation. The variation in phase, which is an amplified form of the energy variation across the bunch, is the main cause of the side bands in the spectrum as well as the broadening of the central spike.

The other two distributions are designed to minimize the energy variation across the bunch. The longitudinal phase space profiles at entrance to the undulator are shown in Fig. 4.9-11. The “M4” distribution is very flat in energy, and has a central current of 850 A, increasing towards the head and tail. The energy spread is roughly 100 keV. The output power (Fig. 4.9-12) rises from 1 GW near the tail of the pulse to over 3 GW at the head of the pulse. The electron bunching also increases, from 0.3 to 0.5. The total energy in the pulse is 0.66 mJ, of which 92% is in a Gaussian mode with $Z_R=8.7$ m and waist location -8.6 m from the end of the FEL. The phase profile (Fig. 4.9-13) has a much weaker quadratic component, and the spectrum (Fig. 4.9-14) is much cleaner although there is a secondary peak at roughly half the amplitude of the main peak. The peak in the spectrum corresponds to 1.0×10^{12} photons per meV

with a FWHM of 10 meV and an rms of 91 meV. Restricting attention to the spectral range shown in the figures, the rms is then 47 meV. The long-bunch examples have small peaks at and beyond ± 200 meV in the spectrum, which are still sufficient to make a significant impact on the rms width.

The “M6” distribution has an extremely flat energy profile, while the current varies from 700 A in the center of the bunch to 900 A near the head and tail. The energy spread is 150 keV in the center of the bunch, and increases rapidly towards the edges of the distribution. The output radiation power (Fig. 4.9-12, right) is significantly reduced to about 400 MW, and has large temporal fluctuations. The FWHM of the power profile is about 400 fs. The total energy in the pulse is 0.14 mJ, of which 89% is in a Gaussian mode with a $Z_R=5.9$ m and waist location -4.9 m from the end of the FEL. While the output eikonal phase (Fig. 4.9-13, right) has very little temporal variation, the peak of the energy spectrum (Fig. 4.9-15) is only 2.6×10^{11} photons per meV, due to the reduced energy output per pulse. The spectrum has a FWHM of 10 meV, with an rms of 72 meV. Narrowing the range of the spectrum to that shown in the figures reduced the calculated rms to 58 meV. While both distributions correct the quadratic phase variation in the output pulse, the “M4” yields significantly improved spectral brightness while the “M6” suffers greatly from the increased energy spread and reduced central current. For both the M4 and M6 distributions, while the sidebands are much smaller in relative amplitude, they occur over a large frequency offset, and thus contribute significantly to the rms calculation.

4.10 Wakefield Calculations for the L2 & L4 Long Pulse Distributions

As was done for the M2 and M6 distributions in Sec. 3.4.3, we also calculated the expected longitudinal wakes for the L2 and L4 pulse distributions that are most relevant to FEL-2. As before, the Al vacuum chamber characteristics were 10.0-mm inner diameter, a surface roughness of 100-nm amplitude with a longitudinal period of 25 microns, and a presumed 10-cm break occurring every 3.4 meters. The resistive wake calculation was based upon an AC conductivity model. Figures 4.10-1 and 4.10-2 show the calculated wakes versus time for the L2 and L4 distributions, respectively. Once again, with the exceptions of possible spikes in the head and tail regions, over an interval exceeding 1.0 ps there is a nearly constant wake of ~ 4 kV/m with temporal fluctuations of ± 2 kV/m or so. Over a 50-m total vacuum chamber length, a fluctuation of ± 100 kV corresponds to less than 0.01% of total energy and is unlikely to cause any significant degradation of emission. On the other hand, if the surface roughness were to grow to ~ 500 -nm amplitude, the fluctuations from this wake component increase 25X to approximately 15 kV/m. Over 40 m of vacuum chamber, this wake would lead to a normalized beam energy fluctuation of 0.05% which could prove quite troublesome. Consequently, it may be important to do some careful studies to help set specifications and tolerances for the FEL-2 vacuum chamber.

4.11 Comparison of Known Advantages and Disadvantages of FB vs. WB

Both the fresh-bunch and whole-bunch approaches to FEL-2 have their own advantages and disadvantages. The fresh-bunch approach in general requires a much shorter final radiator and in general is much less sensitive to the initial incoherent energy spread. The energy and eventual microbunching modulation in each stage is quite strong and this together with the shorter final radiator implies that the signal-to-noise ratio should be relatively good. Moreover, pump-probe experiments are relatively straightforward with this configuration. On the other hand, the fresh-bunch approach is quite sensitive to timing jitter and if the RMS jitter value is significantly greater than ~ 250 fs, the amount of useful pulse can be quite small if one is concerned with minimizing shot-to-shot jitter in the output pulse energy. By definition, the duration of the final output pulse can be no longer than half the electron beam pulse which

implies the ultimate Fourier transform bandwidth will be twice as large than the whole-bunch approach for an unchirped electron pulse. Moreover, the fresh-bunch configuration is somewhat more complicated in that a temporal delay section is needed.

The whole-bunch approach has an important advantage in that timing jitter can be overcome by simply having a seed laser pulse sufficiently long in duration that there is no question it will fully temporally cover the electron beam pulse. Similarly, for an unchirped electron beam pulse, the transform limit of the output spectral bandpass is at least twice as narrow as the fresh-bunch approach. The undulator configuration is somewhat simpler and the required input laser seed power is generally 3-10X smaller. On the other hand, the whole-bunch approach can actually have a significantly worse bandpass than the fresh-bunch approach if there is a great deal of a quadratic chirp in energy on the electron beam. Since the final radiator is in the deep exponential gain regime, the output power can be very sensitive to electron beam parameters such as current and incoherent energy spread. The achievable signal-to-noise ratio is also much worse because SASE has a much greater number of gain lengths in which to grow. The whole-bunch configuration also appears to be more sensitive to initial transverse tilts and offsets in the electron beam.

From an experimental point of view and if time permits, it is best not to make a firm choice between the two approaches until good diagnostic results are obtained from FEL-1. If the initial (-at undulator entrance) incoherent energy spread is "low" (*i.e.*, < 100 keV), this may favor the whole-bunch approach. On the other hand if it is high (*i.e.*, > 200 keV) or if there is a large, unremovable quadratic energy chirp on the e-beam, fresh-bunch would be favored. In fact, since the two configurations are not that all different, one could begin with the fresh-bunch approach and then later remove the temporal delay section and add some additional second stage radiator segments to test out the whole-bunch approach.

4.12 Diagnostic Needs Unique to FEL-2

For many items, the diagnostic needs of FEL-1 and FEL-2 are quite similar. The first stage of FEL-2 can be considered a somewhat shorter version of FEL-1 and the various diagnostics proposed for FEL-1 could be replicated here. For the fresh-bunch approach, the output radiation of the first stage will be used to modulate a portion of the electron bunch so detailed radiation diagnostics could prove extremely useful. In order to tune the delay section between the first and second stages, some sort of cross-correlator (e.g. between the first stage coherent signal and the spontaneous emission from the second stage) might be useful. Some sort of microbunching diagnostic after the second stage modulator is likely to be essential. A diagnostic to resolve in z the build up of coherent radiation in the second stage radiator would also provide very important information. Since the undulator gaps can be opened up, a gross mistuning downstream of the wanted diagnostic point might be sufficient.

For the whole-bunch approach, a microbunching diagnostic at the break between the first and second stage radiators would be useful. It is also important to establish if the coherent bunching at the final wavelength is significantly greater than the initial shot noise value; if not, SASE contamination will be extreme. As in the fresh-bunch approach, a z -resolved power diagnostic is also useful along the final radiator.

=====

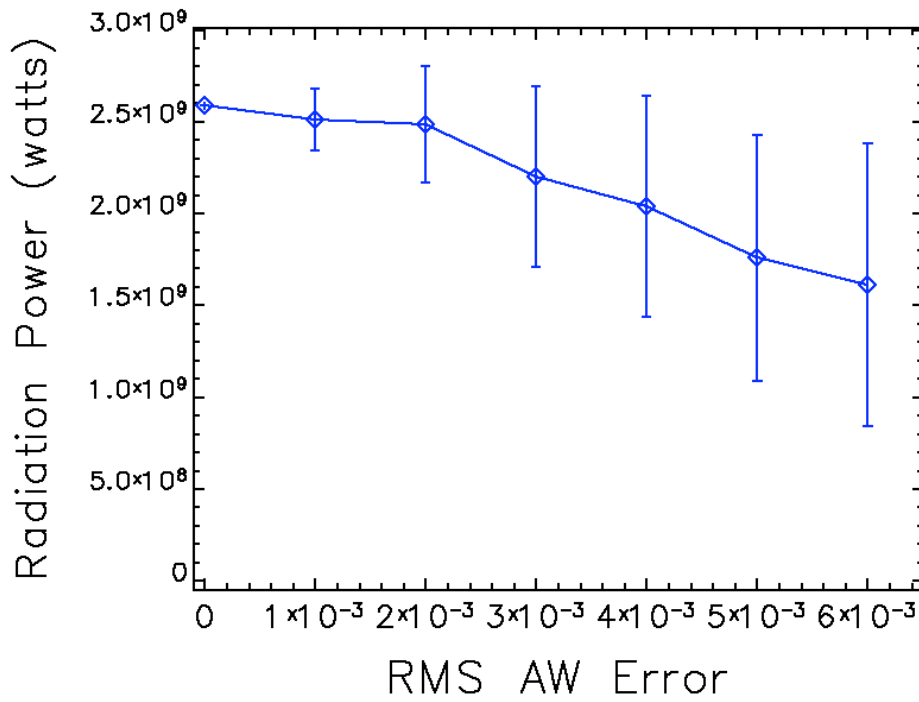


Figure 2.3-1: Output radiation power at 40-nm from FEL-1 in the presence of random undulator segment mistunings at as a function of RMS error in aw. The the diamond symbol and error bars refer to the mean and standard deviation over 64 independent mistunings. The distribution of errors at individual segments follow a one-dimensional Gaussian.

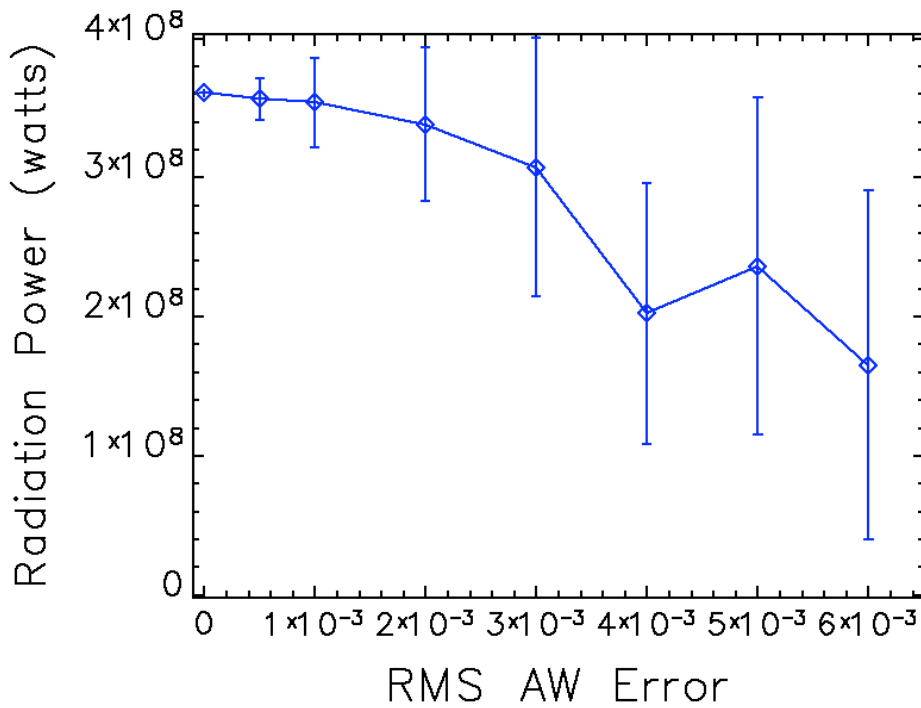


Figure 2.3-2: Output radiation power at 10-nm from FEL-2 (fresh bunch approach) in the presence of random undulator segment mistunings at as a function of RMS error in aw. The the diamond symbol and error bars refer to the mean and standard deviation over 25 independent mistunings. The distribution of errors at individual segments follow a one-dimensional Gaussian.

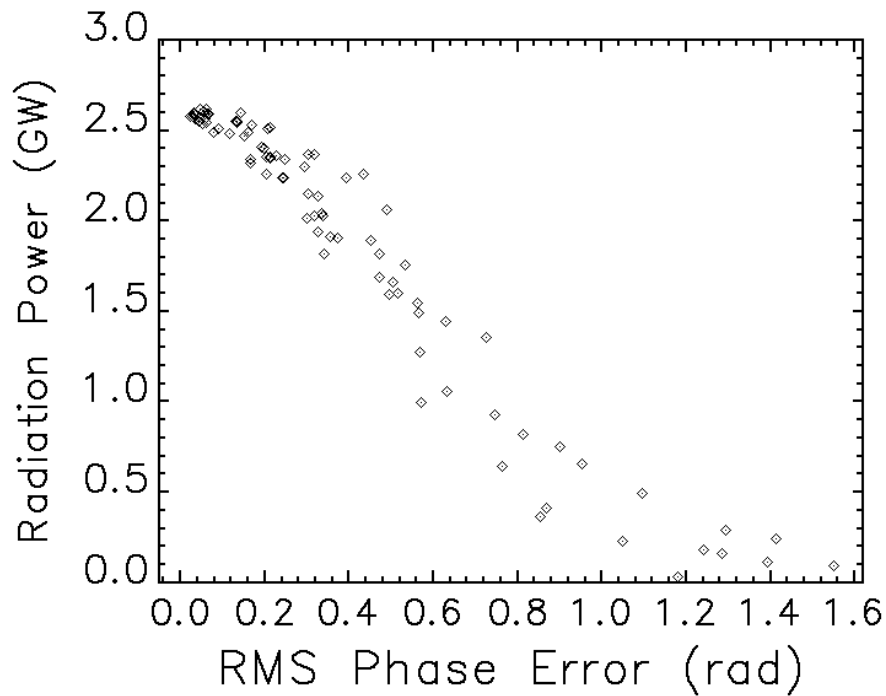


Figure 2.3-3: Output radiation power at 40-nm from FEL-1 in the presence of random undulator pole strength errors as a function of residual RMS phase error as evaluated at all locations within the actual undulator segments. Both transverse tilt and offset corrections were made at two locations within each segment (*i.e.*, two virtual shims), in addition to removing the net longitudinal phase error at the ends of each segment. Each diamond is the result of an independent set of random pole mistunings following a 1-D Gaussian distribution.

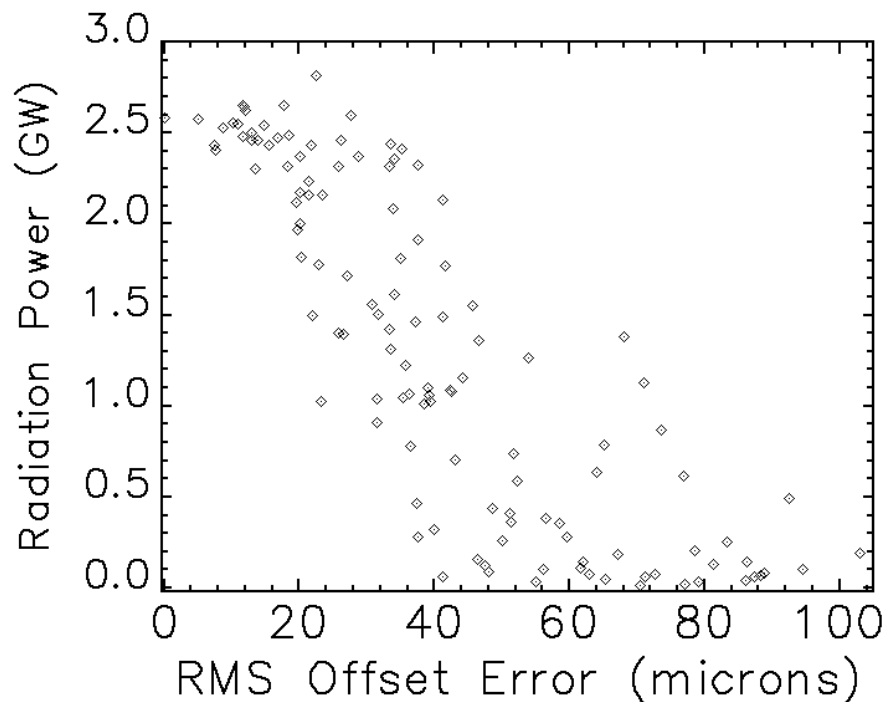


Figure 2.3-4: Output radiation power at 40-nm from FEL-1 in the presence of random undulator pole strength errors as a function of residual RMS orbit offset error as evaluated at all locations within the actual undulator segments. These data come from the same set as that in Fig. 2.3-3.

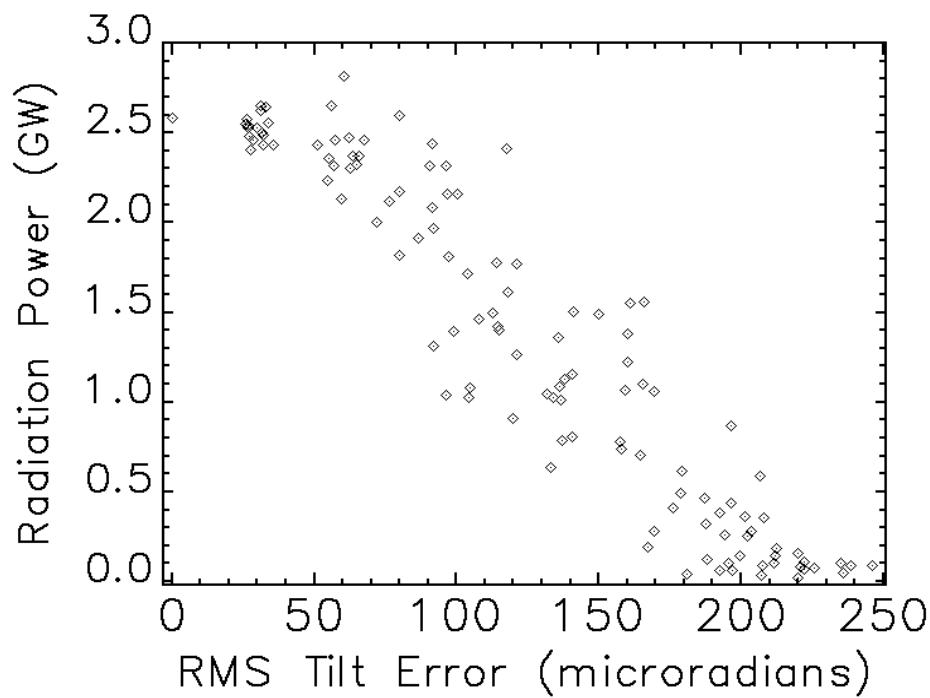


Figure 2.3-5: Same as the preceding two figures except the independent variable is the RMS tilt error of the electron beam orbit.

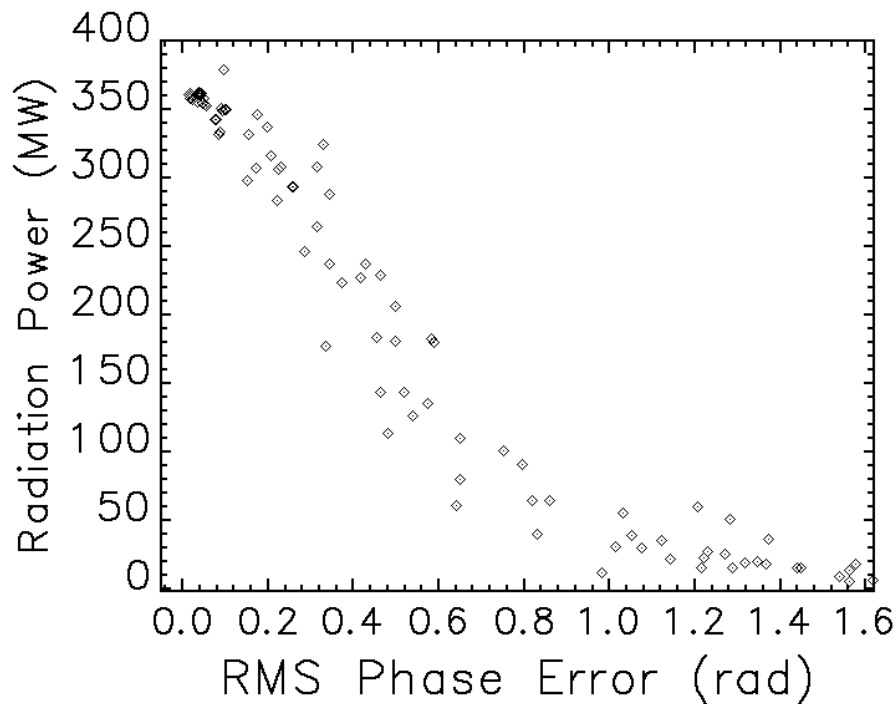


Figure 2.3-6: Output radiation power at 10-nm from FEL-2 (fresh bunch approach) in the presence of random undulator pole strength errors in the final radiator as a function of residual RMS phase error as evaluated at all z-locations in the radiator with the exception of the break spaces.. Both transverse tilt and offset corrections were made at two locations within each final radiator segment (i.e., two virtual shims), in addition to removing the net longitudinal phase error at the ends of each segment. Each diamond is the result of an independent set of random pole mistunings following a 1-D Gaussian distribution.

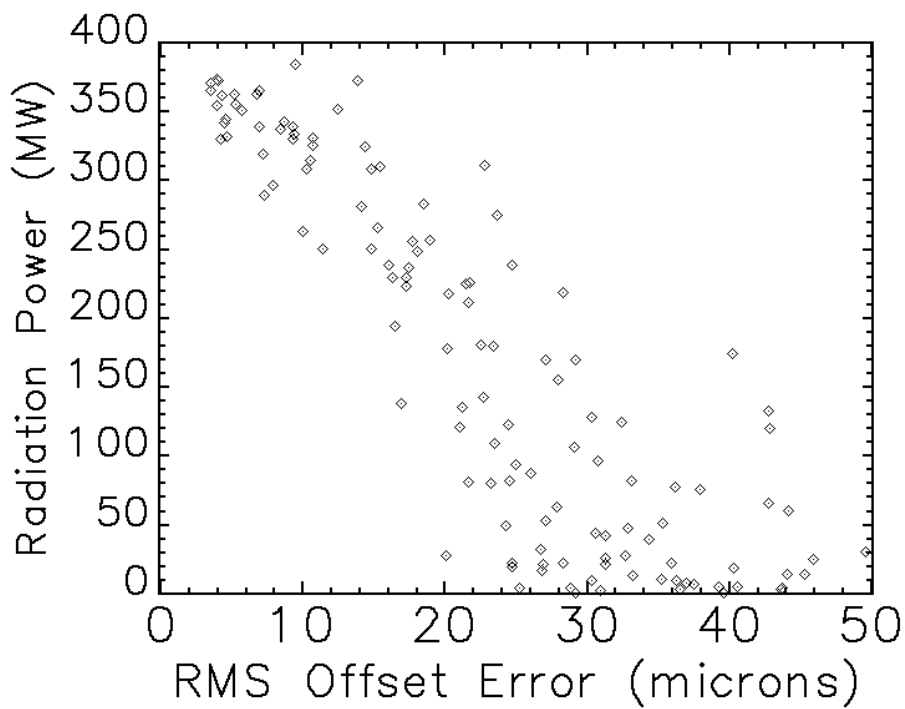


Figure 2.3-7: Output radiation power at 10-nm from FEL-2 in the presence of random undulator pole strength errors in the final radiator as a function of residual RMS orbit offset error as evaluated at all locations within the actual undulator segments. These data come from the same set as that in Fig. 2.3-6.

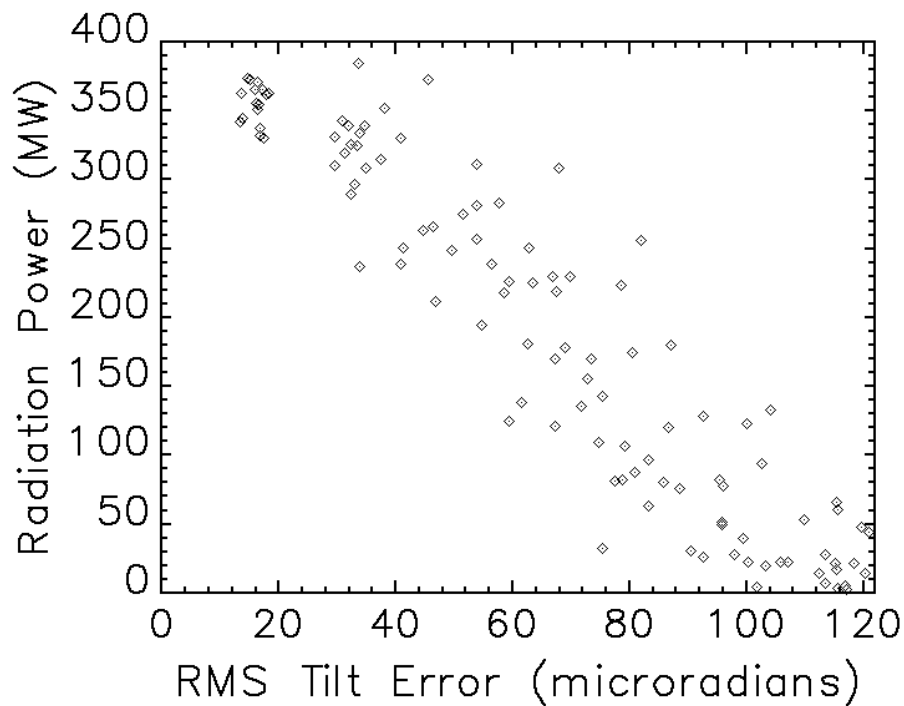


Figure 2.3-8: Same as the preceding two figures except the independent variable is the RMS tilt error of the electron beam orbit in the final radiator of FEL-2.

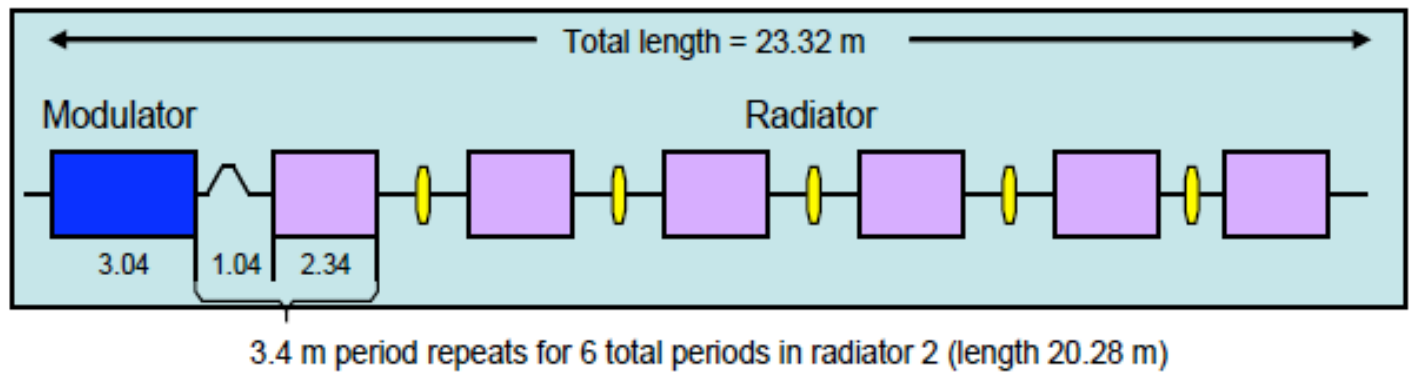


Figure 3.1-1: Nominal undulator layout for FEL-1. The radiator consists of 6 sections, each 2.34 m long separated by 1.04 m drifts containing a focusing quad, phase corrector, and diagnostics. Between the modulator and radiator is a dispersive section whose purpose is to convert energy modulation into strong microbunching.

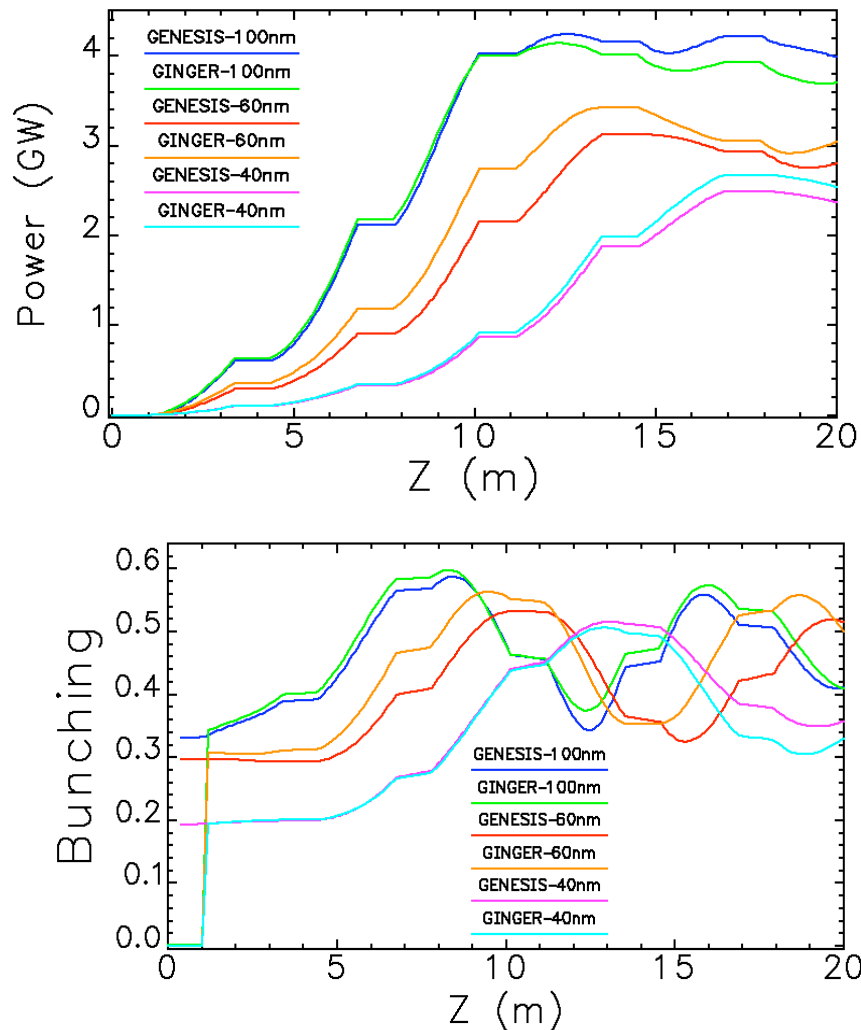


Figure 3.1-2: GENESIS and GINGER results for radiation power $P(z)$ and microbunching fraction $b(z)$ for FEL-1 at 100-, 60-, and 40-nm wavelengths.

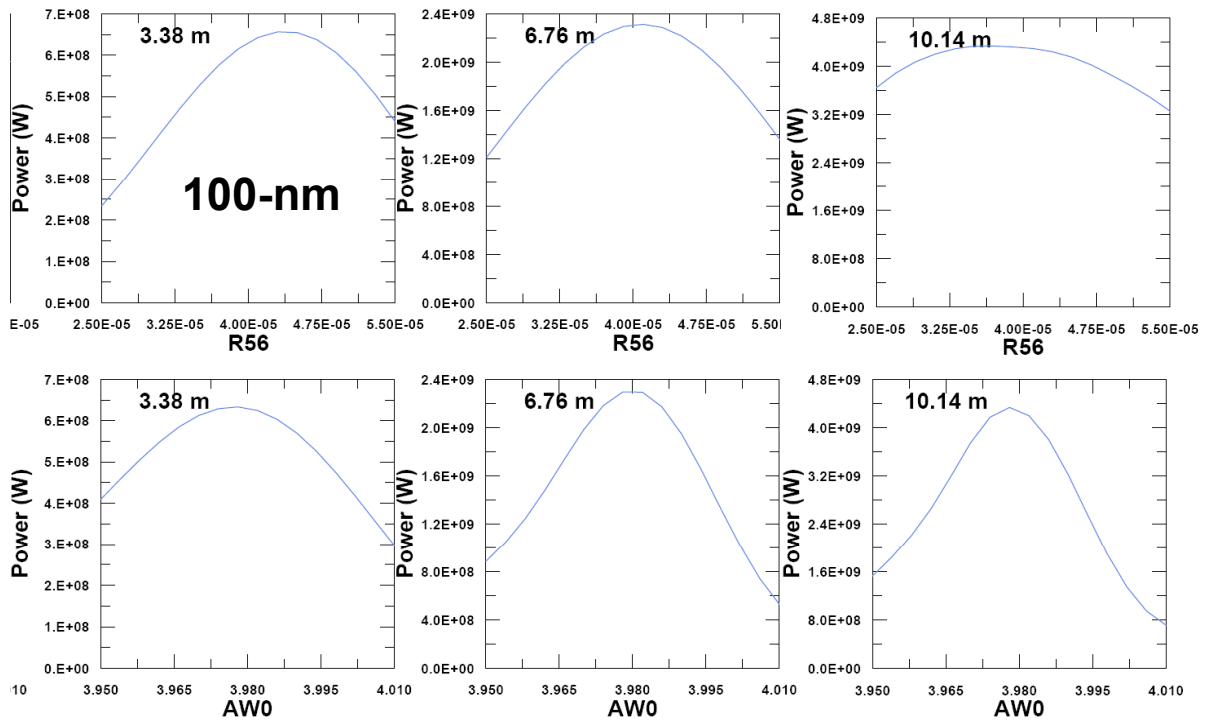


Figure 3.2-1: The upper plot shows the scaling of FEL output power at 100-nm vs. R_{56} after each of the first 3 radiator sections. The lower plot row shows the sensitivity to a_w at the same locations.

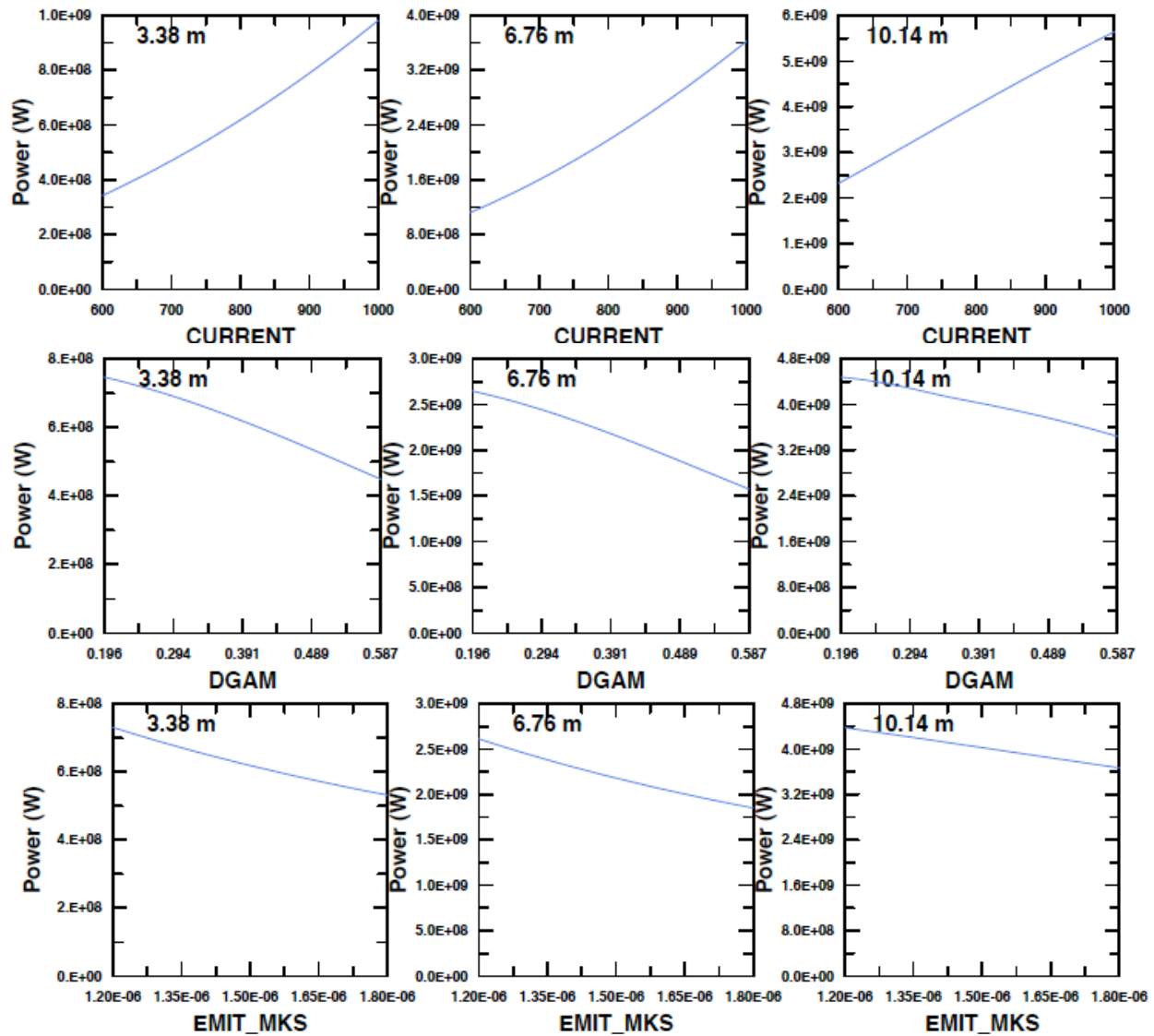


Figure 3.2-2: These plots show sensitivity of FEL performance at 100-nm wavelength to e-beam emittance, incoherent energy spread, and current. The z -locations correspond to the ends of the 1st, 2nd, and 3rd radiator undulator sections.

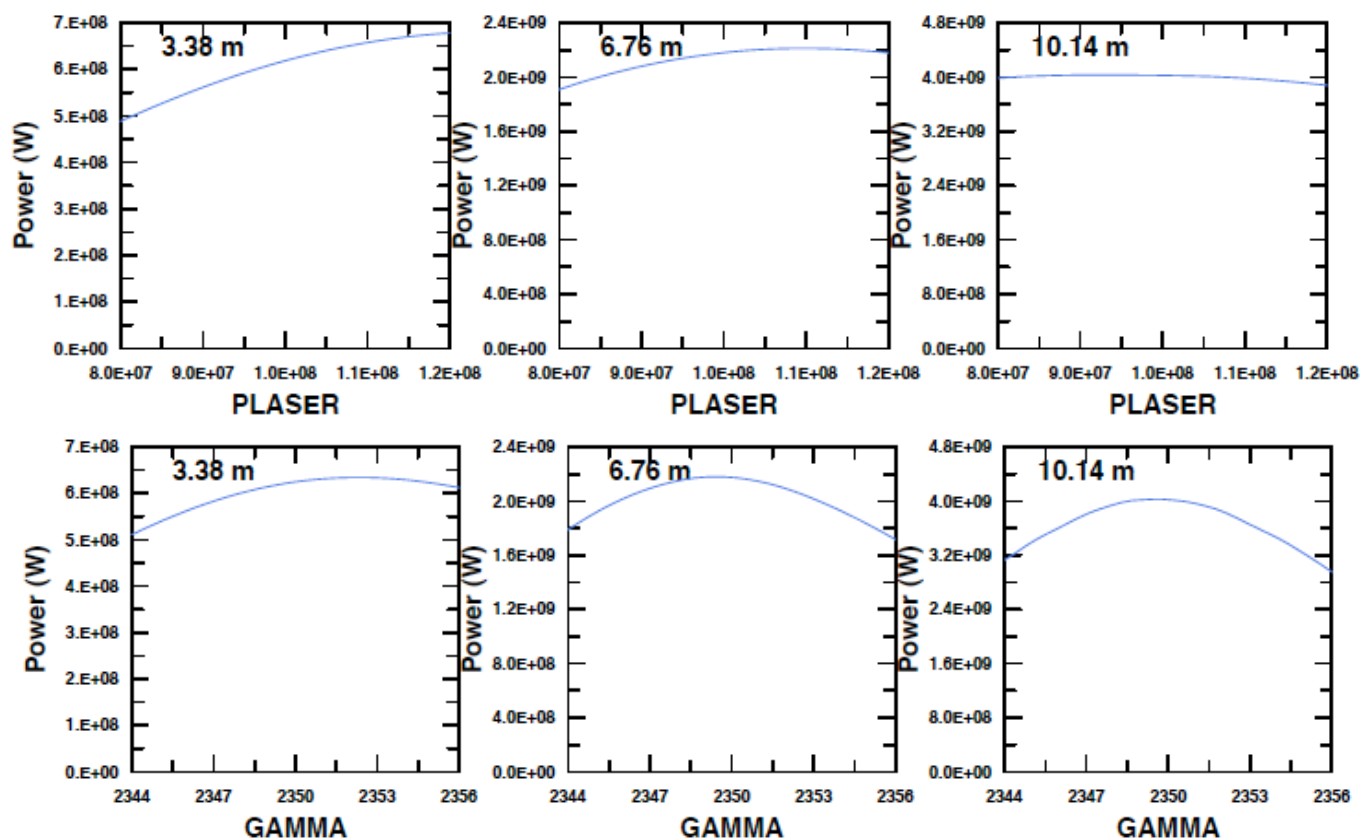


Figure 3.2-3: FEL performance at 100-nm wavelength to input seed power and electron energy. The z-locations correspond to the ends of the 1st, 2nd, and 3rd radiator undulator sections

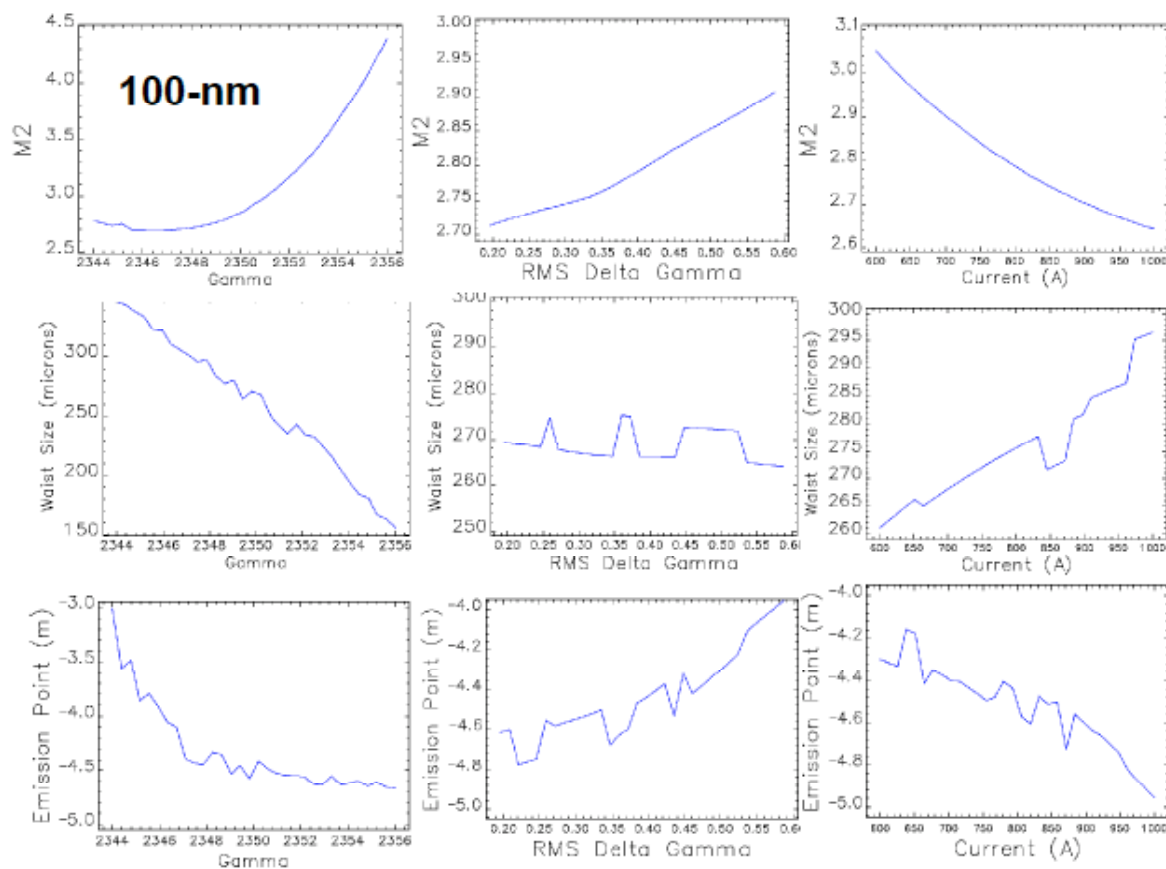


Figure 3.2-4: Far-field diagnostics for mode quality M^2 , virtual emission point location (relative to the end of the 3rd radiator undulator segment), and virtual waist size as a function of e-beam energy, σ_E , and current for FEL-1 at 100-nm output wavelength.

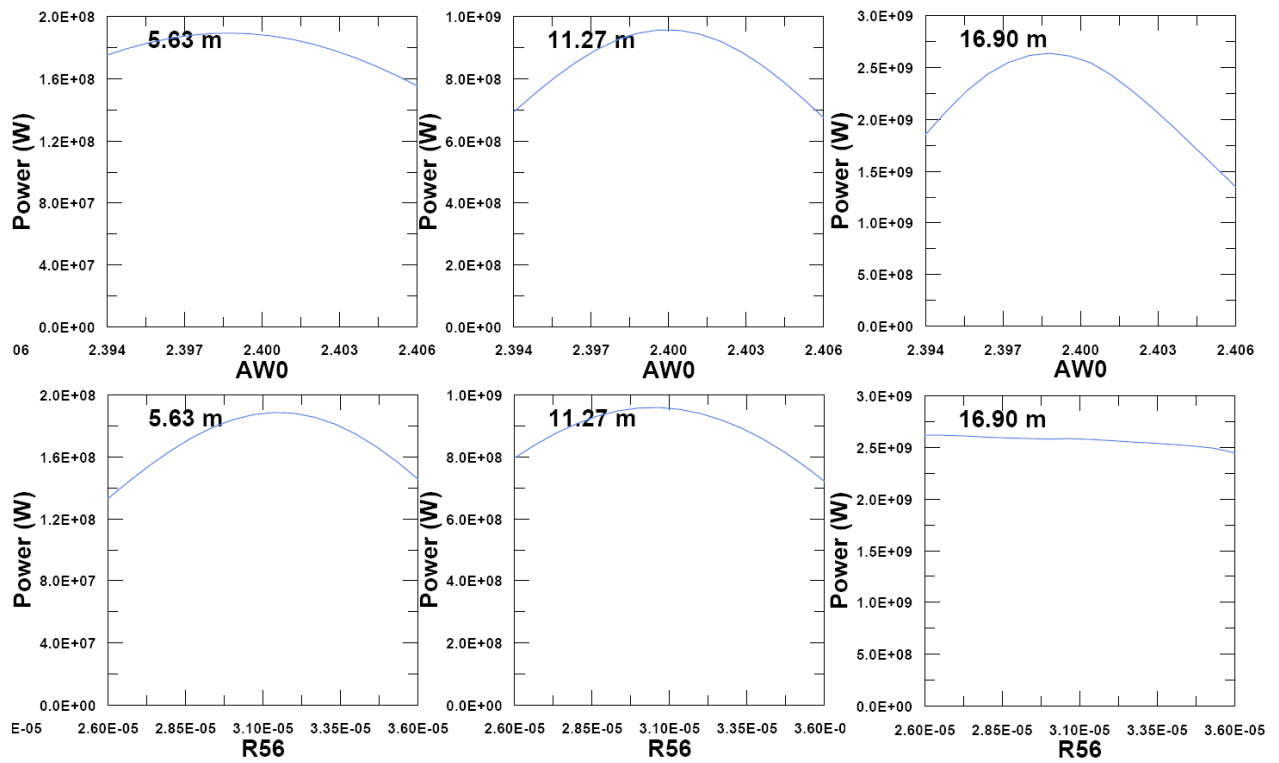


Figure 3.2-5: Optimization of FEL-1 radiator output power at 40-nm output wavelength with respect to R_{56} and a_w . The z -locations of 5.6, 11.3 and 16.9 m correspond to the middle of the 2nd, the middle of the 4th, and end of the 5th undulator sections

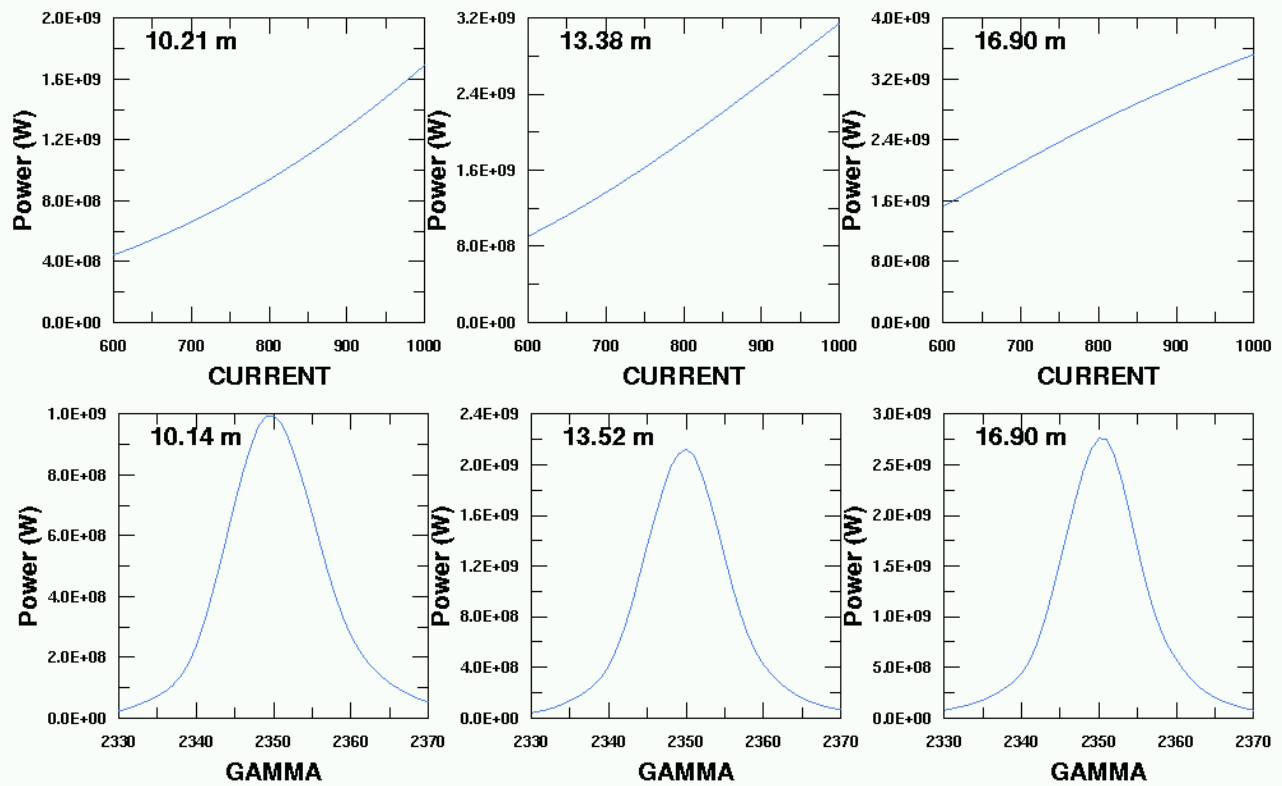


Figure 3.2-6: FEL performance at 40-nm wavelength to electron beam current and energy. The z -locations correspond to the ends of the 3rd, 4th and 5th radiator undulator sections

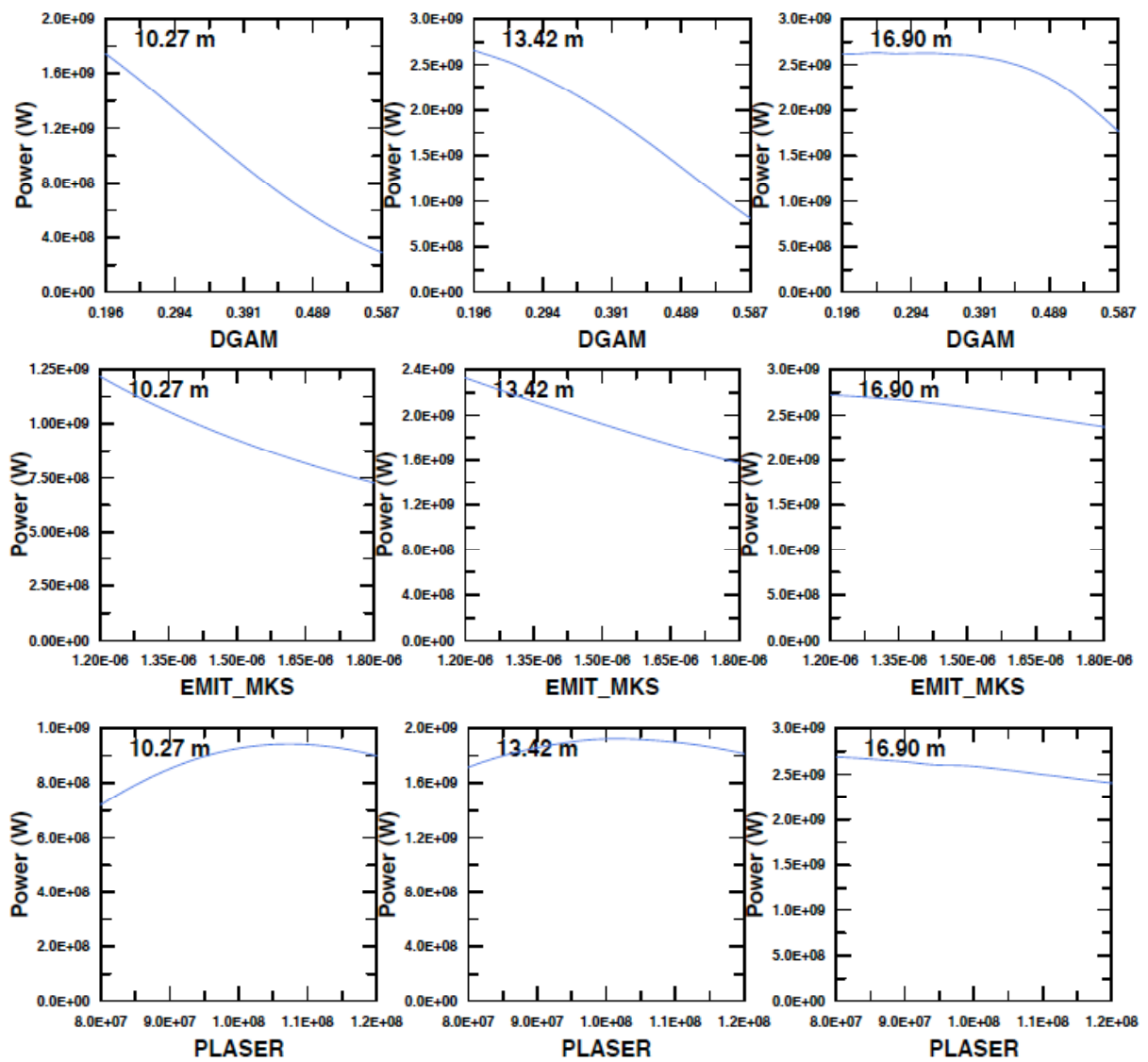


Figure 3.2-7: FEL performance at 40-nm wavelength to electron beam incoherent energy spread and emittance and also to the input seed laser power. The z -locations correspond to the ends of the 3rd, 4th and 5th radiator undulator sections

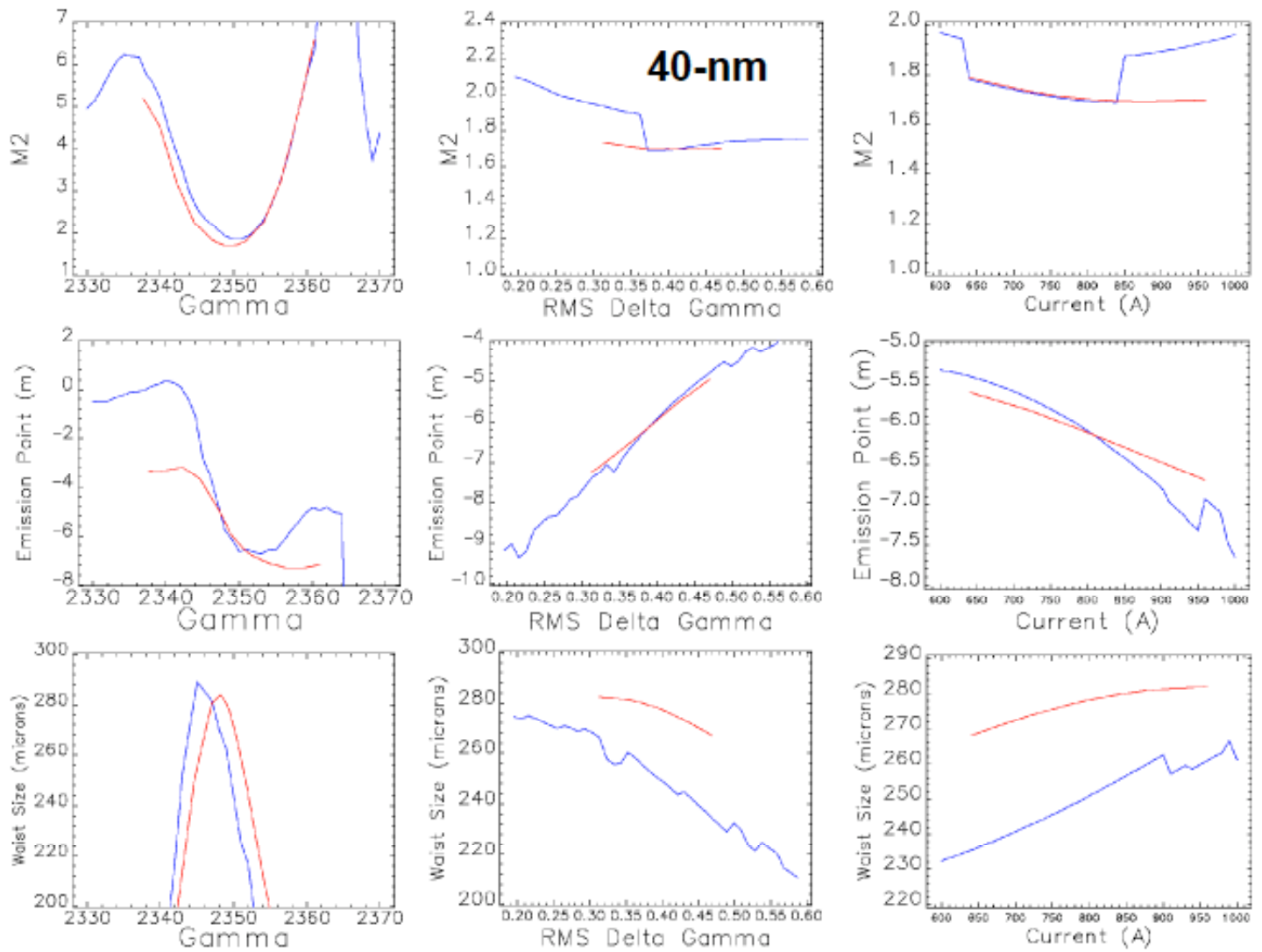


Figure 3.2-8: Far-field diagnostics from GINGER (blue curves) and GENESIS (red curves) results for output mode quality M^2 , virtual emission point location (relative to the undulator exit at 16.9 m), and virtual waist size as a function of several e-beam input parameters.

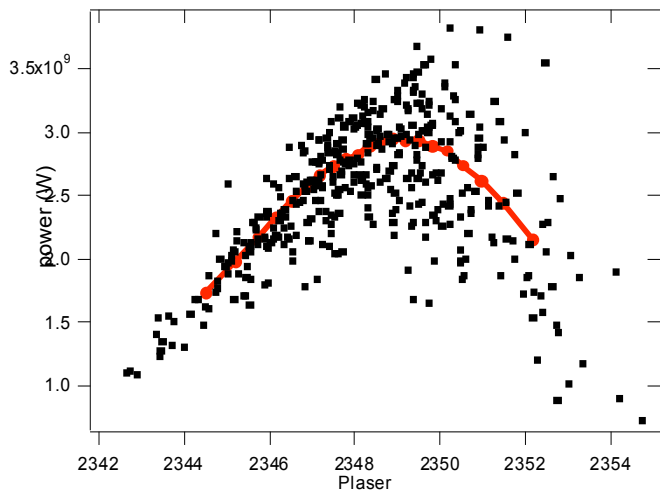


Figure 3.2-9: FEL output power as a function of the electron beam energy (γ) in the case of a single parameter only (red curve) and multiparameter (black dots) variation.

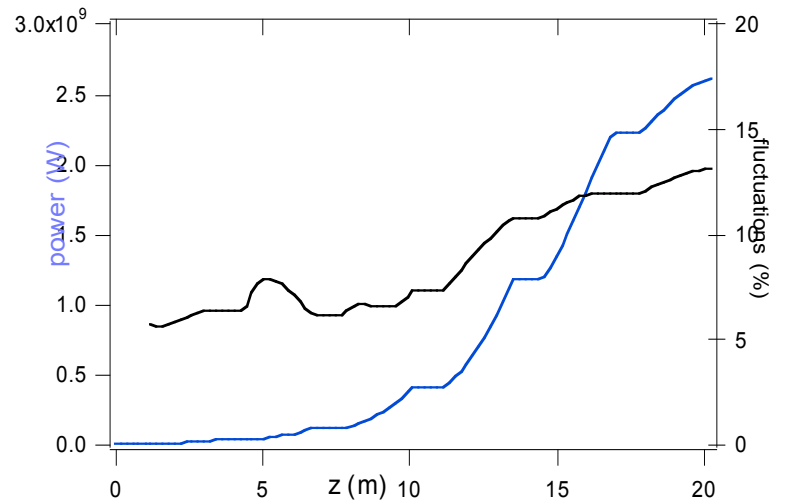


Figure 3.2-10: Statistical average (blue curve) and normalized fluctuation level (black curve) of the radiation power along the radiator for jitter in the beam energy γ whose distribution had a Gaussian width of 10%.

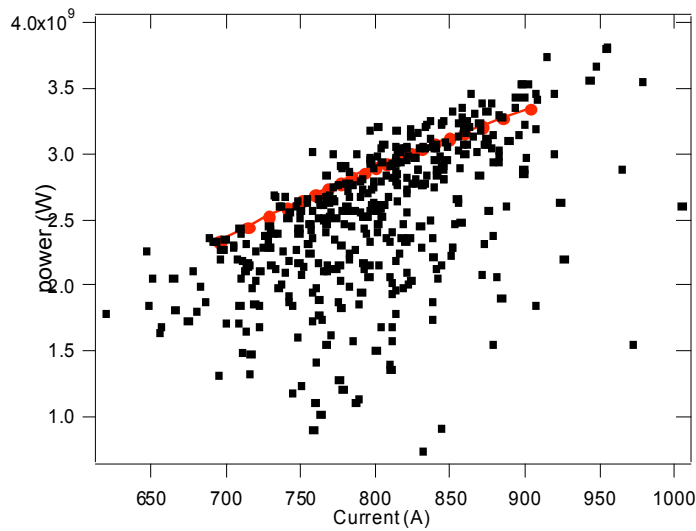


Figure 3.2-11: FEL output power as a function of beam current in the case of a single parameter only (red curve) and multiparameter (black dots) variation.

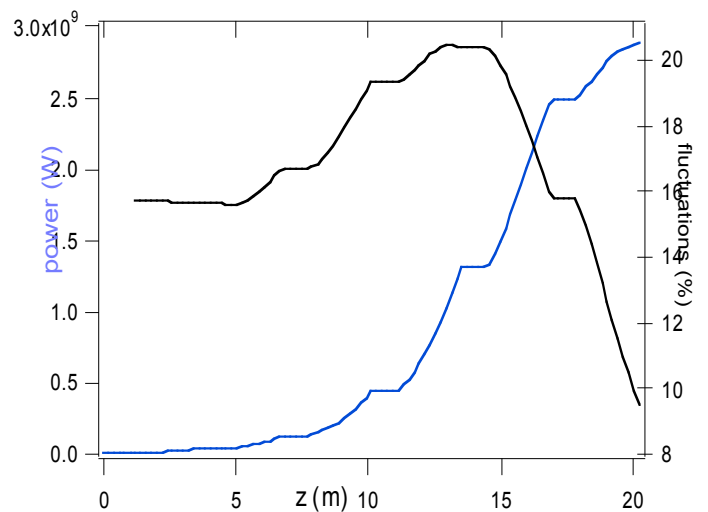


Figure 3.2-12: Statistical average (blue curve) and normalized fluctuation level (black curve) of the radiation power along the radiator for jitter in the beam current whose distribution had a normalized Gaussian width of 8%.

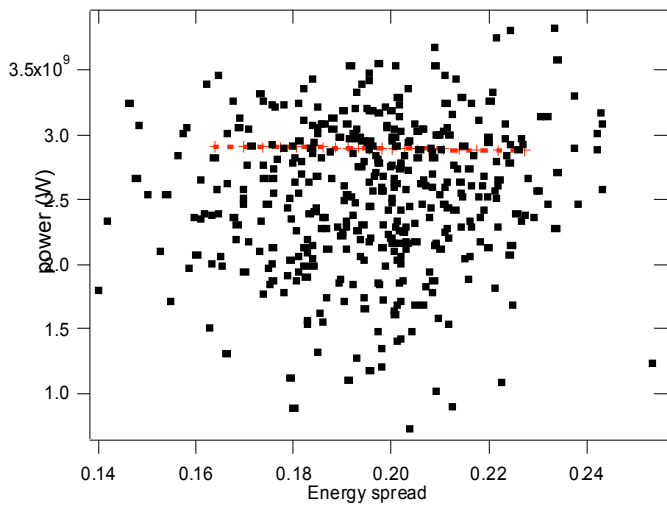


Figure 3.2-13: FEL output power as a function of instantaneous beam energy spread σ_y in the case of a single parameter only (red curve) and multiparameter (black dots) variation.

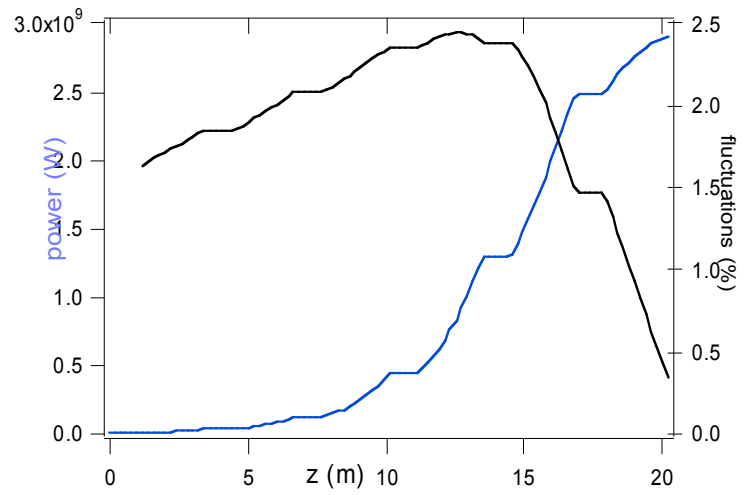


Figure 3.2-14: Statistical average (blue curve) and normalized fluctuation level (black curve) of the radiation power along the radiator for jitter in the beam energy spread whose distribution had a normalized Gaussian width of 10%.

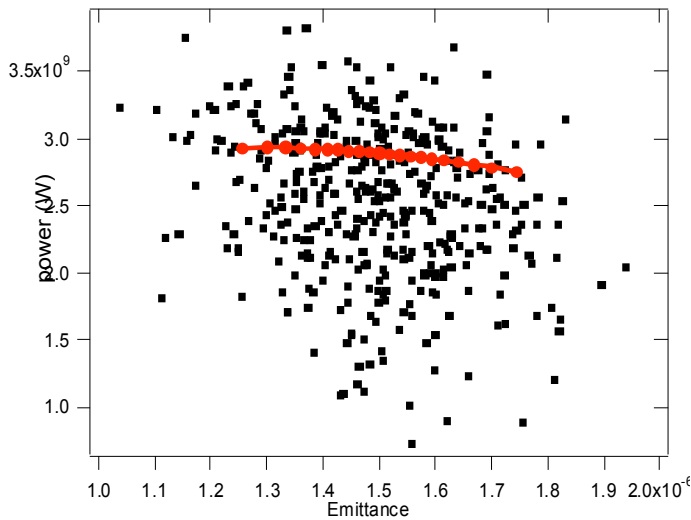


Figure 3.2-15: FEL output power as a function of normalized transverse beam emittance in the case of a single parameter only (red curve) and multiparameter (black dots) variation.

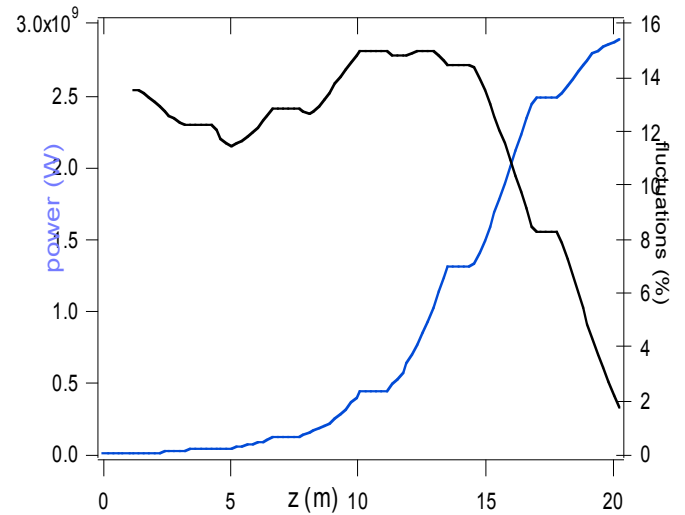


Figure 3.2-16: Statistical average (blue curve) and normalized fluctuation level (black curve) of the radiation power along the radiator for jitter in the beam emittance whose distribution had a normalized Gaussian width of 10%.

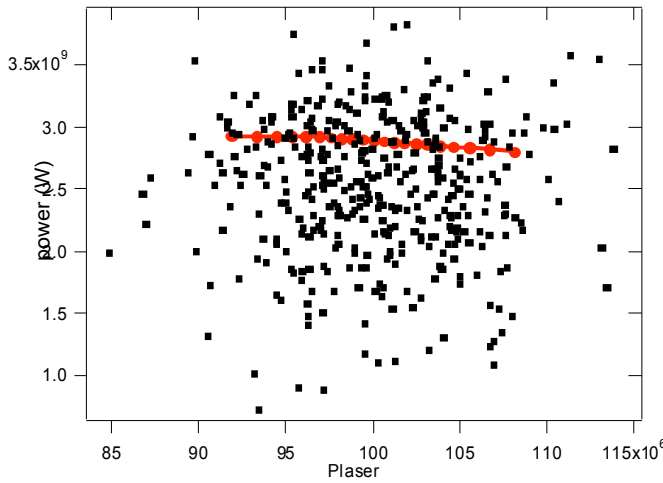


Figure 3.2-17: FEL output power as a function of seed laser power in the case of a single parameter only (red curve) and multiparameter (black dots) variation.

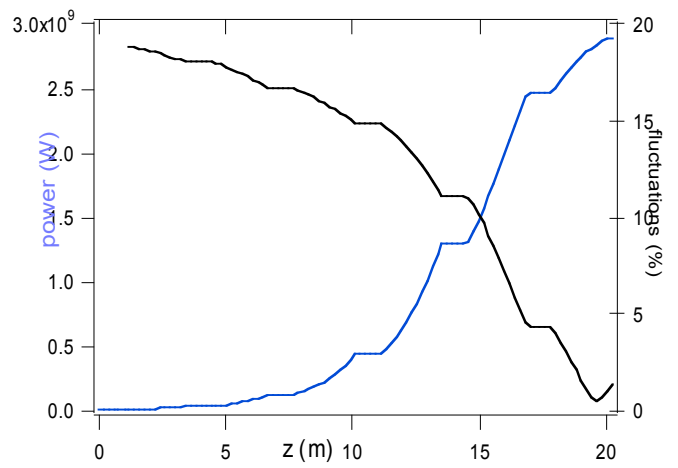


Figure 3.2-18: Statistical average (blue curve) and normalized fluctuation level (black curve) of the radiation power along the radiator for jitter in the input seed laser power whose distribution had a normalized Gaussian width of 10%.

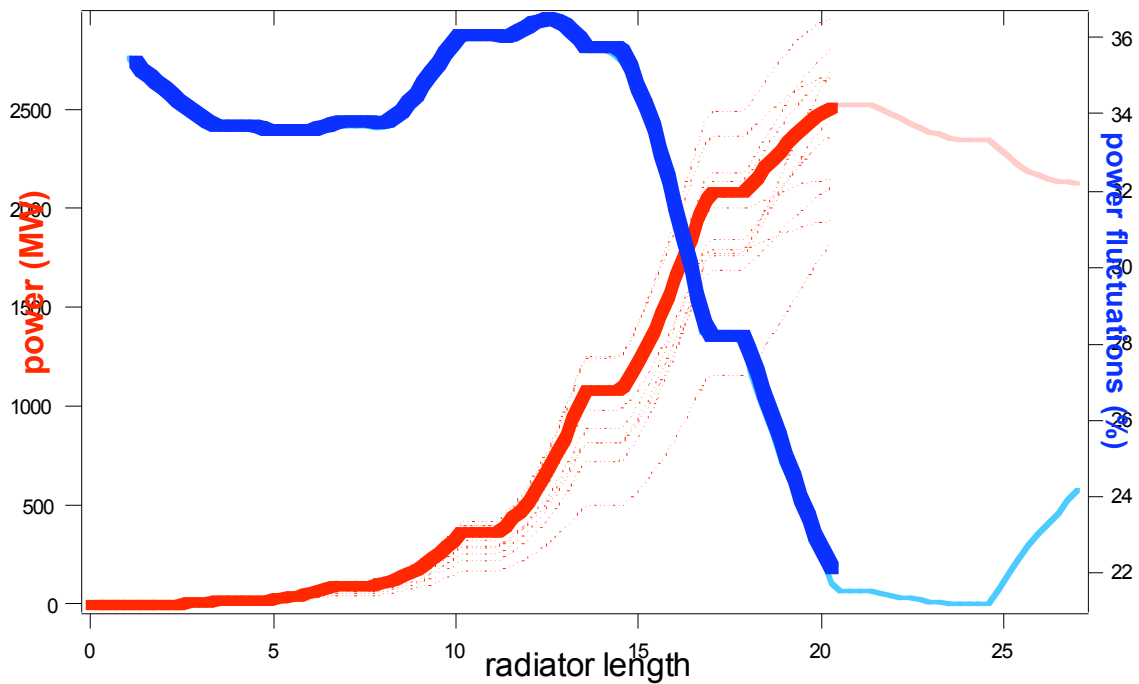


Figure 3.2-19: Statistical averages over the 400 multiparameter jittered runs of the FEL-1 40-nm power along the radiator (red solid line) and the sigma of the corresponding distribution (blue solid line). The dotted lines show evolution for some of the run for which the radiator was extended to ~28-m length in order to show that there is no advantage in using a longer radiator.

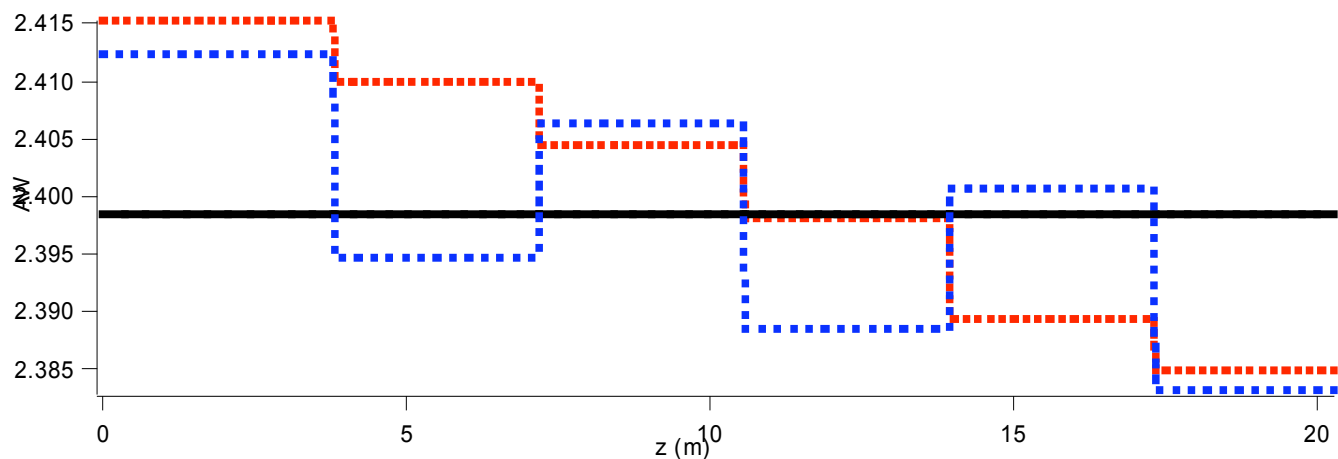


Figure 3.2-20: Some examples of the radiator magnetic strength taper used to reduce the FEL-1 output sensitivity to electron beam energy: black line: constant nominal value; red dotted line, linearly decreasing a_w values; blue-dashed line, alternating strength configuration.

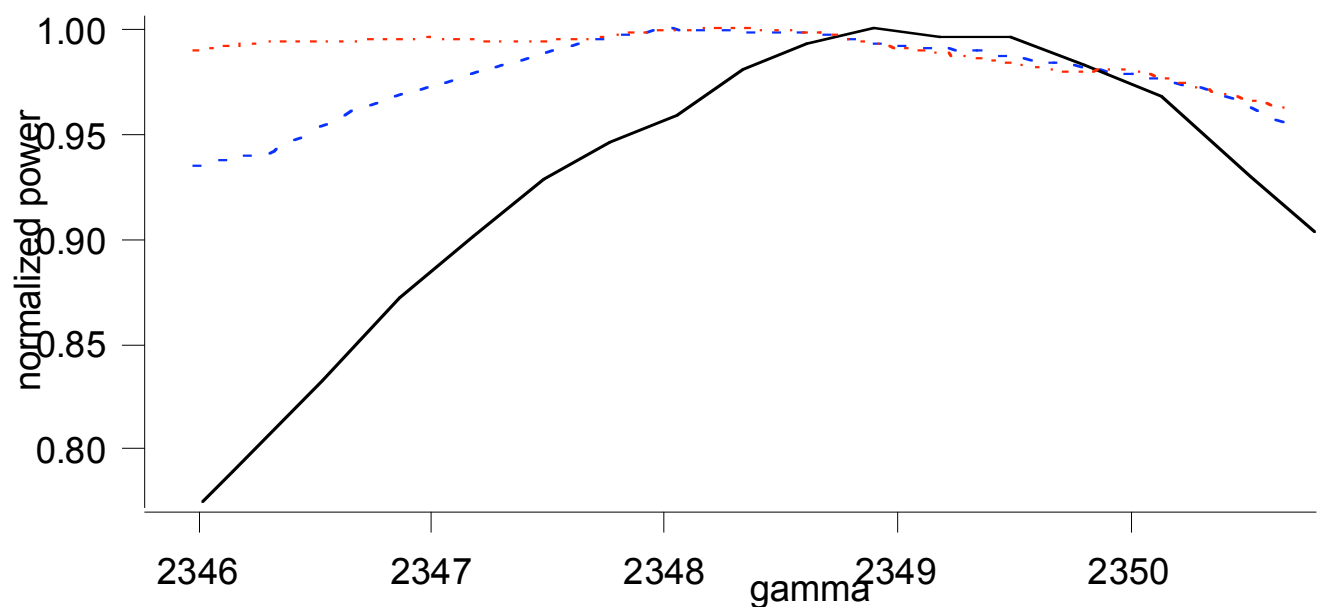


Figure 3.2-21: Some Output power vs γ obtained with the radiator configuration presented in Fig. 3.2-20: black line: constant nominal value; red dotted line, linearly decreasing a_w values; blue-dashed line, alternating strength configuration.

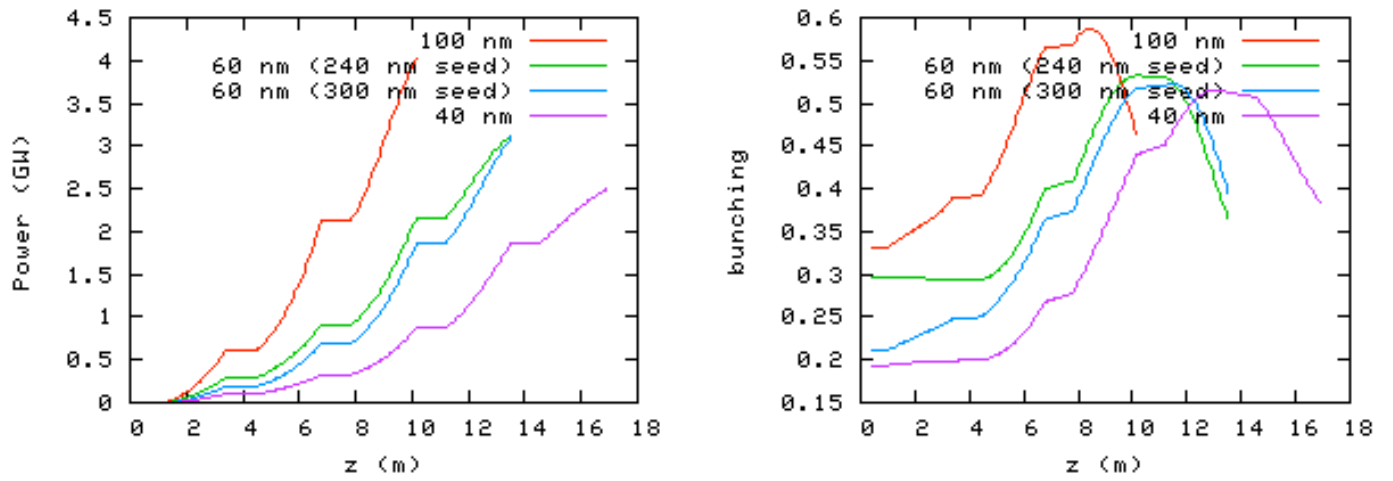


Figure 3.3-1: Power (left) and bunching (right) for different wavelengths and configurations as a function of position within the FEL-1 radiator.

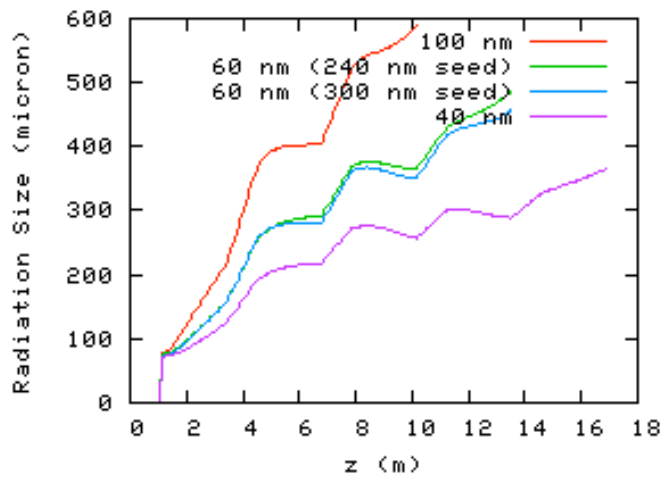


Figure 3.3-2: Radiation spot size for different FEL-1 configurations as a function of position within the radiator.

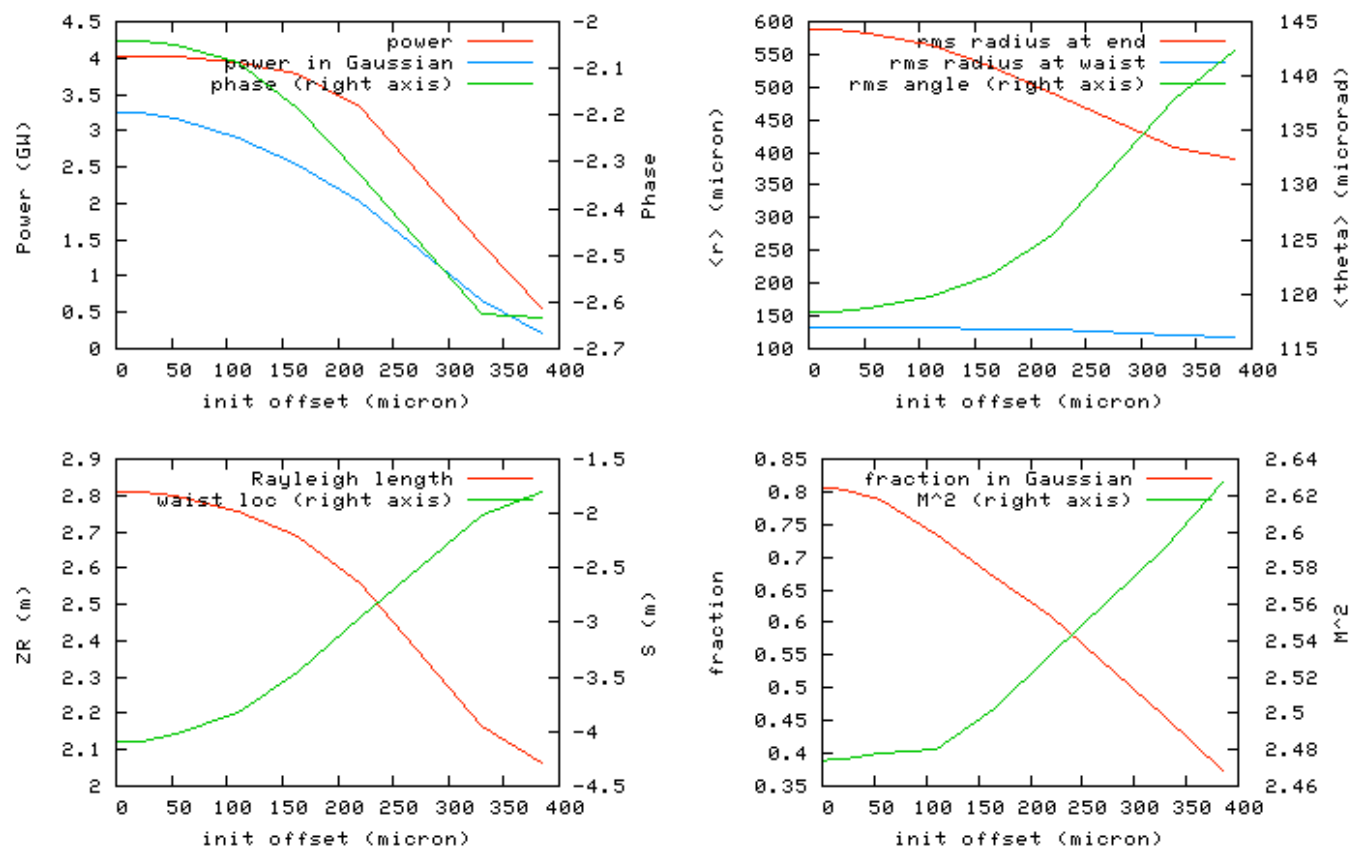


Figure 3.3-3: Sensitivity of 100-nm wavelength output to initial electron beam transverse offset.

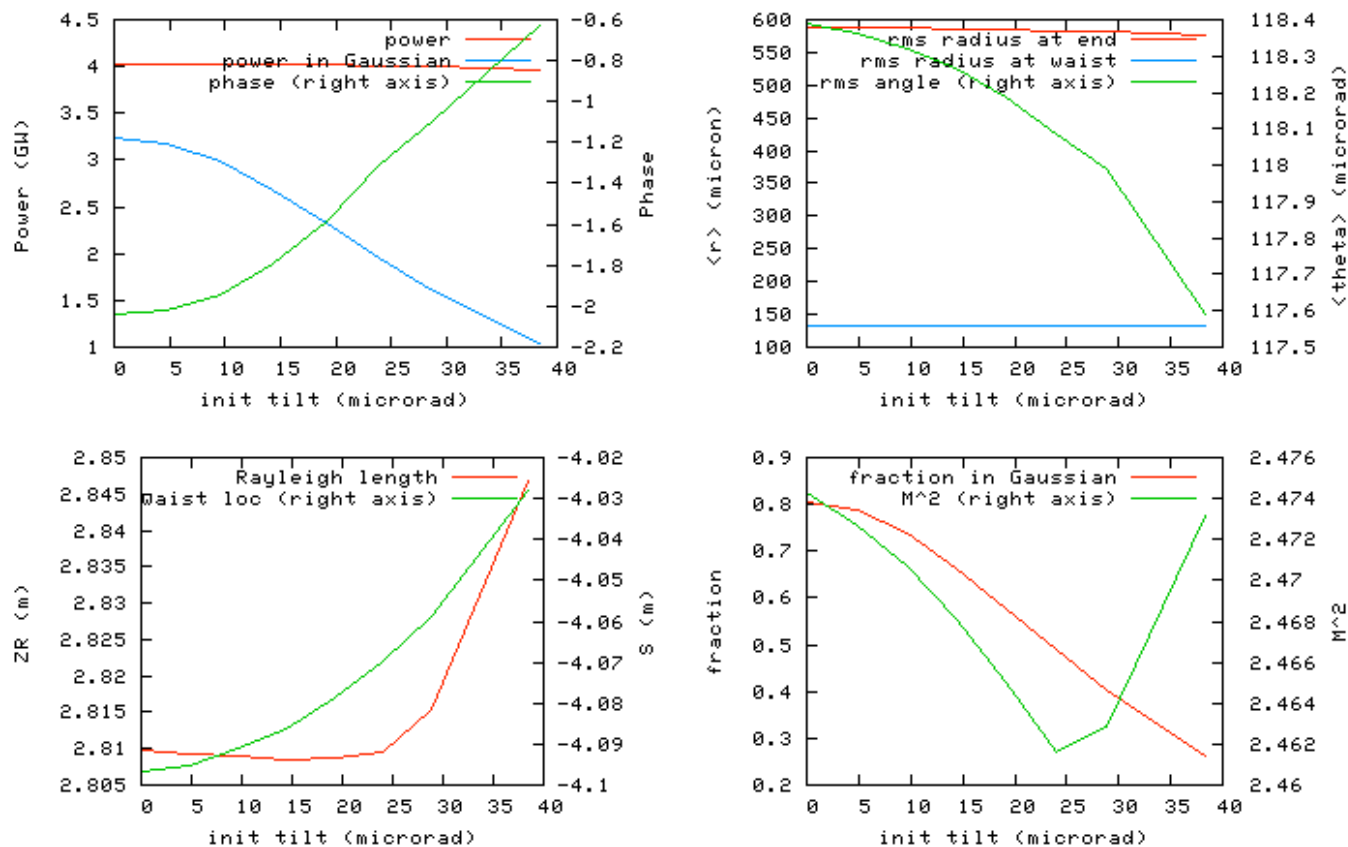


Figure 3.3-4: Sensitivity of 100-nm wavelength output to input electron beam tilt.

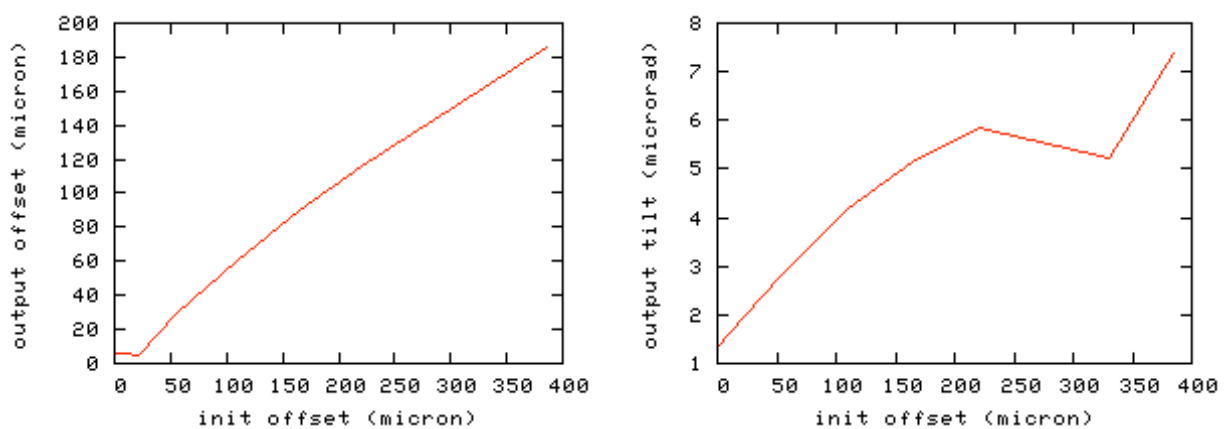


Figure 3.3-5: 100-nm output radiation misalignment as a function of initial electron beam offset.

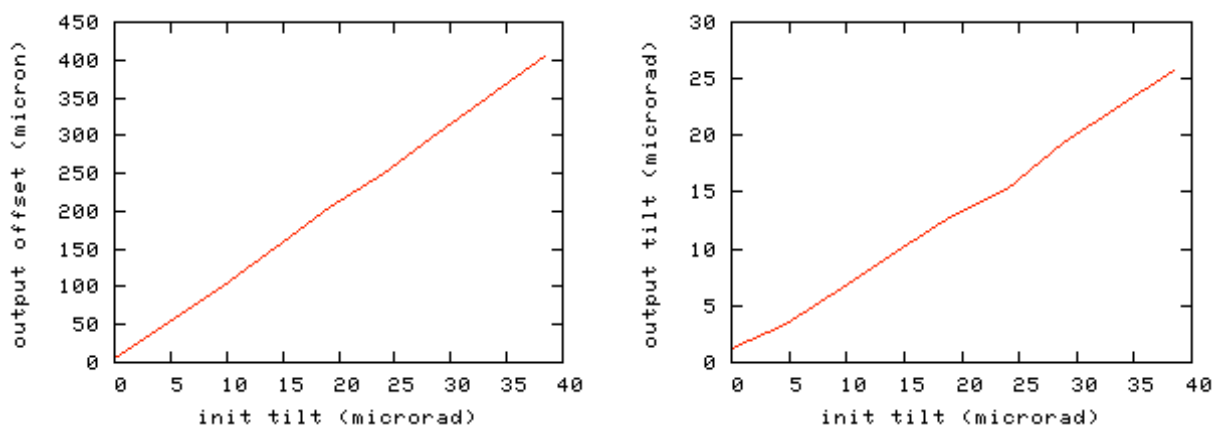


Figure 3.3-6: 100-nm output radiation misalignment as a function of initial electron beam tilt.

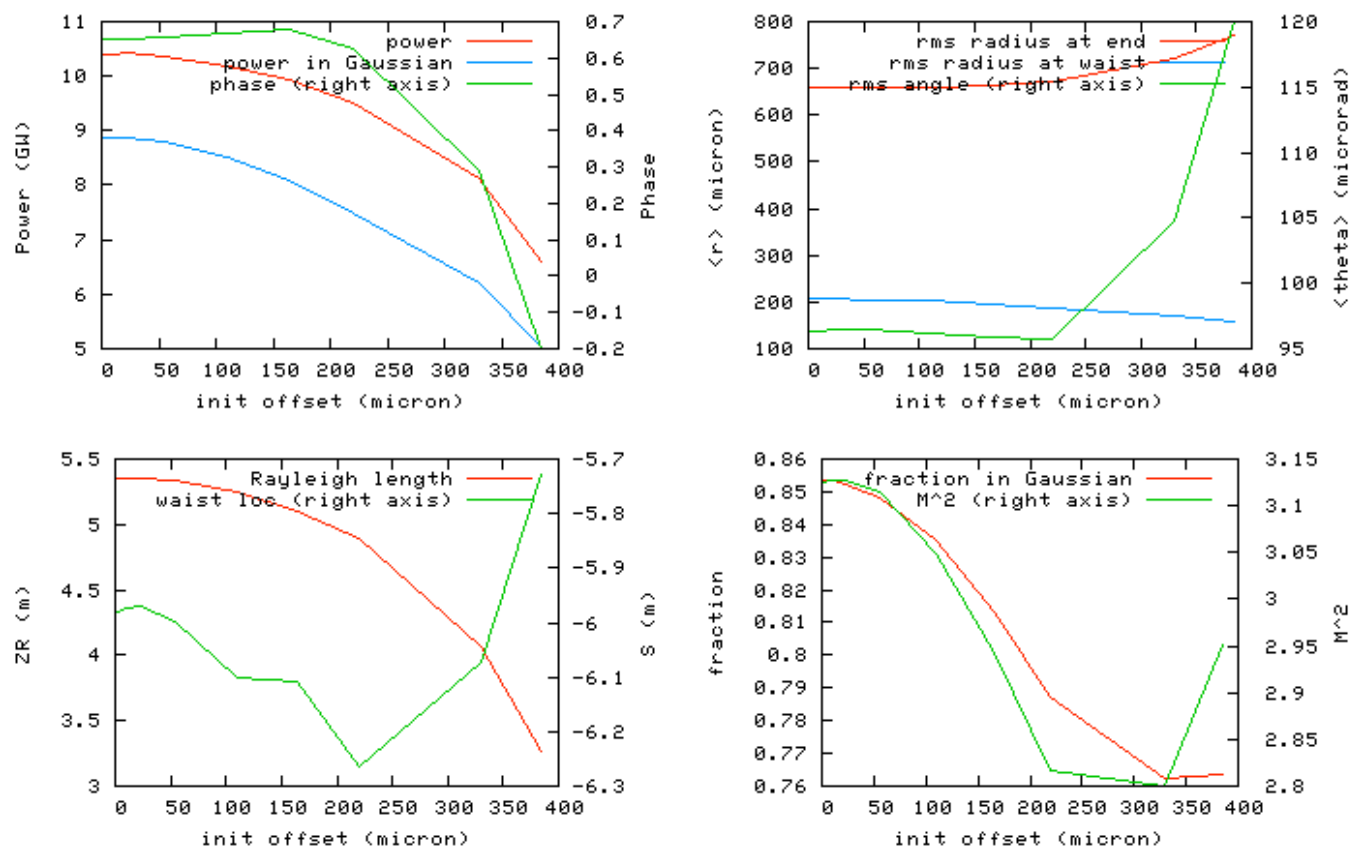


Figure 3.3-7: Sensitivity of output to input electron beam offset, tapered wiggler 100-nm case.

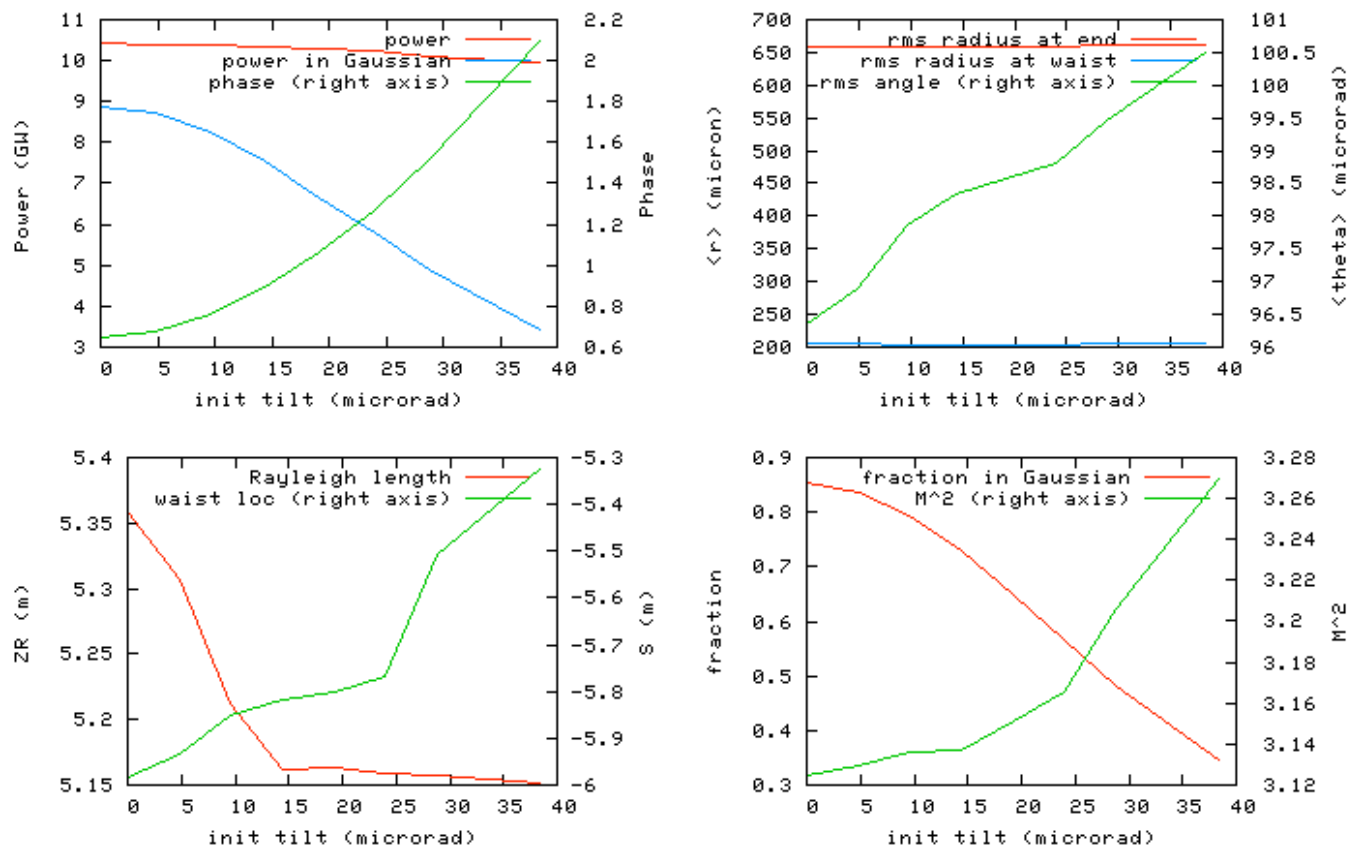


Figure 3.3-8: Sensitivity of output to input electron beam tilt, tapered wiggler 100-nm case.

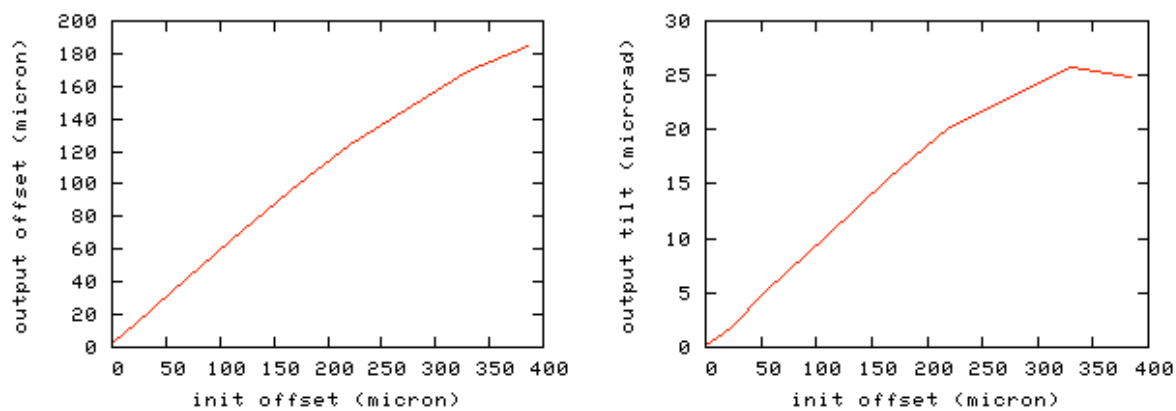


Figure 3.3-9: Predicted output radiation misalignment from as a function of initial electron beam offset, tapered wiggler 100-nm case.

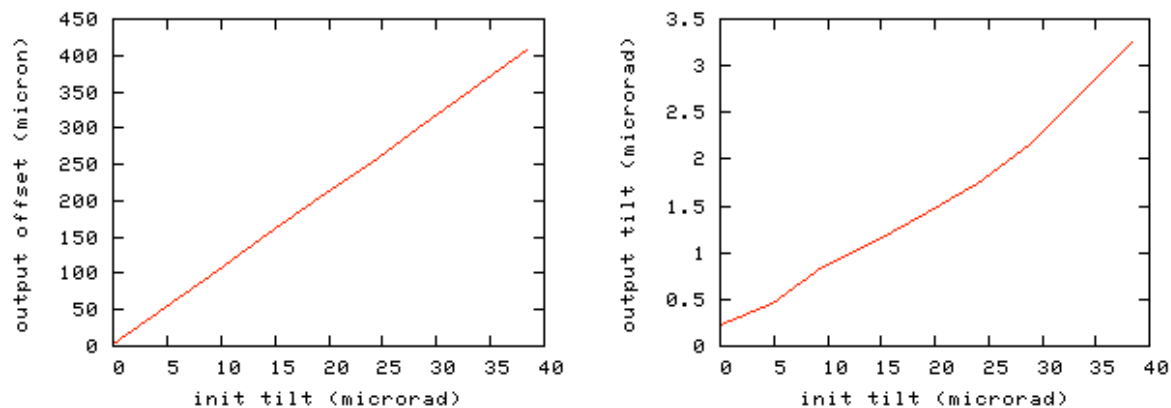


Figure 3.3-10: Output radiation misalignment from an electron beam initial tilt, tapered wiggler 100-nm case.

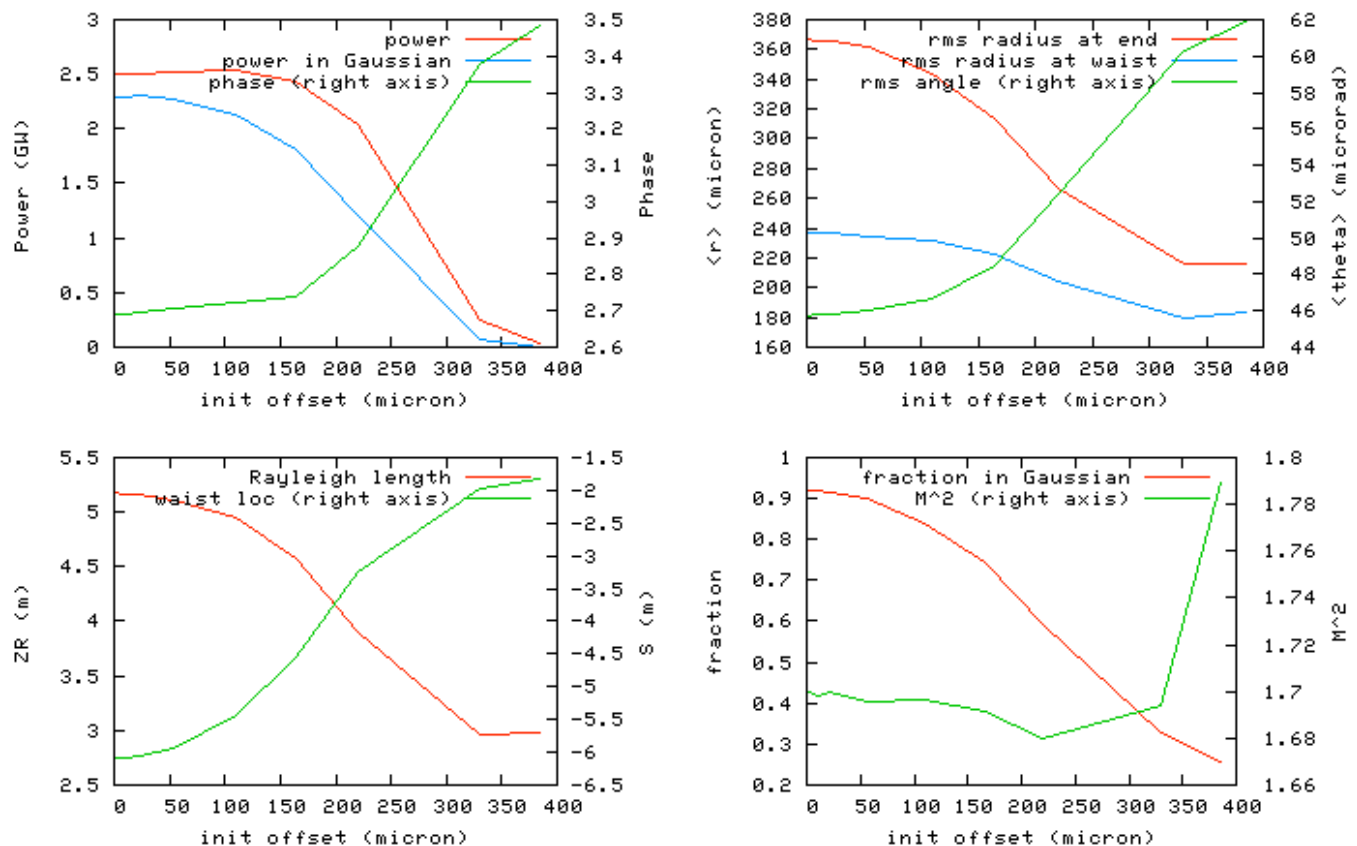


Figure 3.3-11: Sensitivity of 40-nm output mode characteristics to input electron beam offset.

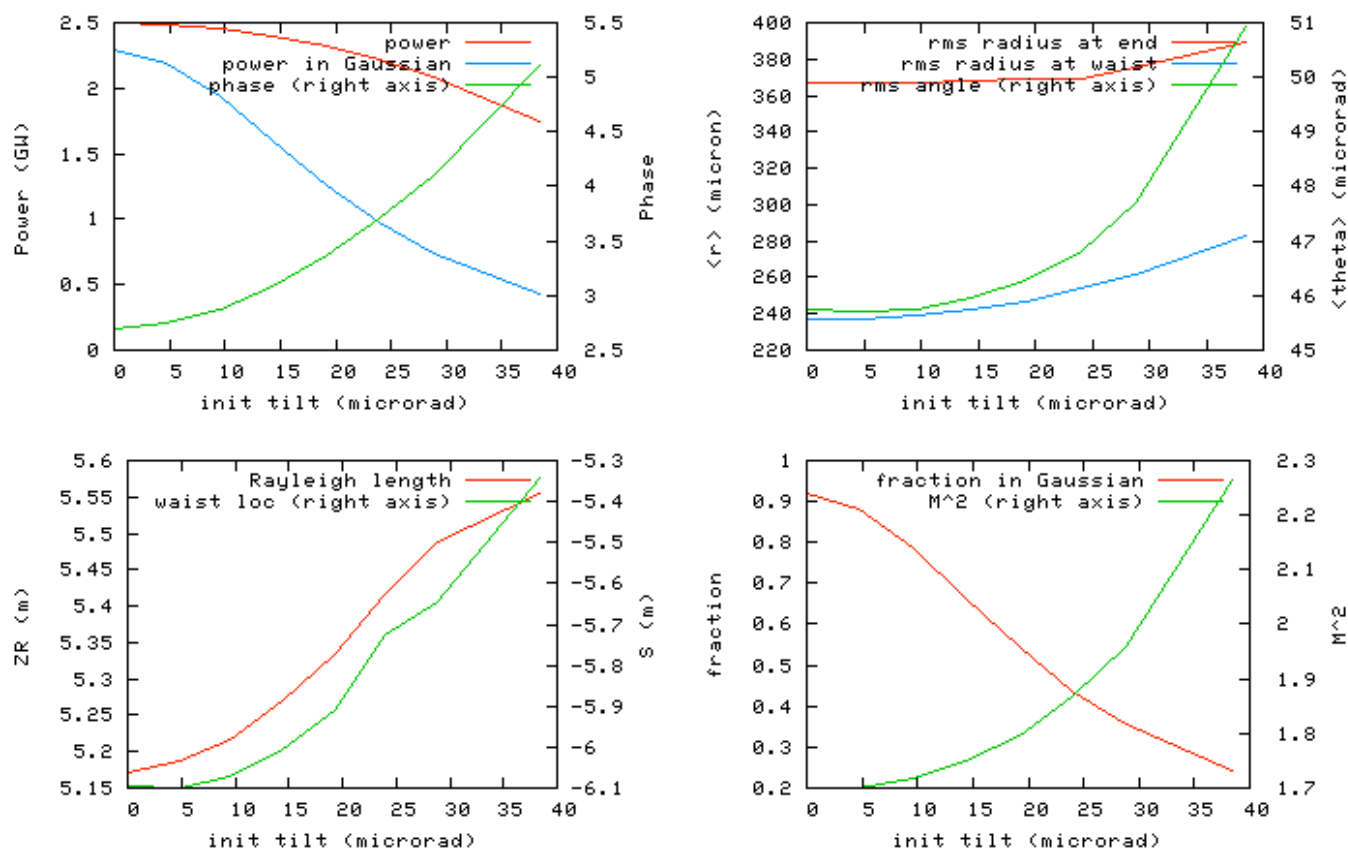


Figure 3.3-12: Sensitivity of 40-nm output mode characteristics to initial electron beam tilt.

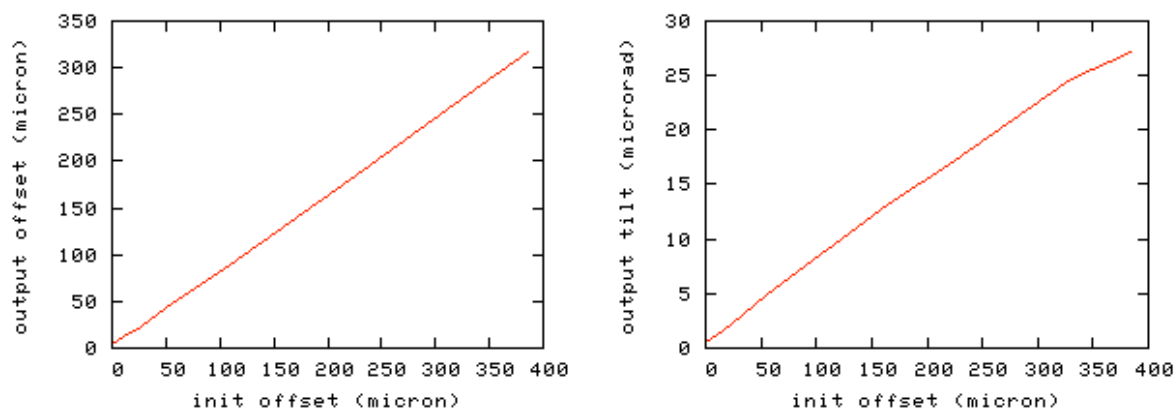


Figure 3.3-13: Predicted 40-nm output radiation misalignment as a function of initial electron beam offset.

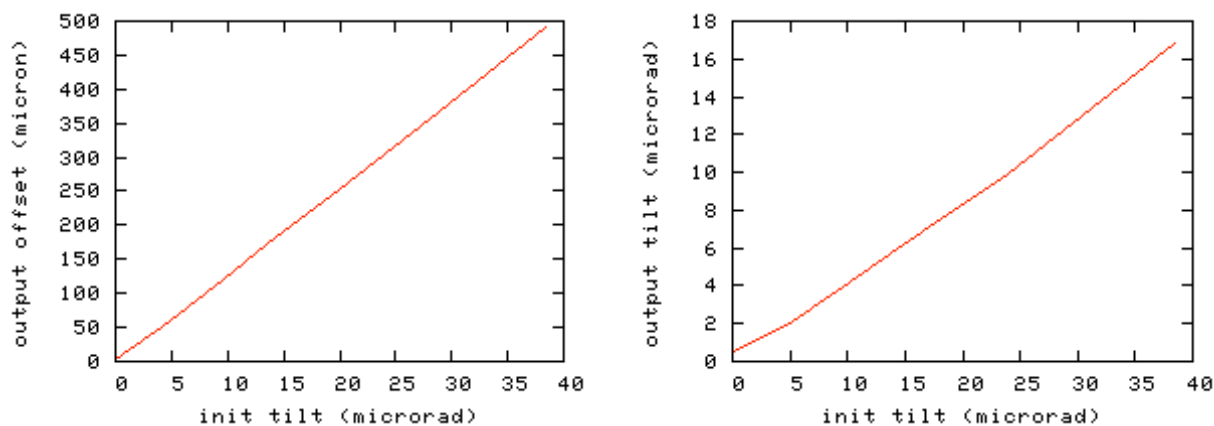


Figure 3.3-14: Predicted 40-nm output radiation misalignment from initial electron beam tilt.

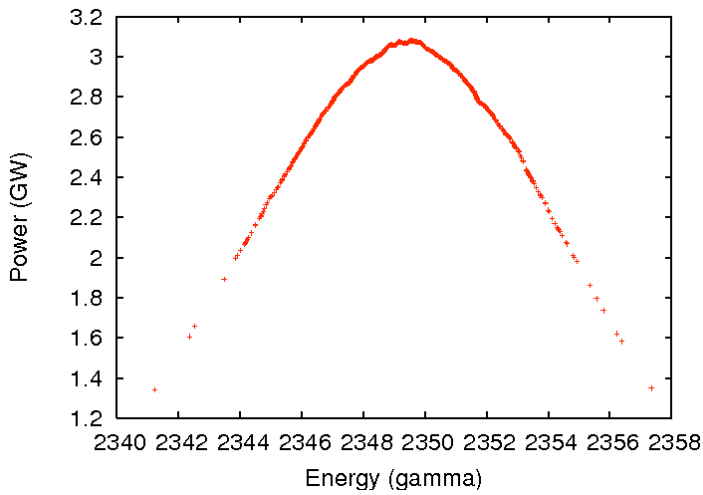


Figure 3.3-15: Output power versus electron beam energy at 100 nm.

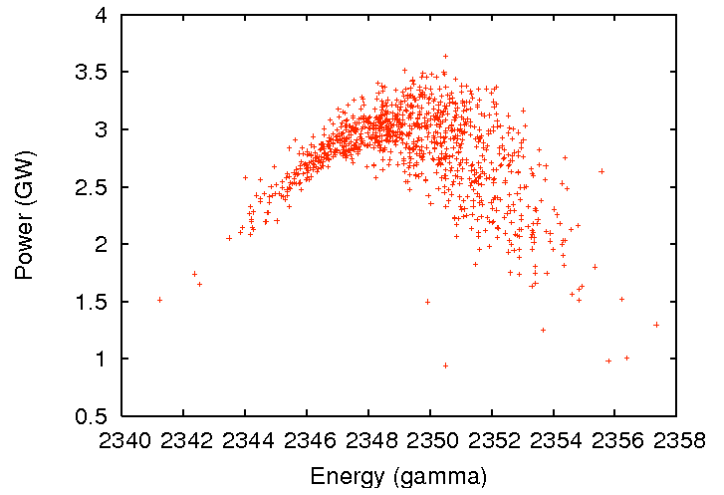


Figure 3.3-16: Output power versus energy at 100 nm for simultaneous variation of all beam parameters and the laser seed power following the distributions given in Table 3-4.

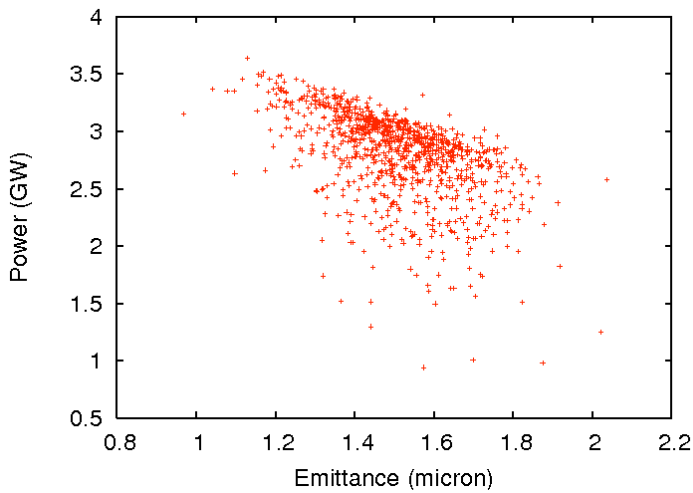


Figure 3.3-17: Output power versus emittance at 100 nm for the same cases plotted in Fig. 3.3-16

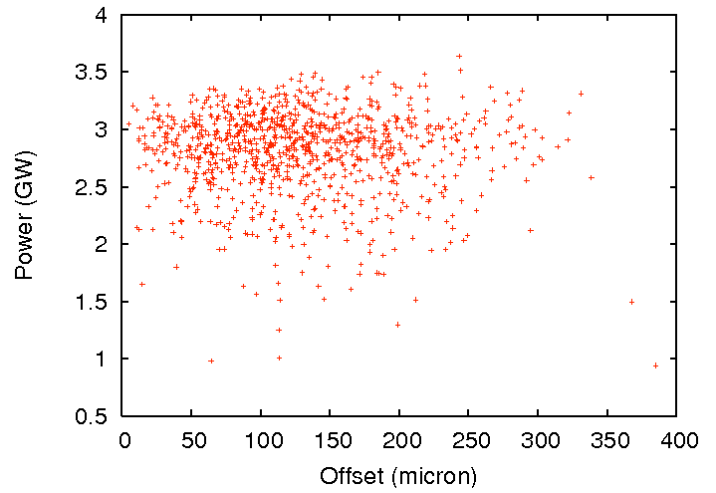


Figure 3.3-18: Output power versus position offset at 100 nm; same cases as presented in Figs. 3.3-16 and 3.3-17.

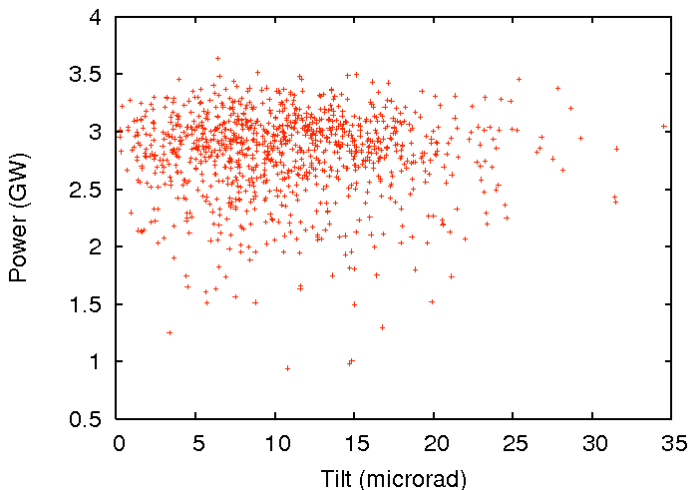


Figure 3.3-19: Output power versus initial tilt at 100 nm.

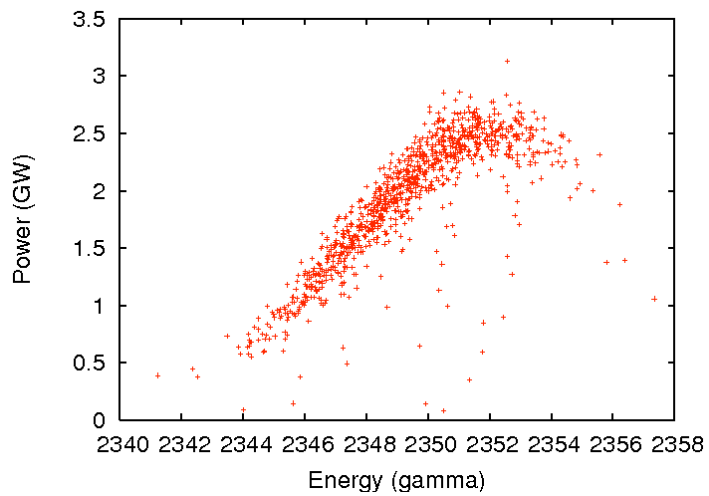


Figure 3.3-20: Output power versus energy at 40 nm with all input e-beam parameters and the laser seed power simultaneously varying as given in Table 3-4.

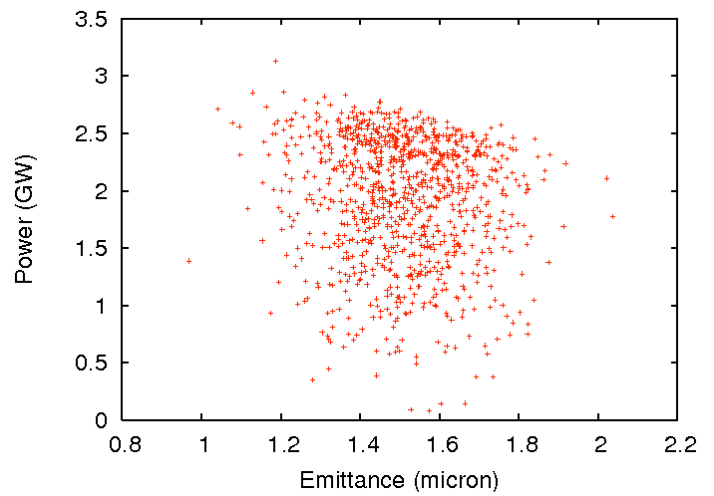


Figure 3.3-21: Output power versus emittance at 40 nm for the same runs as plotted in Fig. 3.3-20.

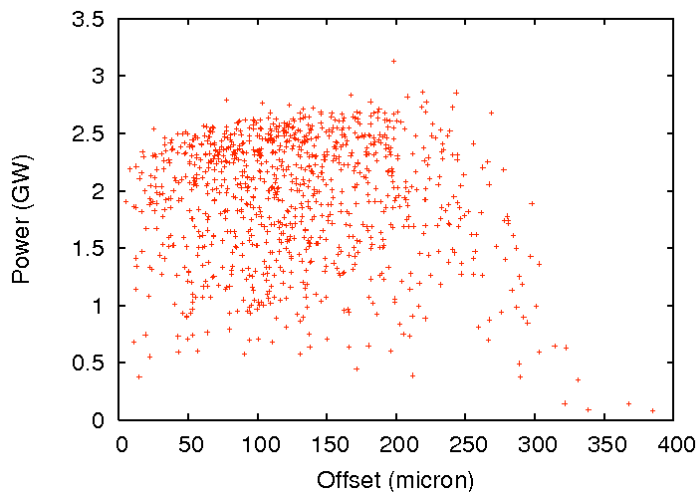


Figure 3.3-22: Output power versus initial transverse e-beam position at 40 nm for the same runs as plotted in Fig. 3.3-20.

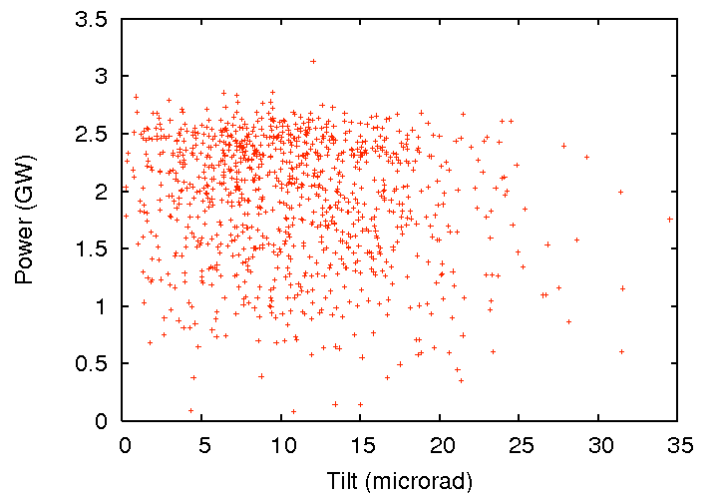


Figure 3.3-23: Output power versus initial transverse e-beam tilt at 40 nm for the same runs as plotted in Fig. 3.3-20.

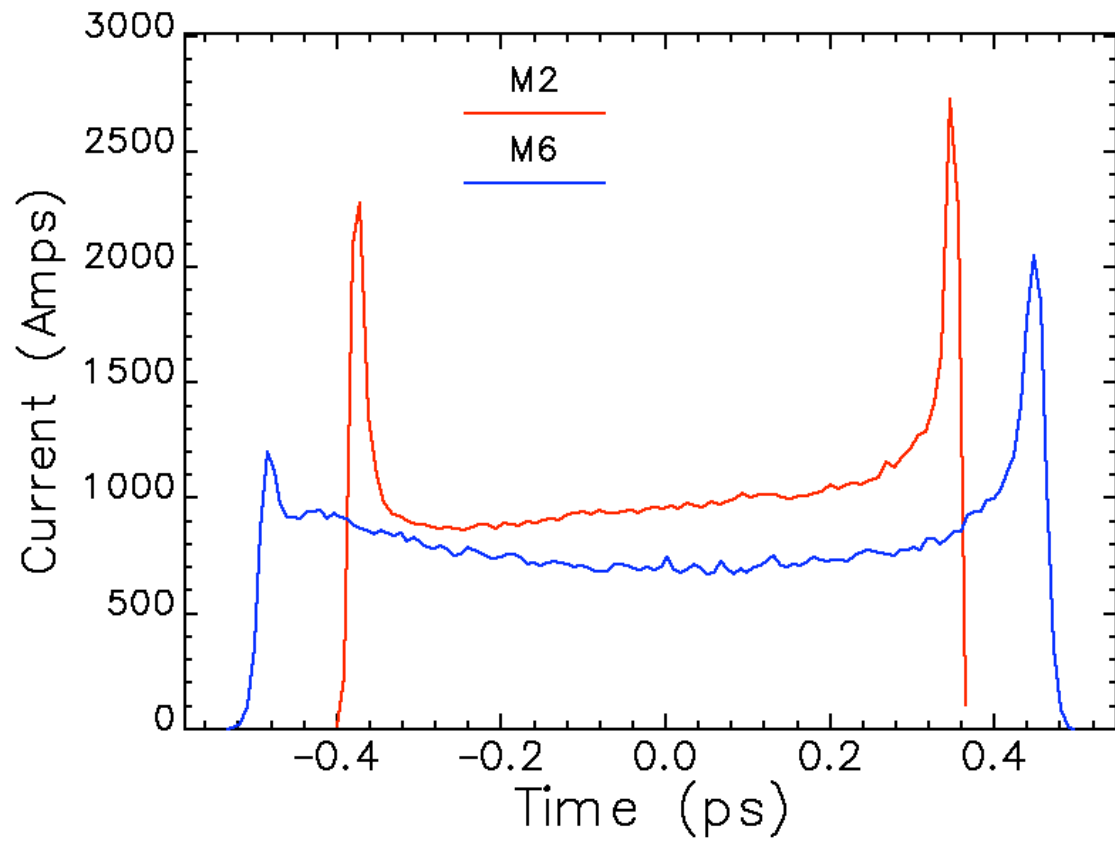


Figure 3.4-1: Current versus time profiles for the “M2” and “M6” medium-duration bunch distributions as computed by the injector and S2E groups.

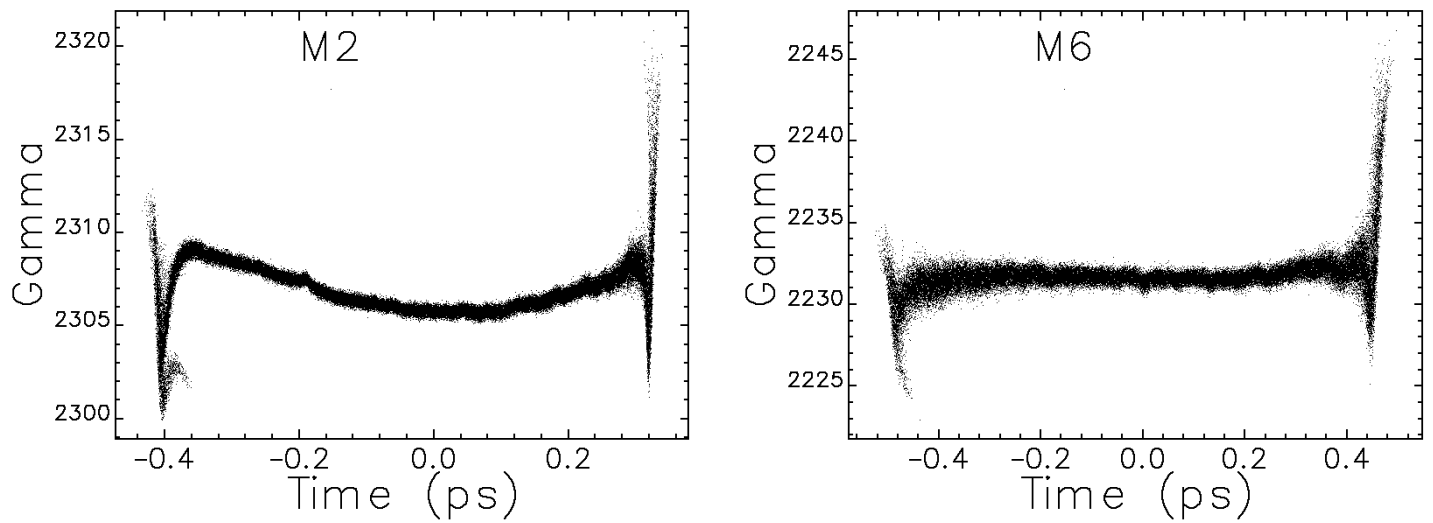


Figure 3.4-2: Longitudinal phase space scatter plots for the “M2” and “M6” medium-duration bunch distributions as computed by the injector and S2E groups.

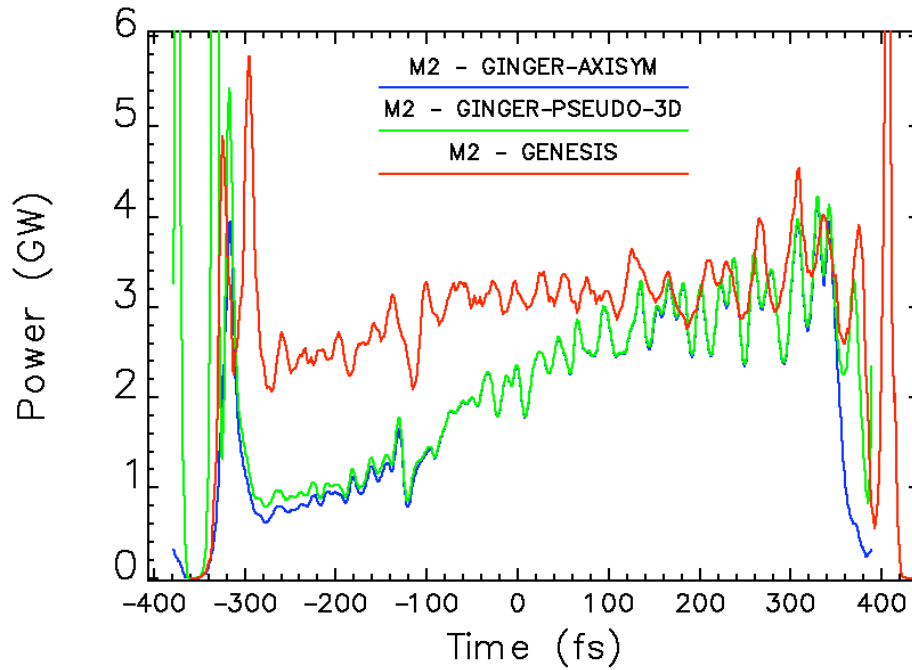


Figure 3.4-3: Time-dependent 40-nm output power results as computed by GINGER and GENESIS for the M2 S2E distributions in which a temporally constant input power laser seed covered the full electron beam pulse. We show a both a normal case (“GINGER-AXISYM”) in which the normal axisymmetric field solver of GINGER is used and a second case (PSEUDO-3D) where the radiation field is forced *artificially* to have the same instantaneous centroid position as the e-beam.

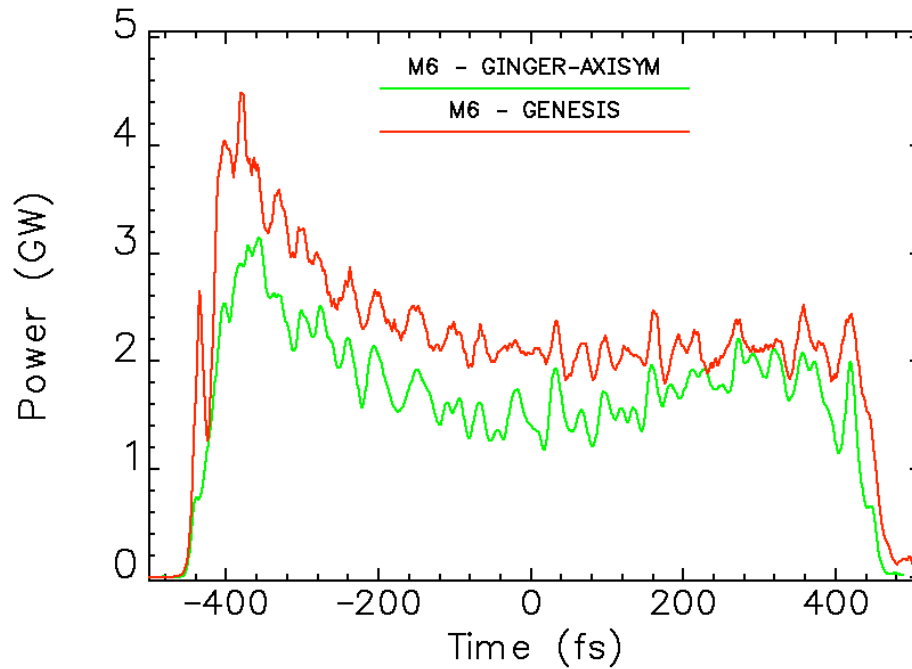


Figure 3.4-4: Time-dependent 40-nm output power results as computed by GINGER and GENESIS for the M6 S2E distributions in which a temporally constant input power laser seed covered the full electron beam pulse.

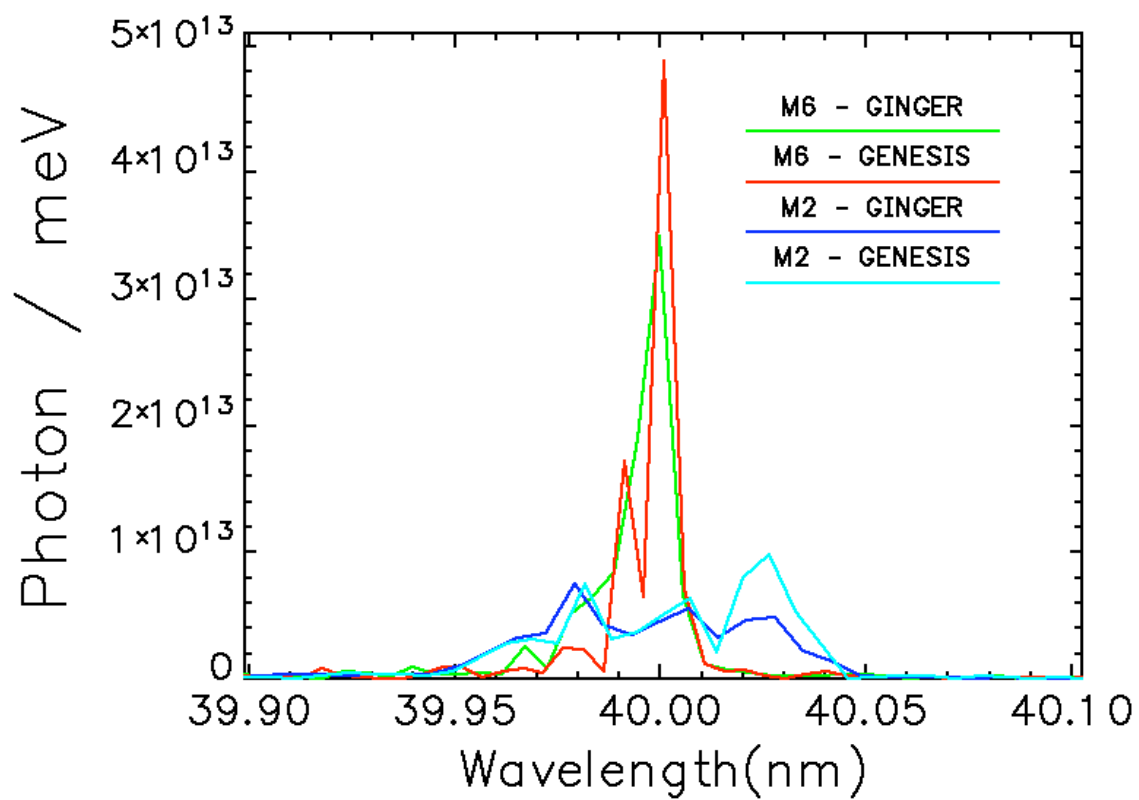


Figure 3.4-5: Spectral intensity (photons/meV/pulse) from GINGER and GENESIS time-dependent results for FEL-1 at 40-nm output wavelength for the M2 and M6 S2E medium pulse distributions. The spectral resolution in both codes was ~ 5 meV.

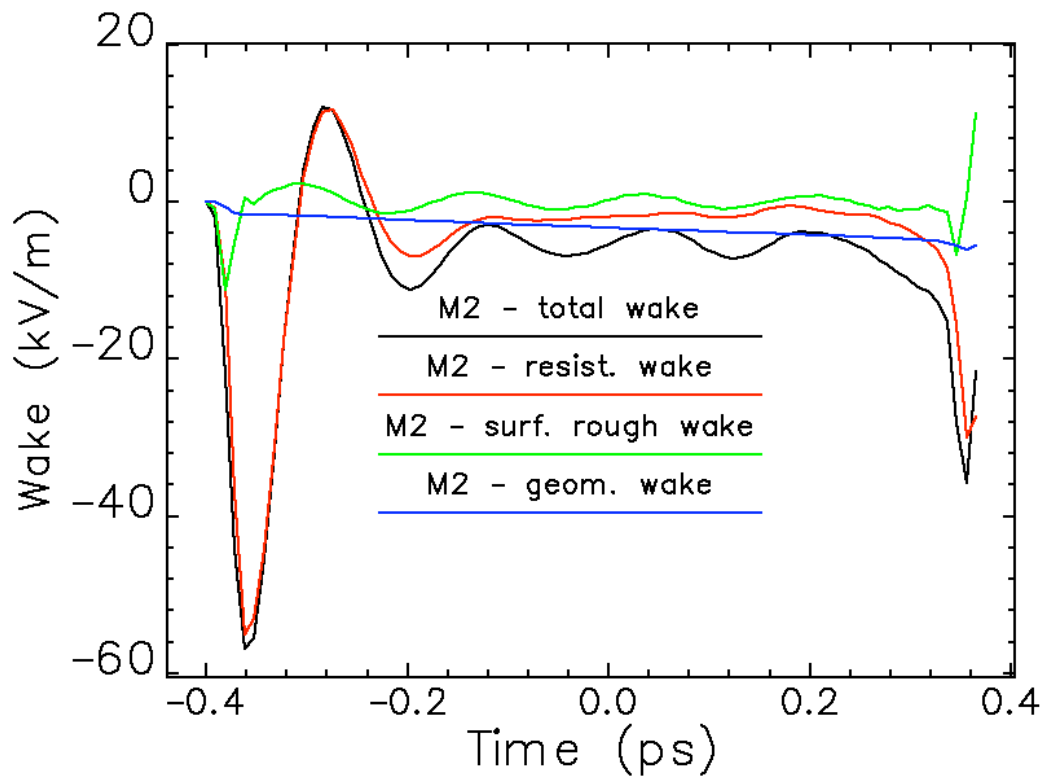


Figure 3.4-6: Time-dependent longitudinal wake results for the M2 S2E distribution. In addition to the total wake (black line), the individual components of the resistive, surface roughness, and geometric cross-section interruption wakes are also plotted.

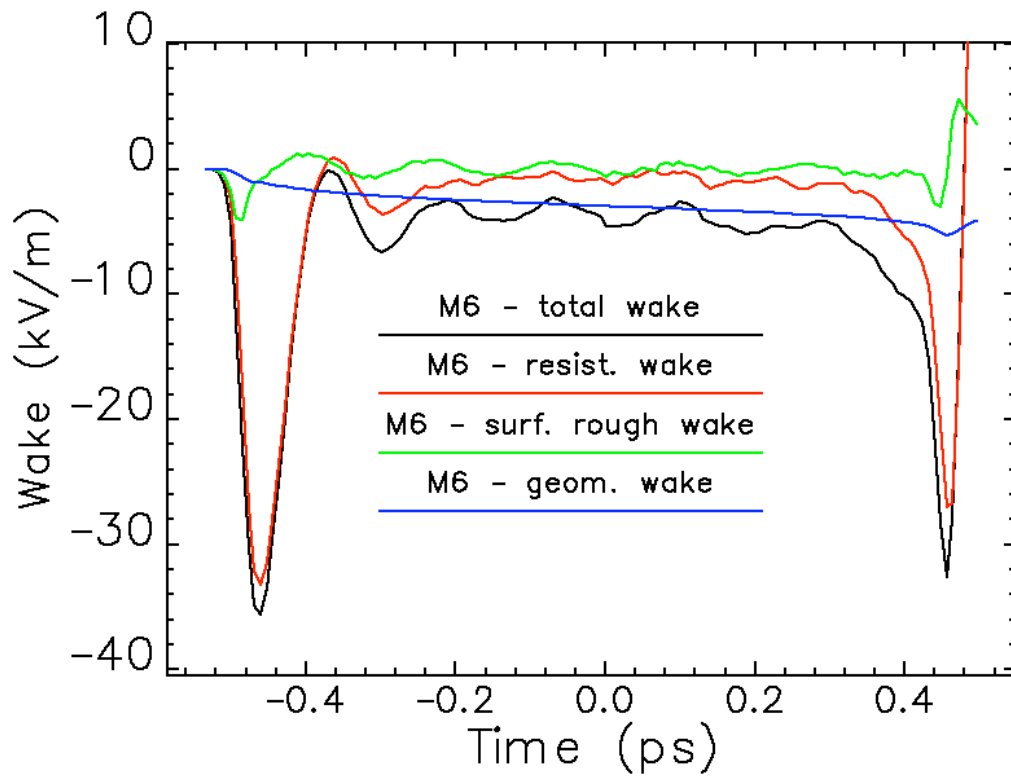


Figure 3.4-7: Time-dependent longitudinal wake results for the M6 S2E distribution. In addition to the total wake (black line), the individual components of the resistive, surface roughness, and geometric cross-section interruption wakes are also plotted.

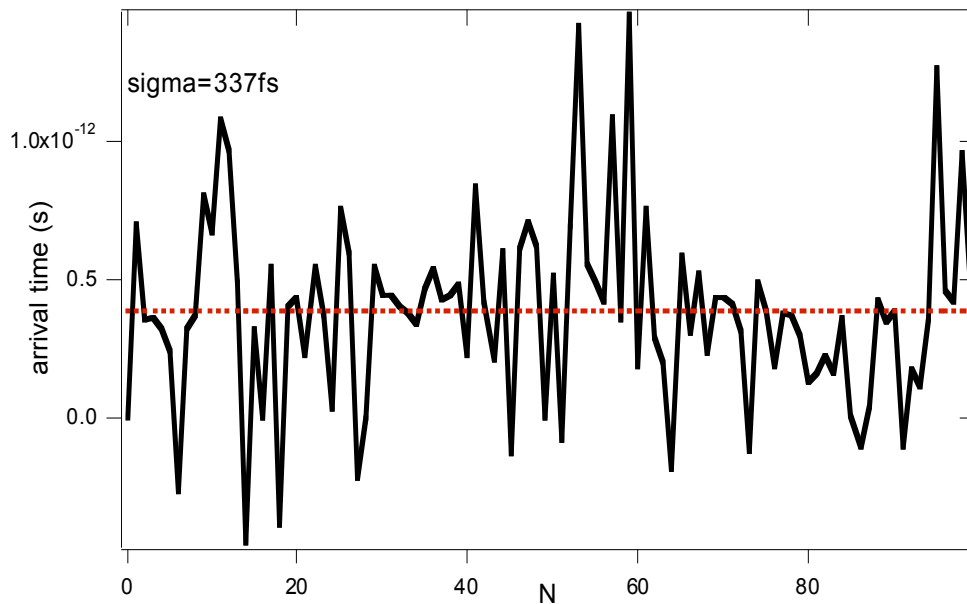
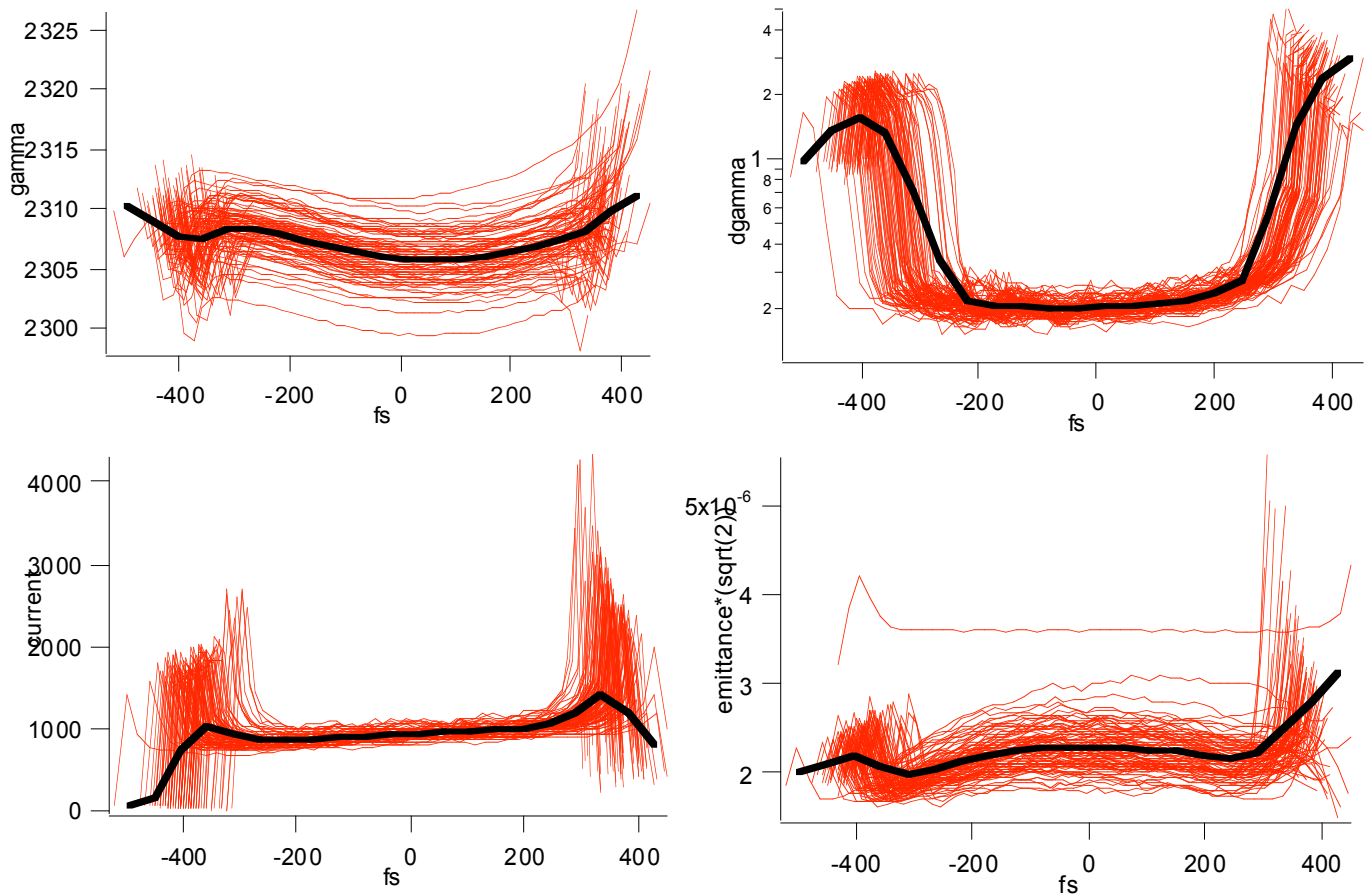


Figure 3.5-1: Mean arrival times for the macroparticles for each of the individual files produced by the injector and linac groups for the medium bunch case for FEL-1. Each file has its own completely independent set of random accelerator cell voltage and timings.



Figures 3.5-2 - 3.5-5: Plots of envelope quantities versus times for the 100 medium bunch jitter files. The time coordinate for each file has been re-centered at the mean arrival time of the individual macroparticles. The black line is the time-resolved average for each envelope quantity over all the files.

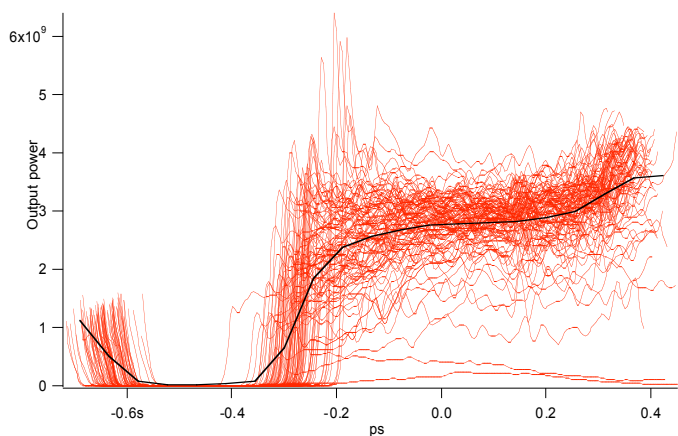


Figure 3.5-6: Plots of instantaneous output power at 40-nm wavelength for the 100 jitter files; the black line is the average power over all the files.

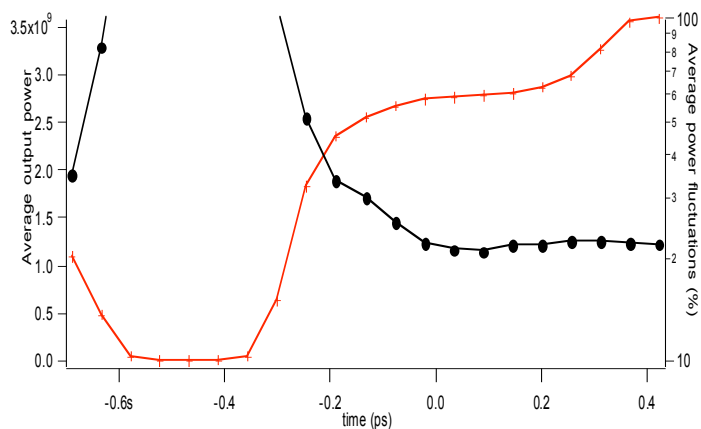


Figure 3.5-7: Average output power (red line) and standard deviation as a function of time for the runs plotted in Fig. 3.5-6..

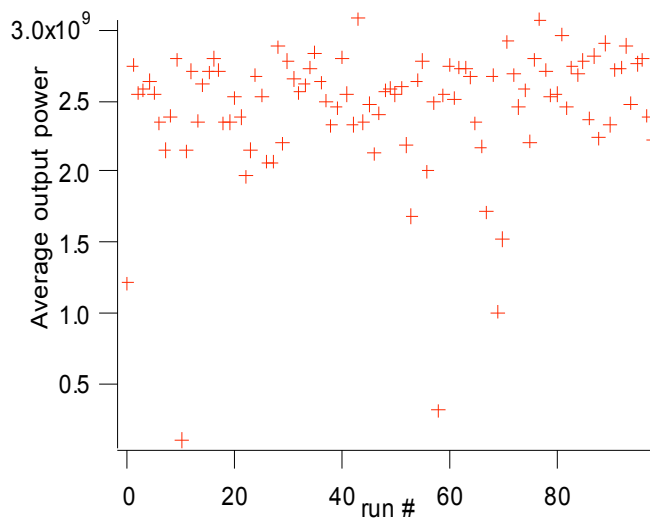


Figure 3.5-8: Average output power over a time window of ± 350 fs as predicted by GINGER for the 100 jitter files plotted in the previous figures.

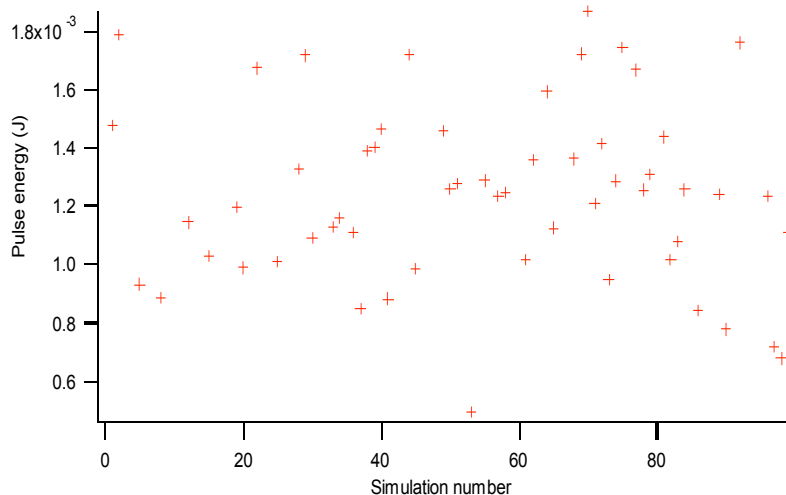


Figure 3.5-9: Average output energy over a time window of ± 300 fs as predicted by GENESIS for the 100 jitter files.

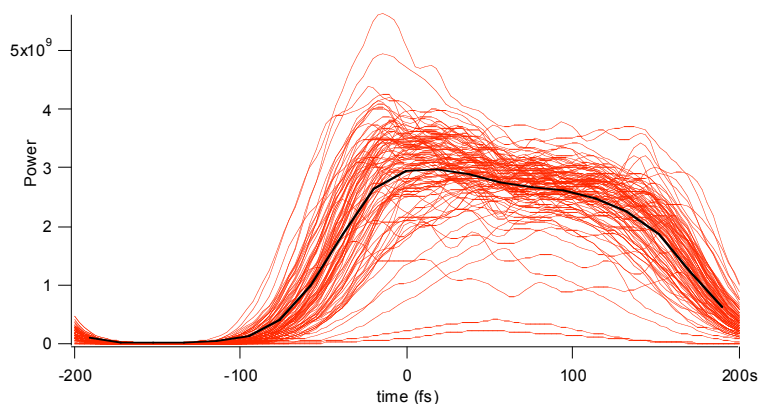


Figure 3.5-10: Instantaneous output power as predicted by GINGER for the 100 jitter files for which the input radiation seed was a 100-fs Gaussian pulse.

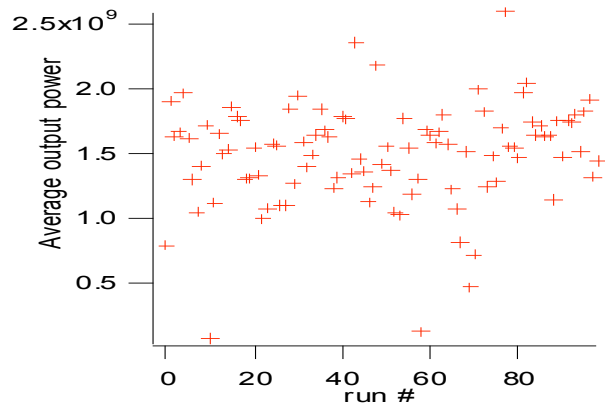


Figure 3.5-11: Average output power over a time window of ± 100 fs using the results of Fig. 3.5-10.

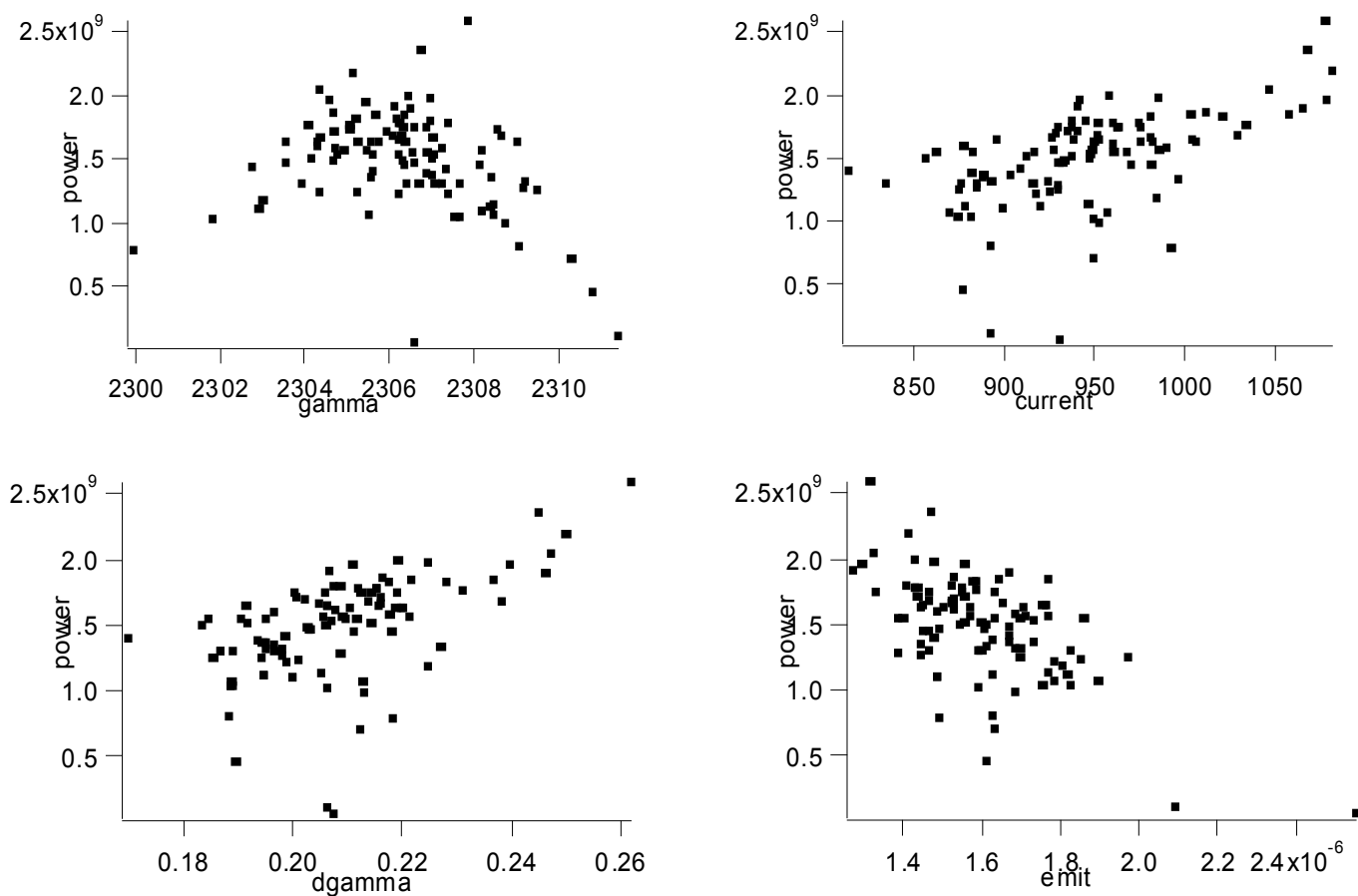


Figure 3.5-12: Scatterplots of average output power over a time window of ± 100 fs versus different electron beam parameters time-averaged over a time window $[-200$ fs, $+200$ fs] using the simulation results plotted in Fig. 3.5-11.

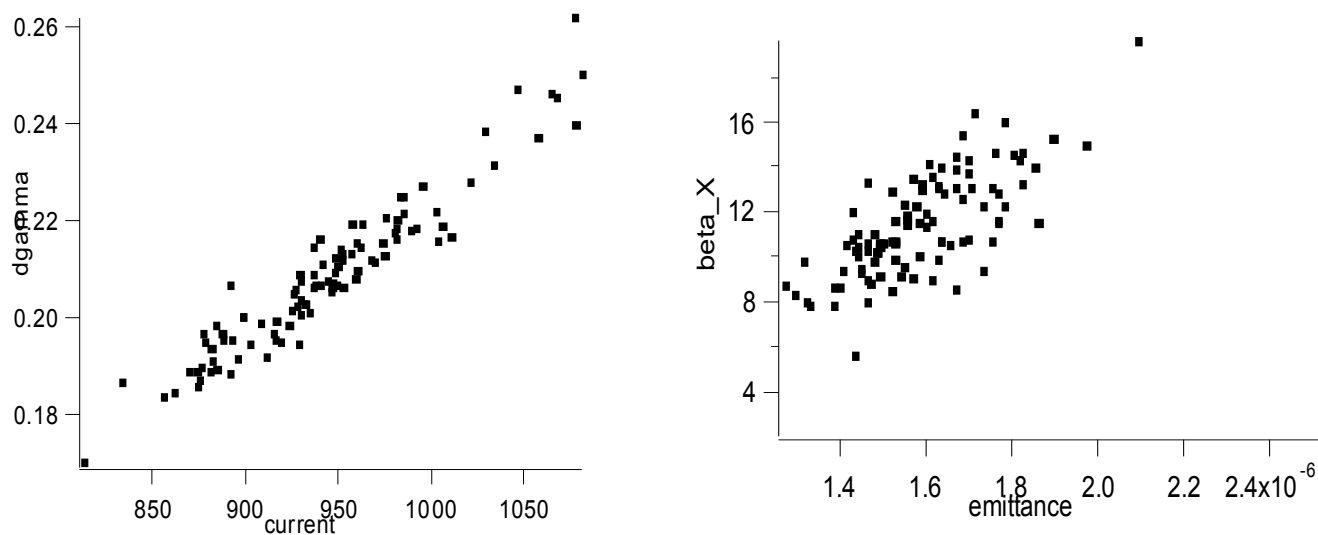


Figure 3.5-13: Scatterplots of showing the correlations in the 100 jitter files of incoherent energy spread with current and the β_x Twiss function with emittance; all quantities are averaged over a time window $[-200$ fs, $+200$ fs].

FEL-2 (40-10 nm): Fresh-Bunch Configuration

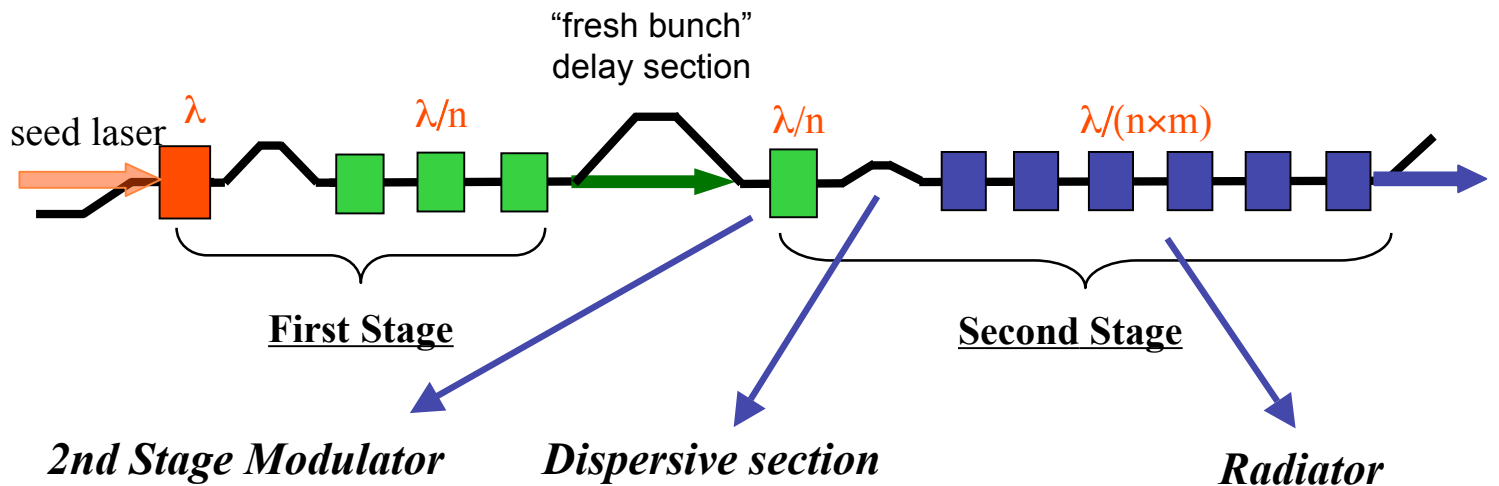


Figure 4.1-1: Schematic diagram of layout for the fresh bunch approach to FEL-2

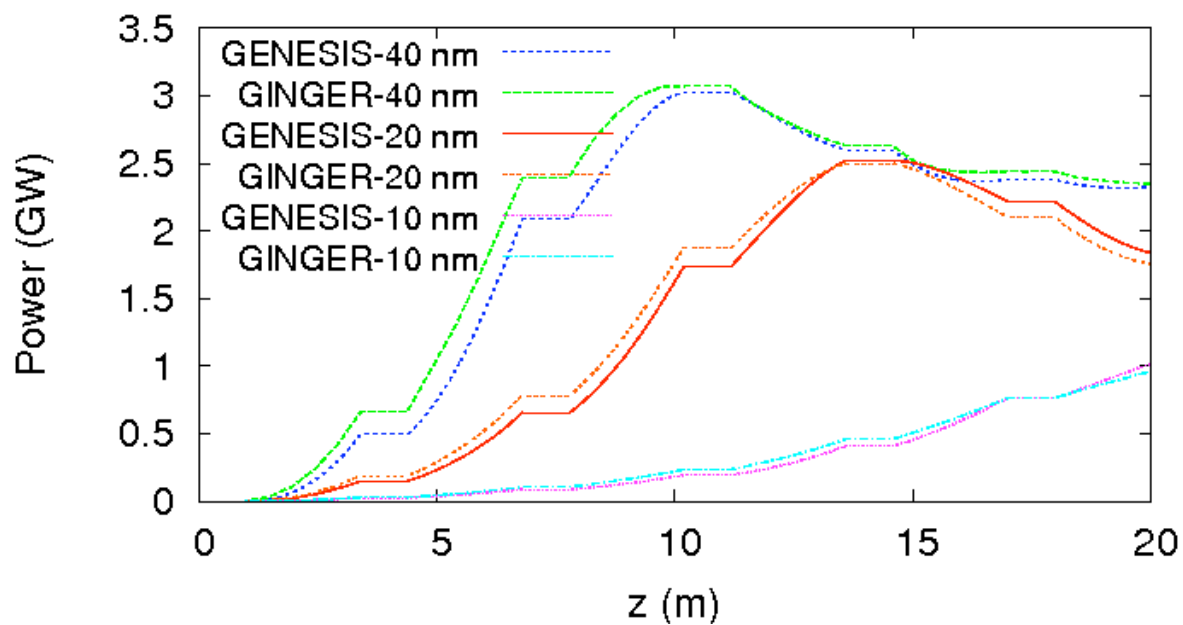


Figure 4.2-1: GENESIS and GINGER time-independent simulations for the FEL-2 radiation power at various wavelengths versus z in the final radiator for the nominal fresh bunch approach. These runs were done with 800-A beam current and 200-keV initial incoherent energy spread.

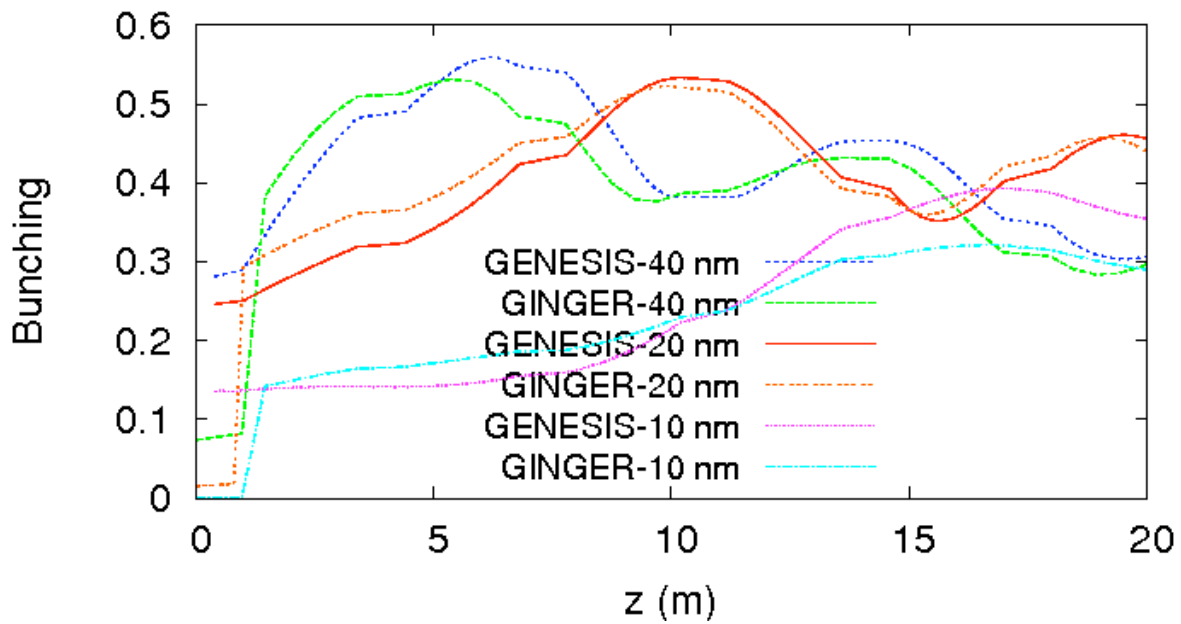


Figure 4.2-2: GENESIS and GINGER time-independent simulations for coherent microbunching at various wavelength versus z in the final radiator for the nominal fresh bunch approach for FEL-2. The beam current was 800 A in this example.

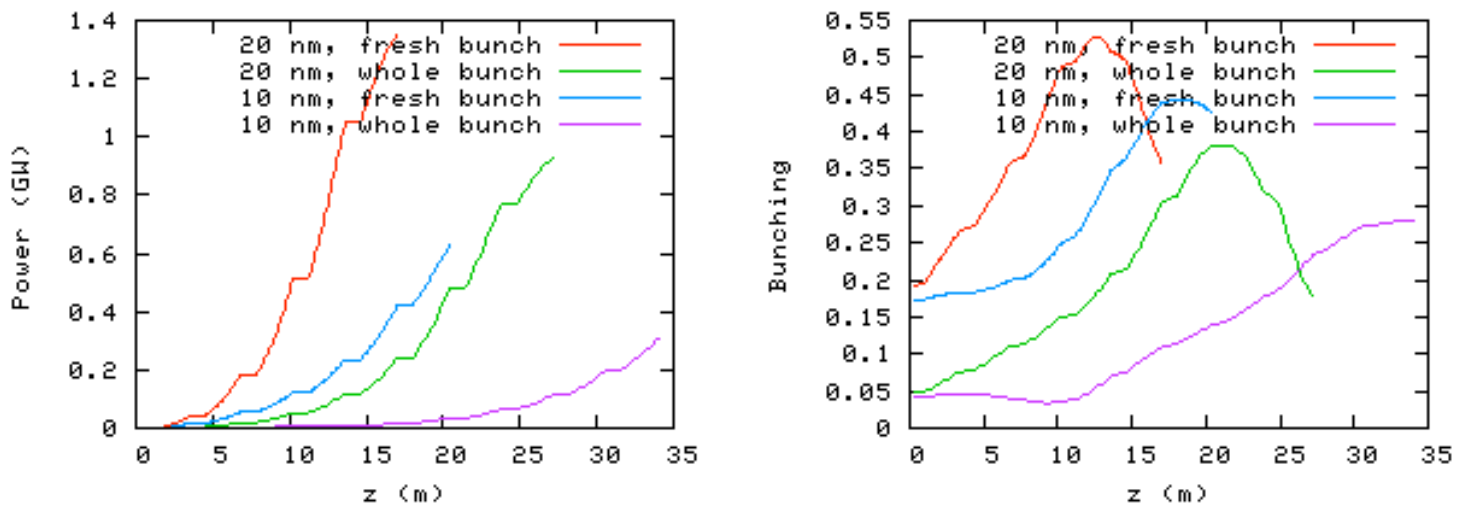


Figure 4.2-3: Power and bunching for different FEL-2 configurations as a function of longitudinal position within the radiator.

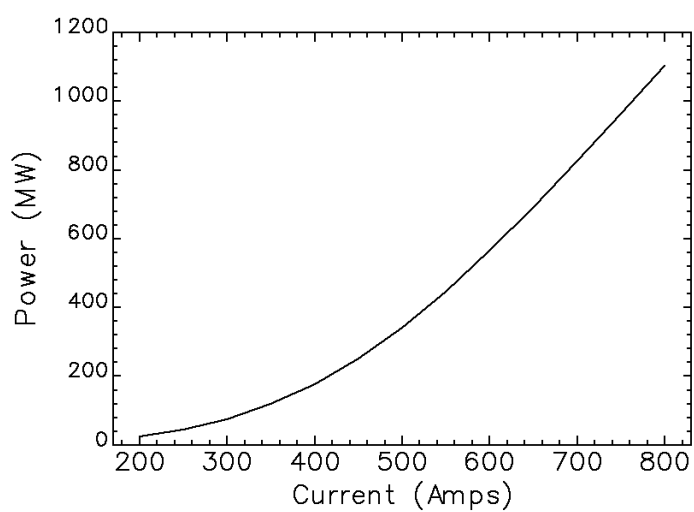
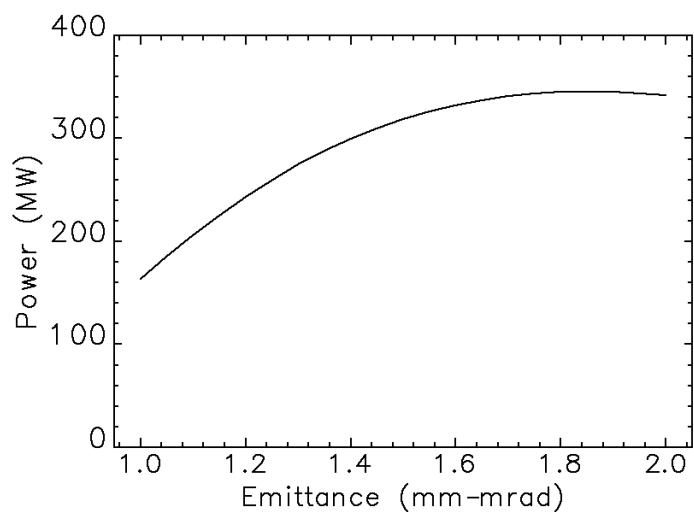
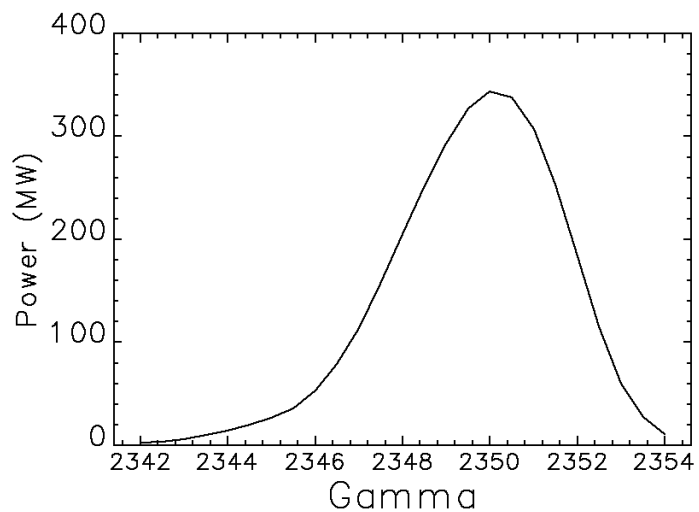
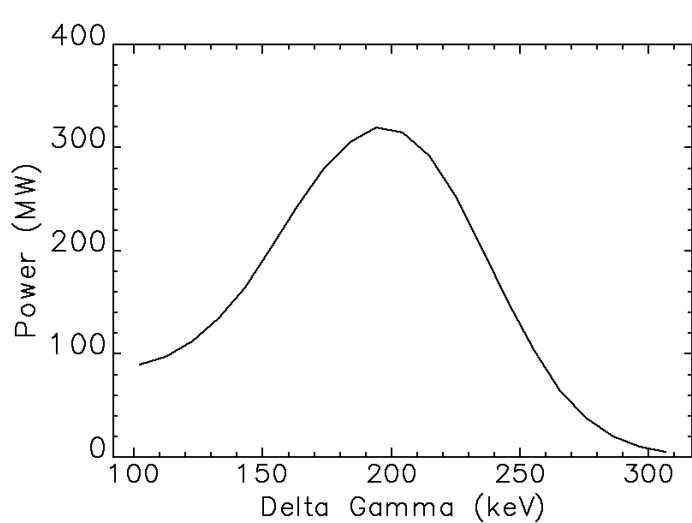


Figure 4.2-4: GINGER time-independent scan of FEL-2 (fresh bunch) output versus individual variation of various input electron beam parameters.

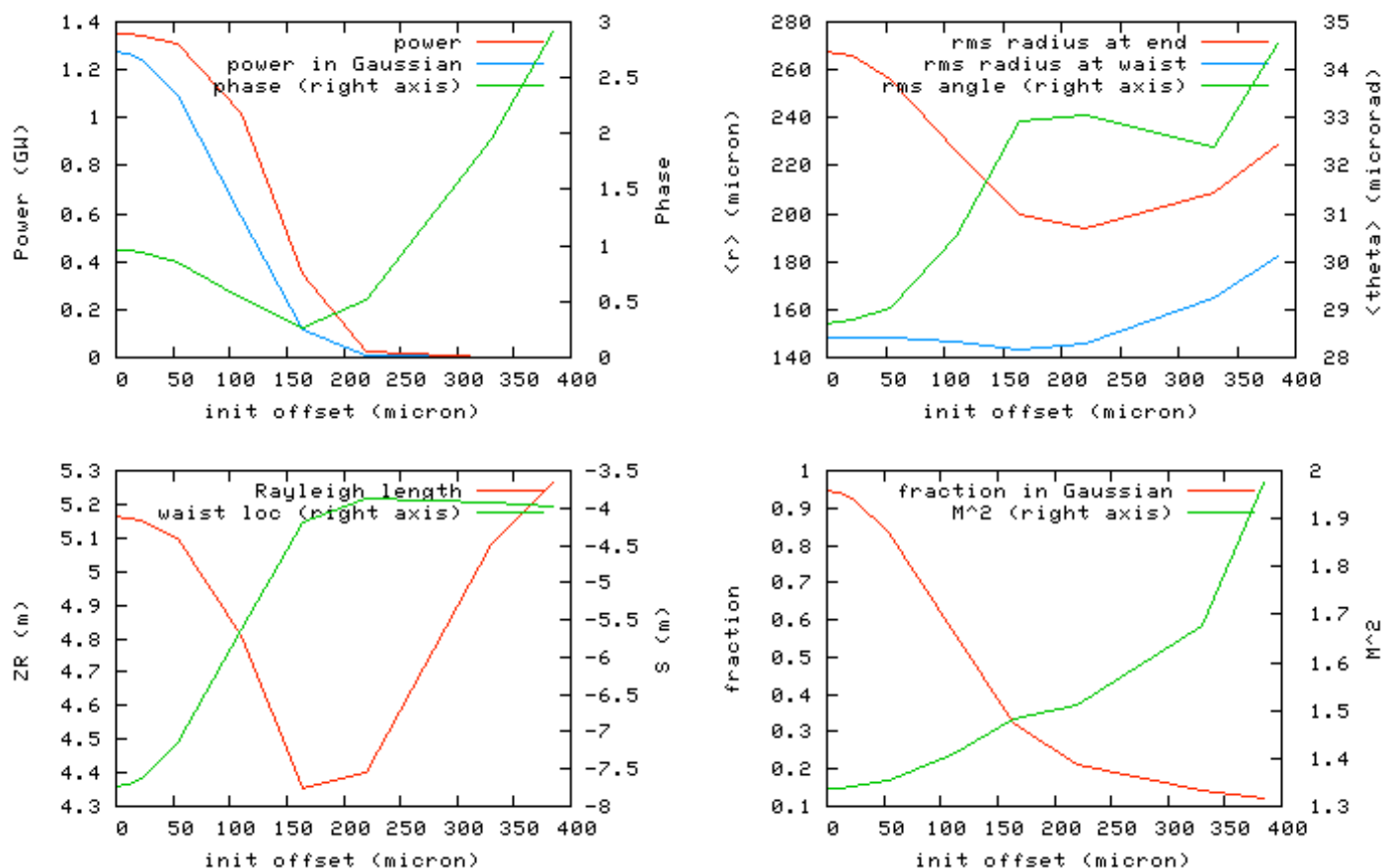


Figure 4.3-1: Sensitivity of final 20-nm power output to initial electron beam transverse offset for the FEL-2 fresh-bunch configuration.

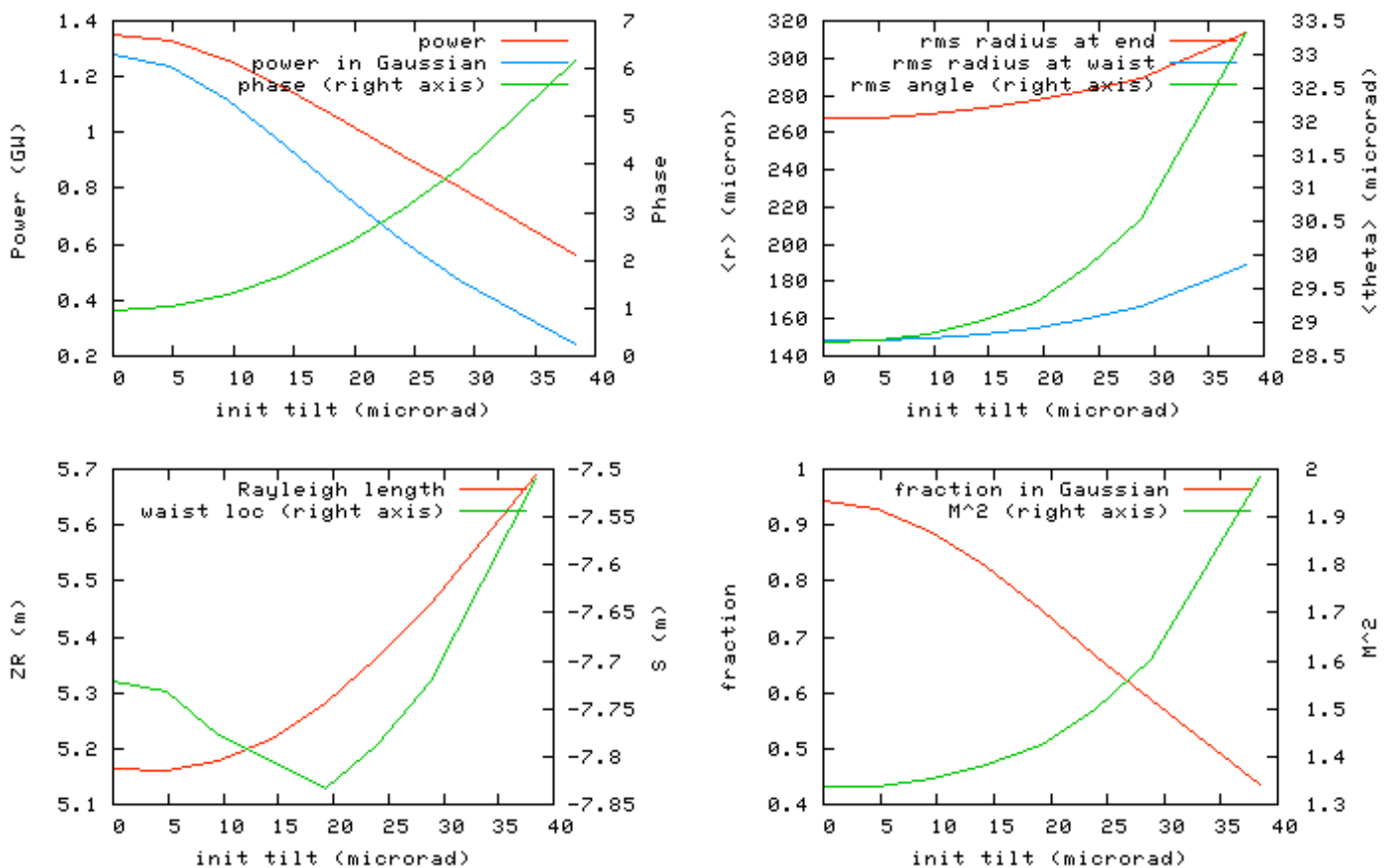


Figure 4.3-2: Sensitivity of final 20-nm power output to initial electron beam tilt for the FEL-2 fresh-bunch configuration.

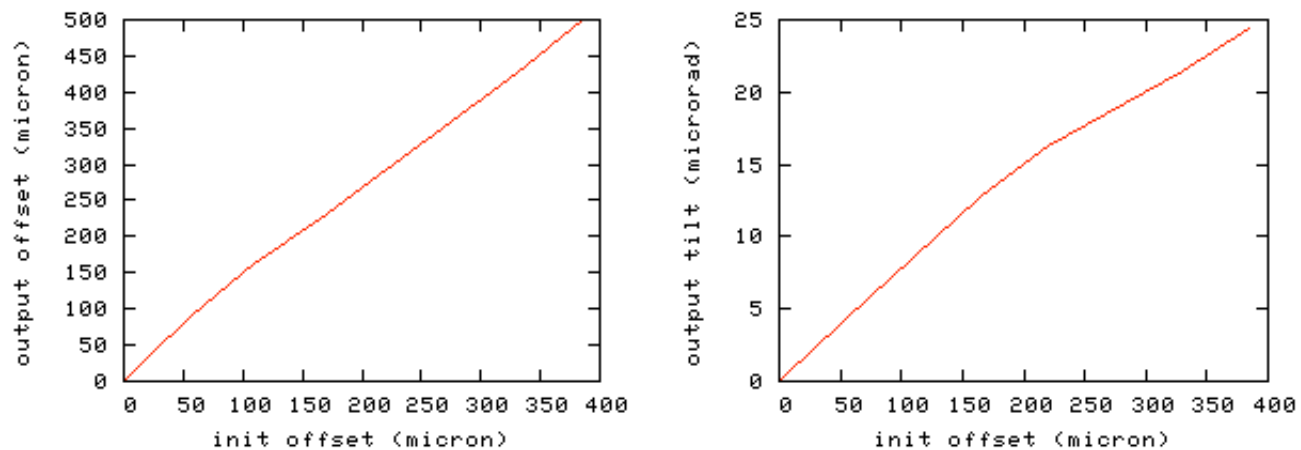


Figure 4.3-3: 20-nm output radiation misalignment as a function of initial electron beam offset in the fresh-bunch approach to FEL-2.

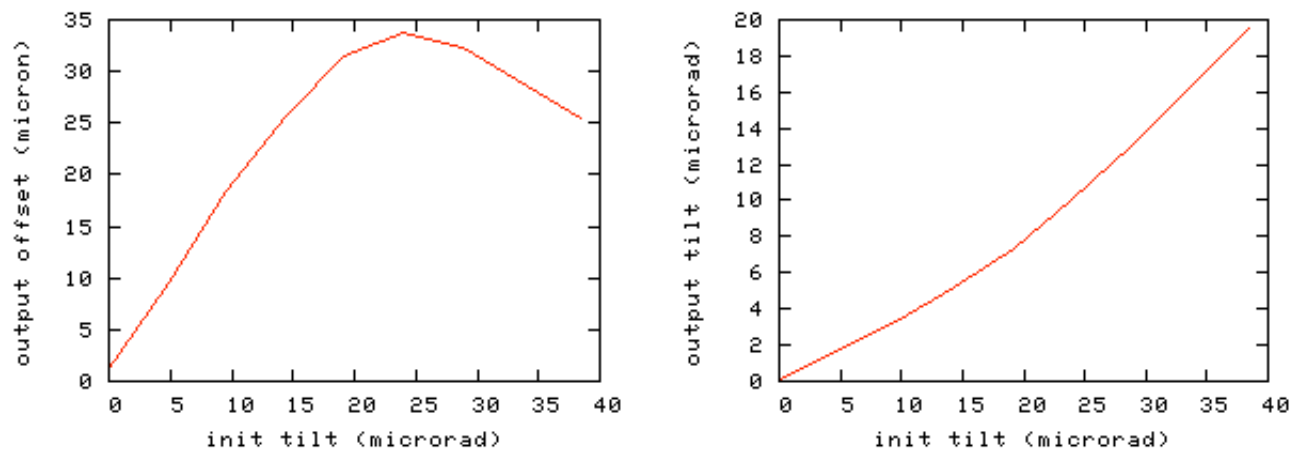


Figure 4.3-4: 20-nm output radiation misalignment as a function of initial electron beam tilt for the fresh-bunch approach to FEL-2.

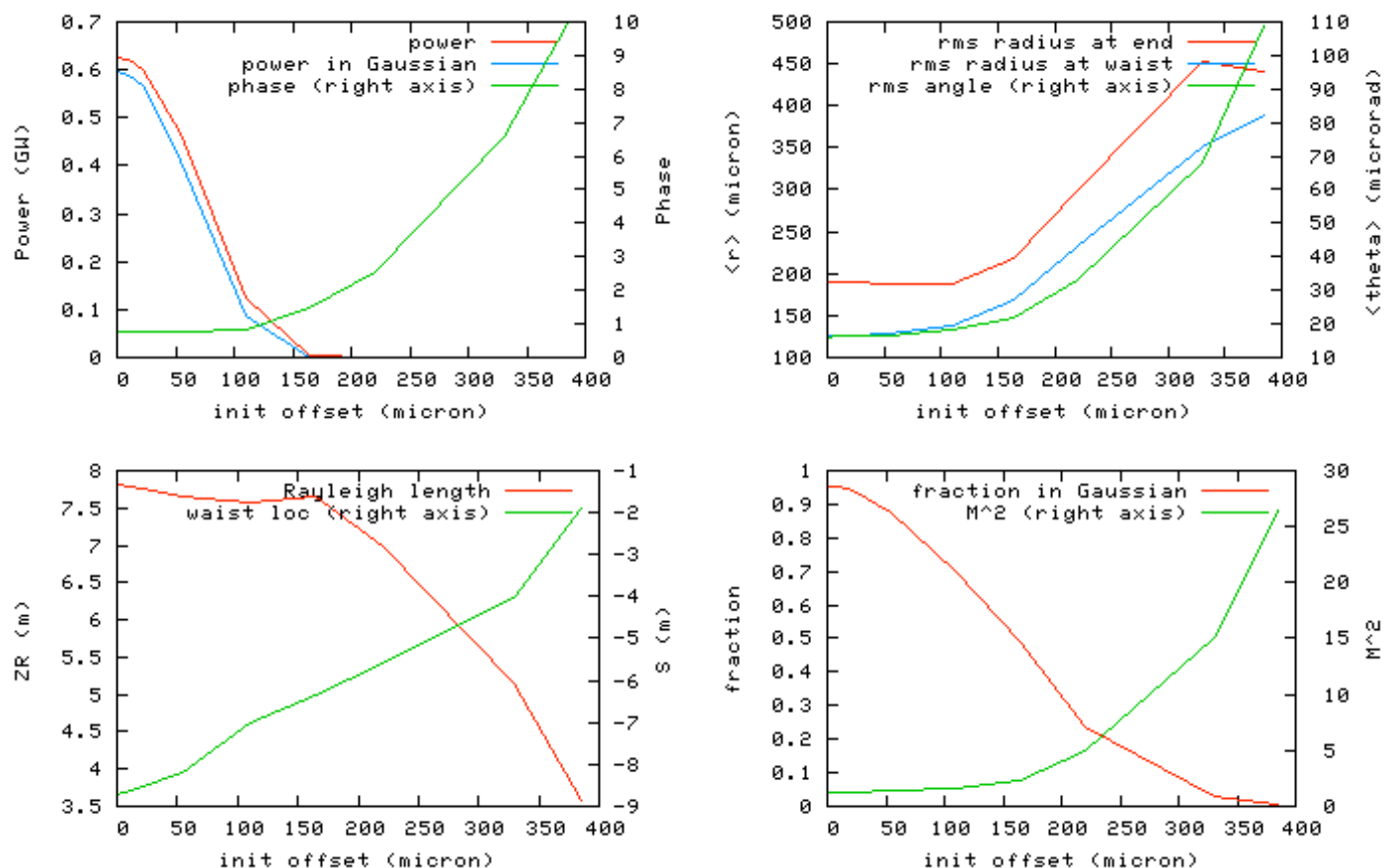


Figure 4.3-5: Sensitivity of 10-nm power output to initial electron beam offset for the fresh-bunch approach for FEL-2.

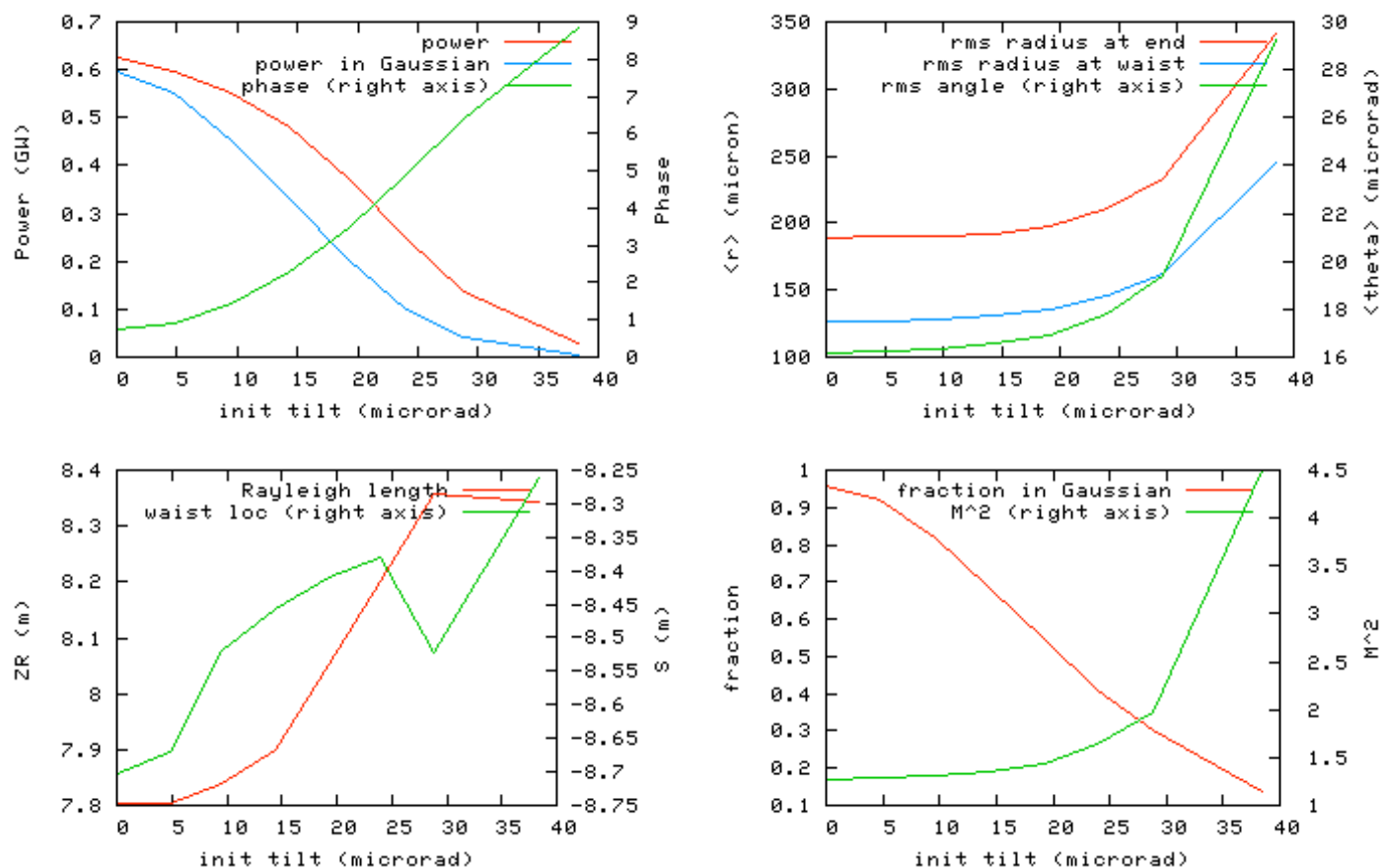


Figure 4.3-6: Sensitivity of 10-nm power output to initial electron beam tilt for the fresh-bunch approach for FEL-2.

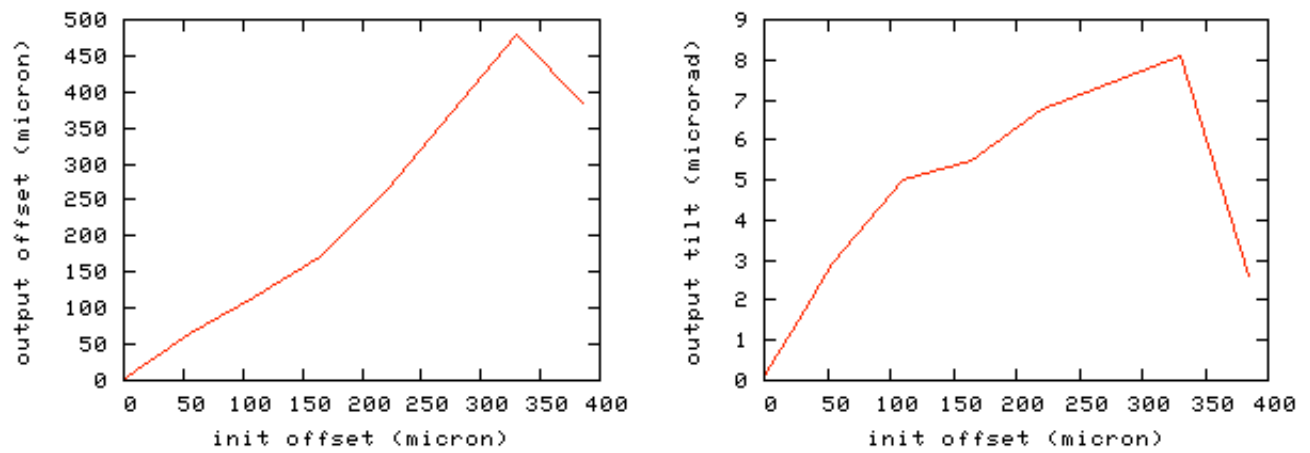


Figure 4.3-7: Output radiation misalignment as a function of initial e-beam offset for the 10-nm fresh-bunch case.

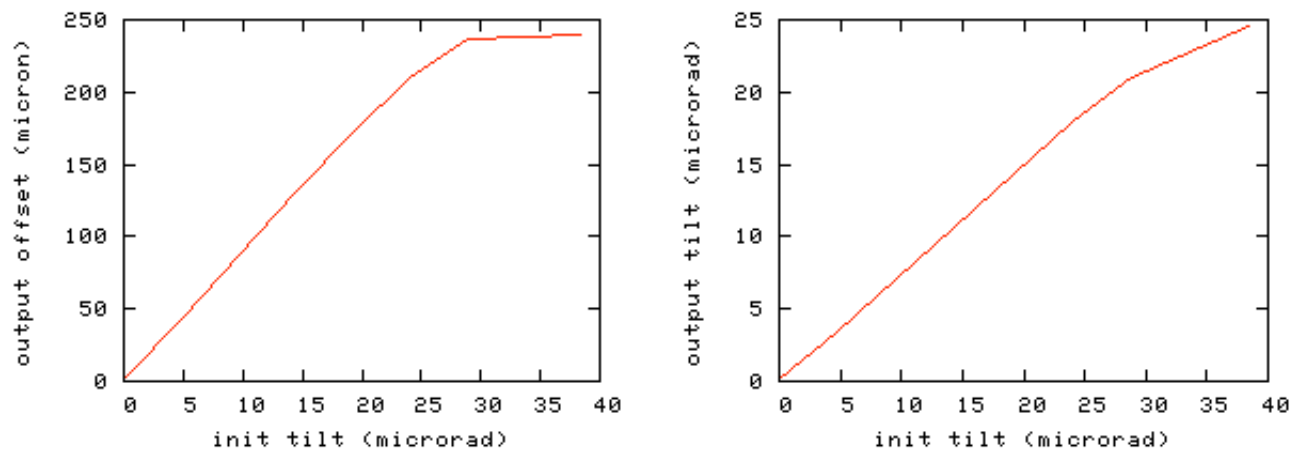


Figure 4.3-8: Output radiation misalignment as a function of initial e-beam tilt for the 10-nm fresh-bunch case.

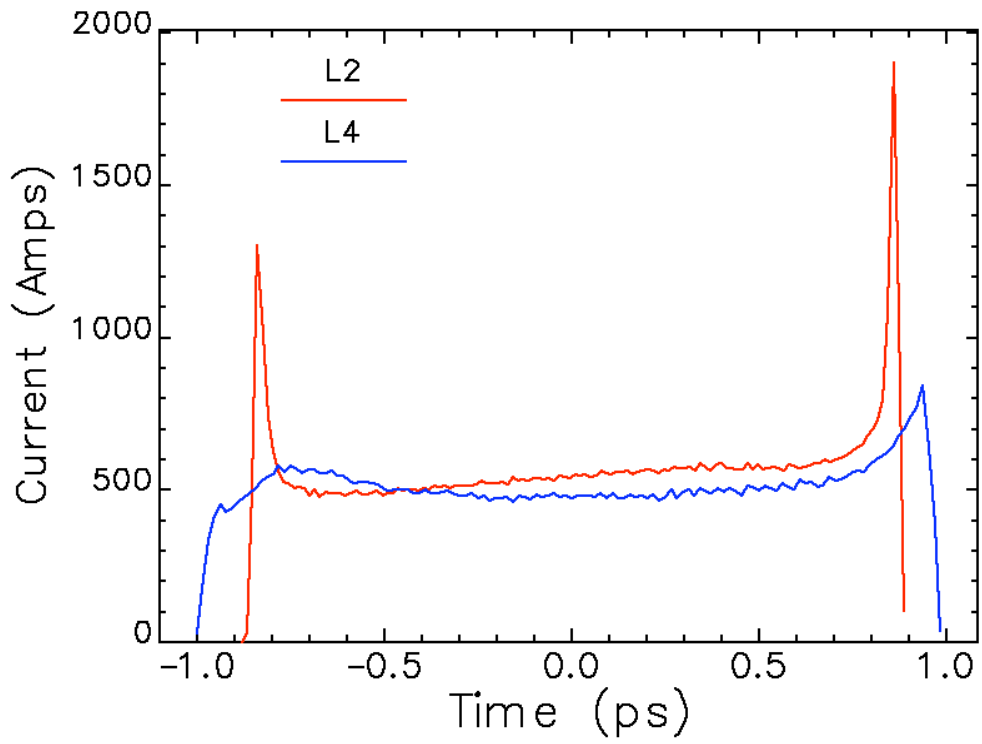


Figure 4.4-1: Current versus time profiles for the “L2” and “L4” long-duration bunch distributions as computed by the injector and S2E groups.

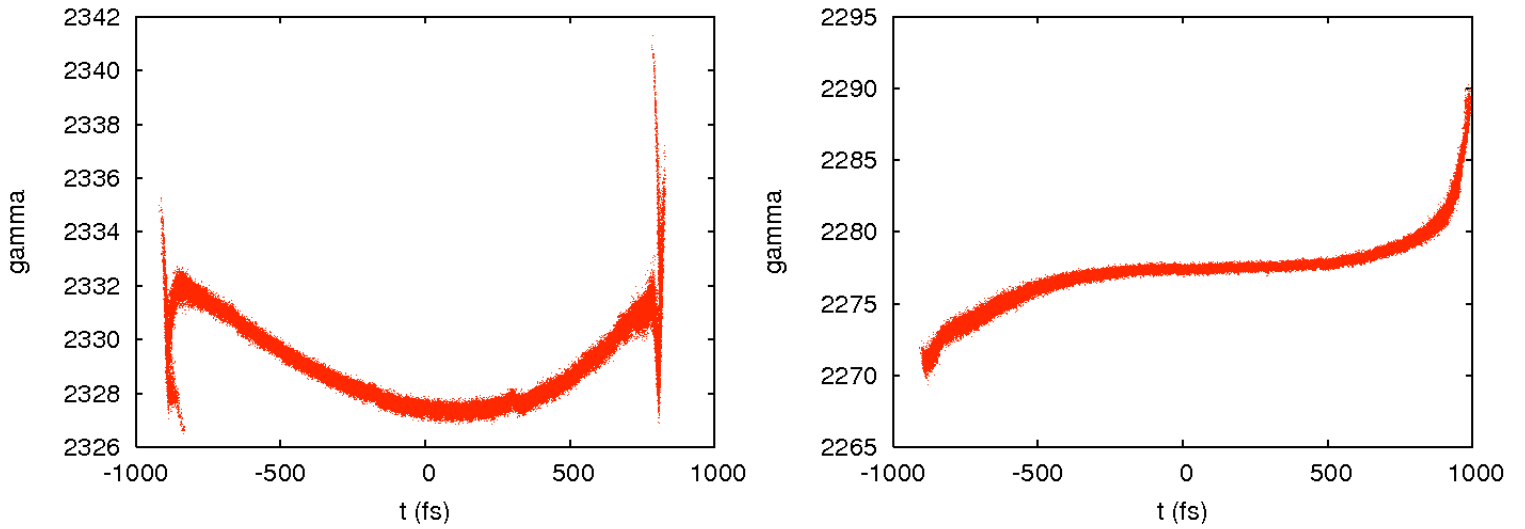


Figure 4.4-2: Longitudinal phase space profiles for the “L2” (left) and “L4” (right) long-duration bunch distributions.

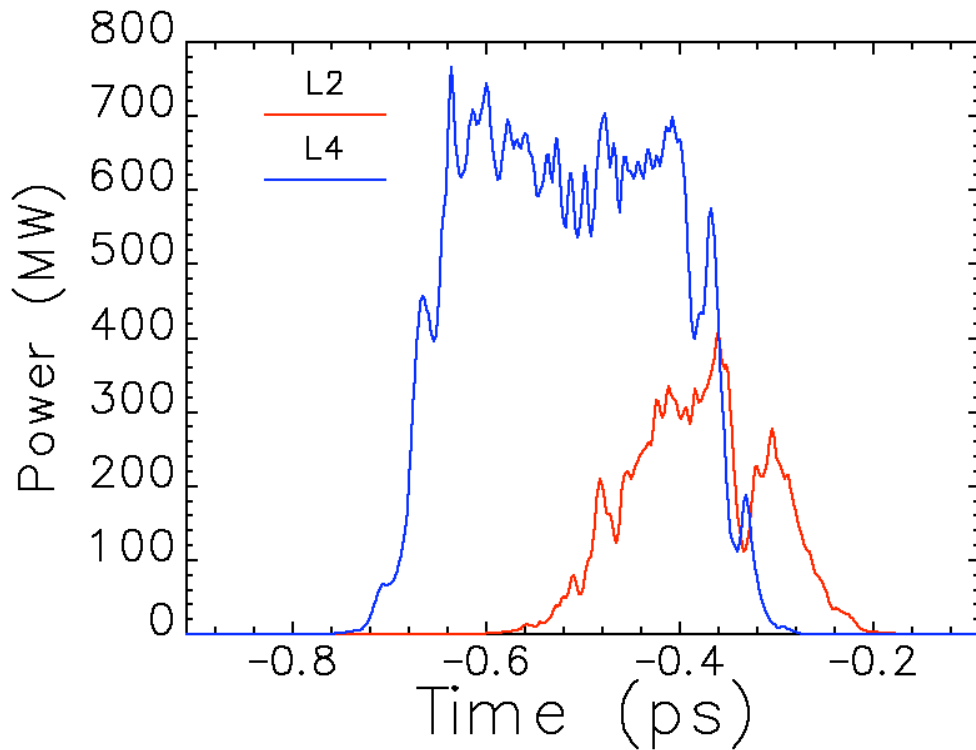


Figure 4.4-3: Output power at 10-nm wavelength versus time for the fresh bunch approach to FEL-2 predicted by GINGER using the L2 and L4 S2E long pulse electron distributions. The initial laser seed was a 200-fs (RMS) Gaussian pulse.

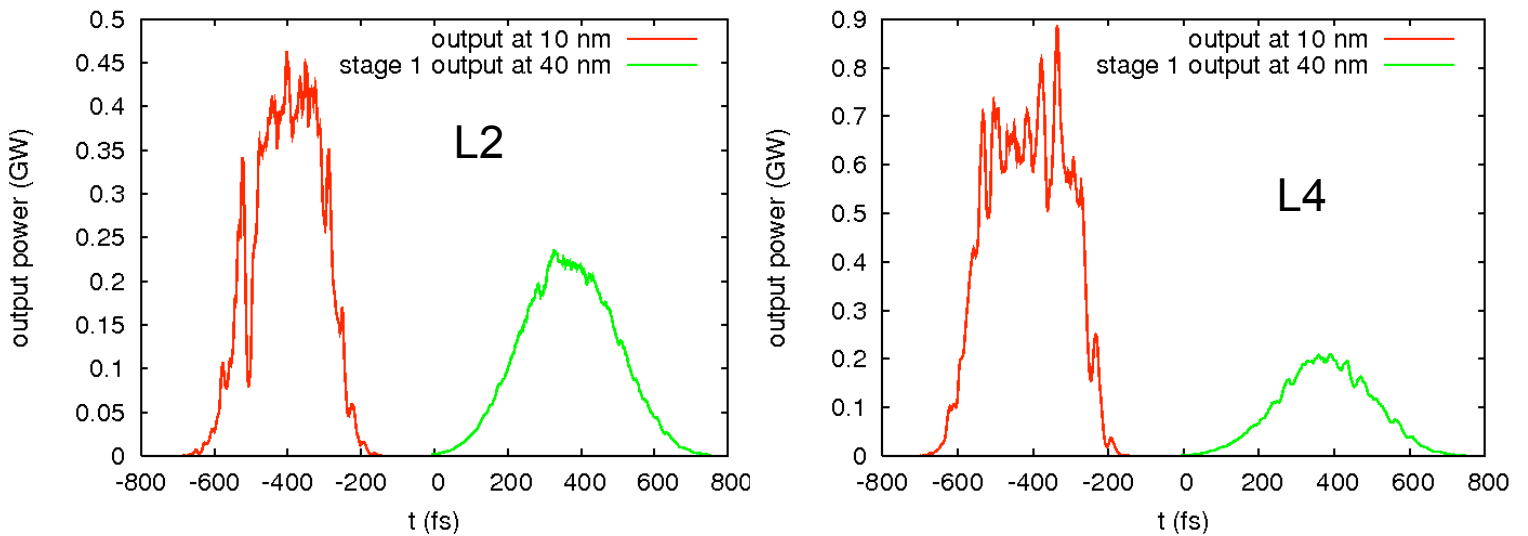


Figure 4.4-4: GENESIS simulation results for output power versus time at 10 nm wavelength (red curves) for the “L2” (left) and “L4” (right) distributions using the fresh-bunch configuration. The output power from the stage 1 radiator is also shown (green curves).

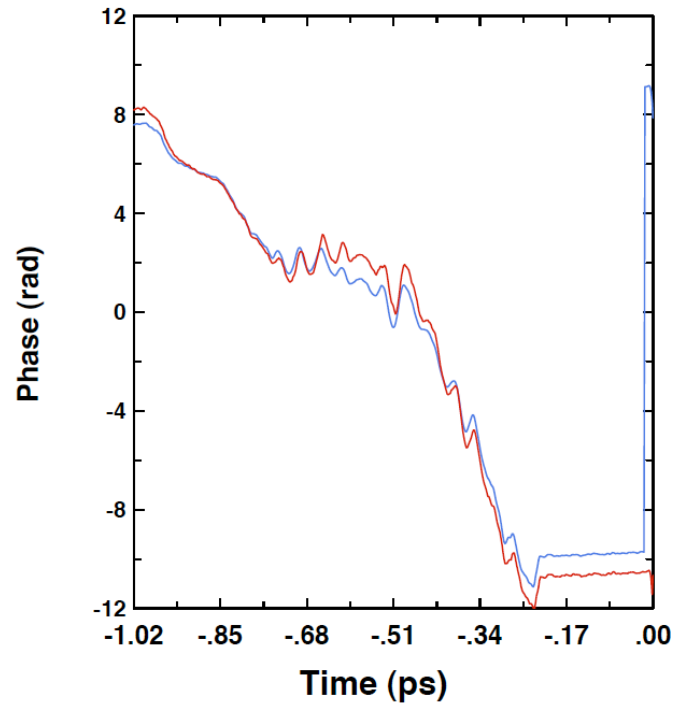
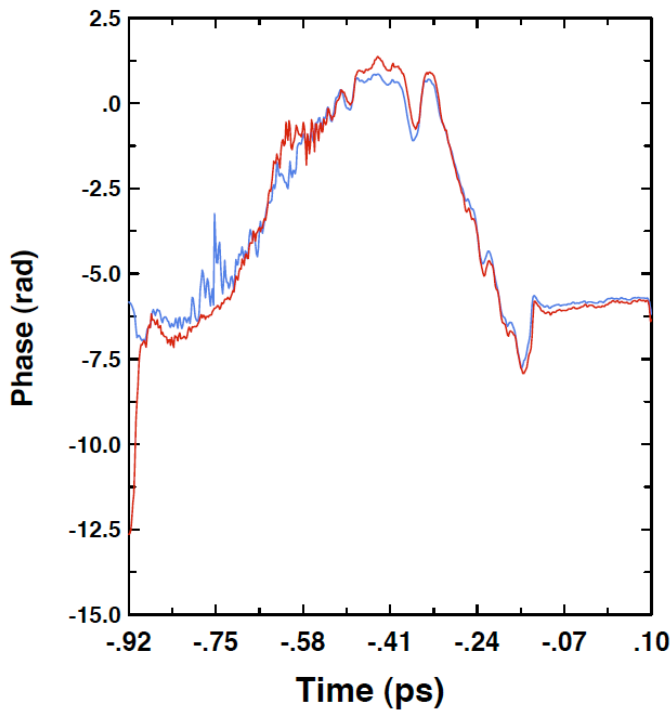


Figure 4.4-5: Output on-axis, far field eikonal phase (blue) at 10-nm wavelength and microbunching phase (red) versus time as predicted by GINGER using the L2 (left) and L4 (right) S2E long pulse electron distributions for FEL-2 with the fresh-bunch configuration.

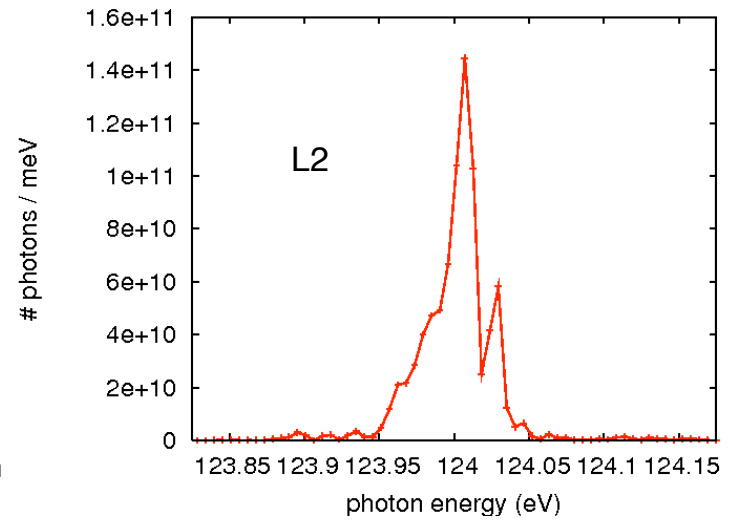
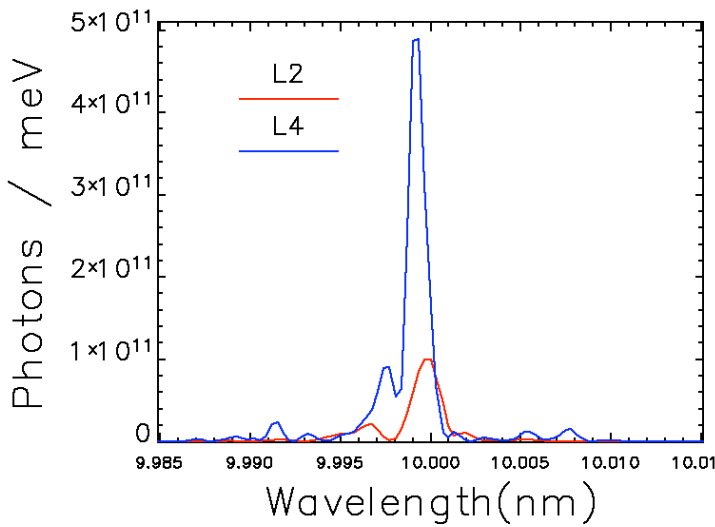
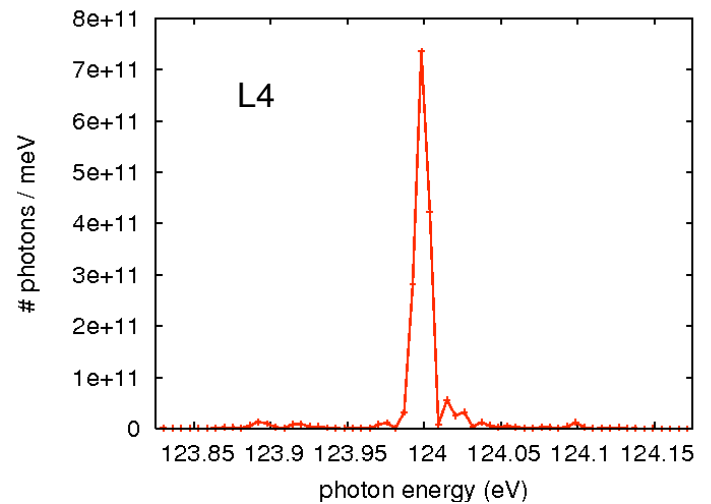


Figure 4.4-6: Output power spectra at 10-nm wavelength for the fresh bunch approach to FEL-2 predicted by GINGER (above plot) and GENESIS (right plots) using the L2 and L4 S2E long pulse electron distributions.



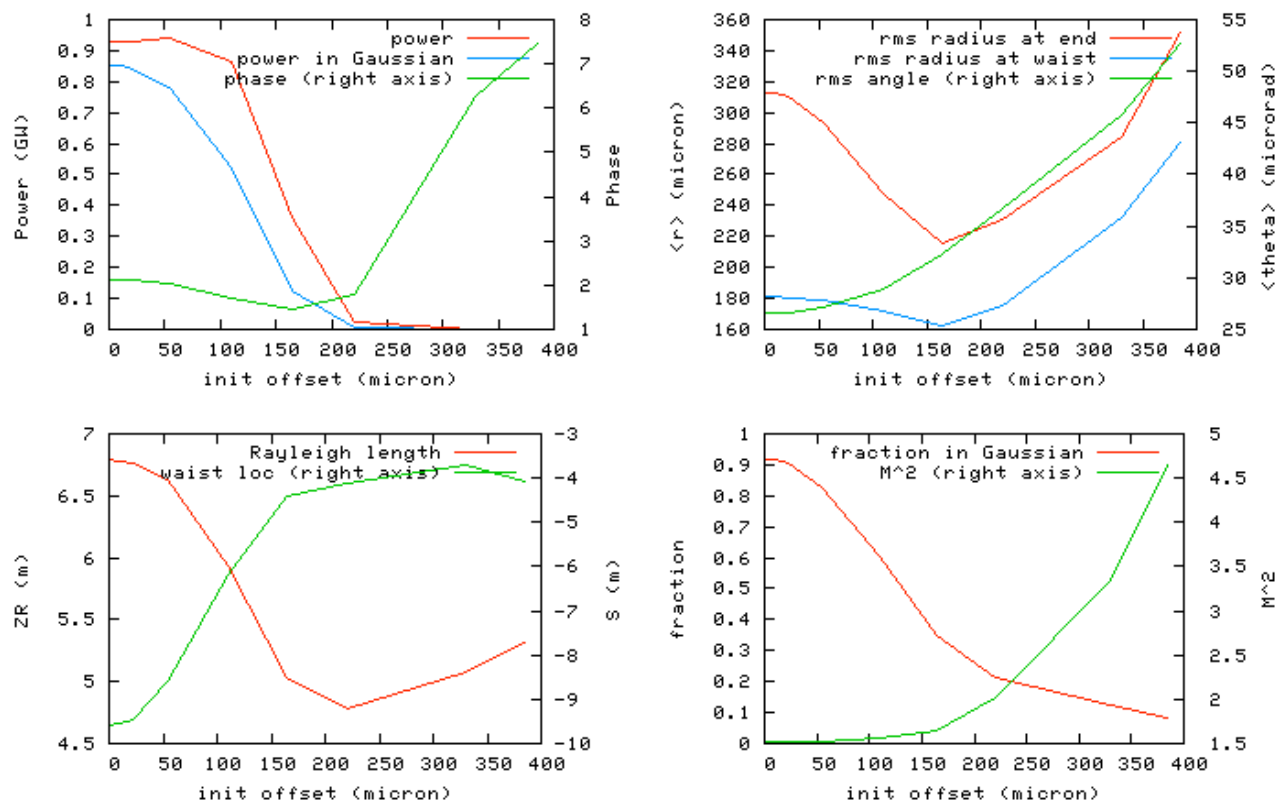


Figure 4.7-1: Sensitivity of output to electron beam offset, 20-nm whole-bunch case.

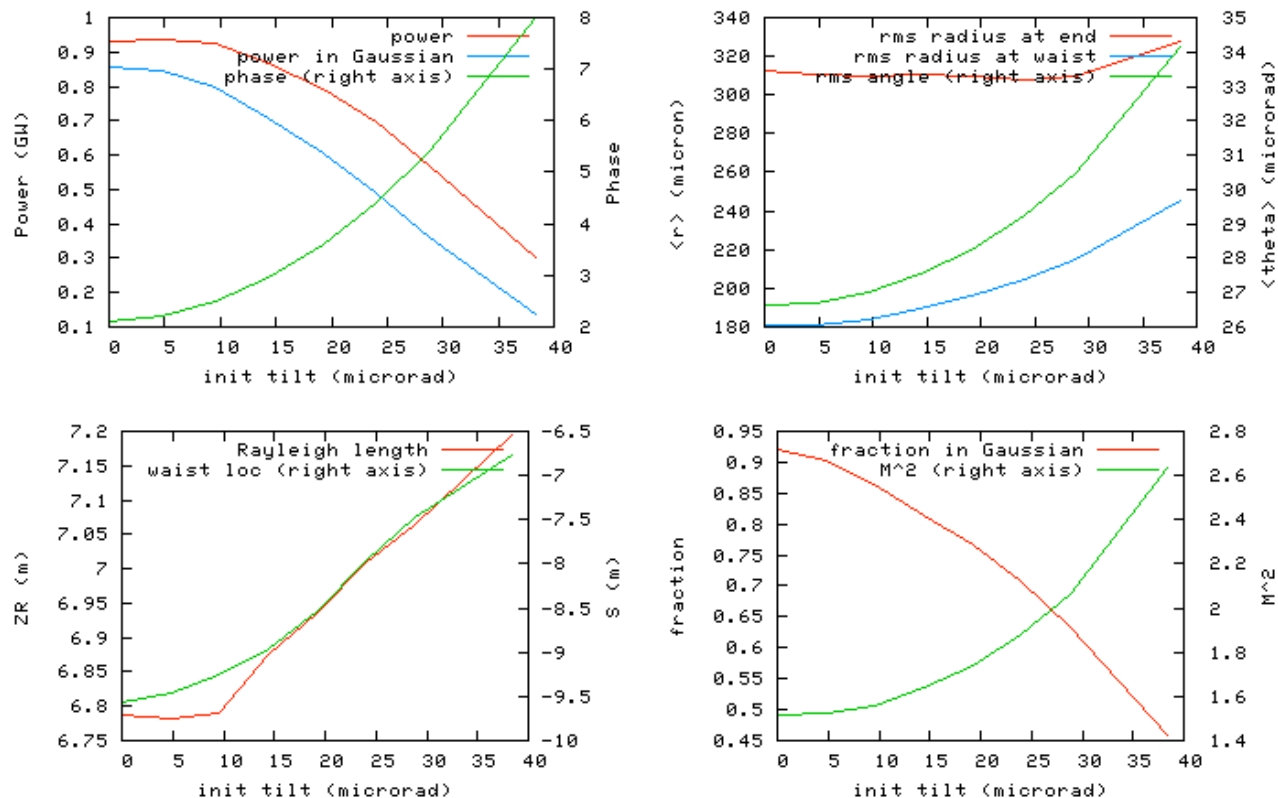


Figure 4.7-2: Sensitivity of output to electron beam tilt, 20-nm whole-bunch case.

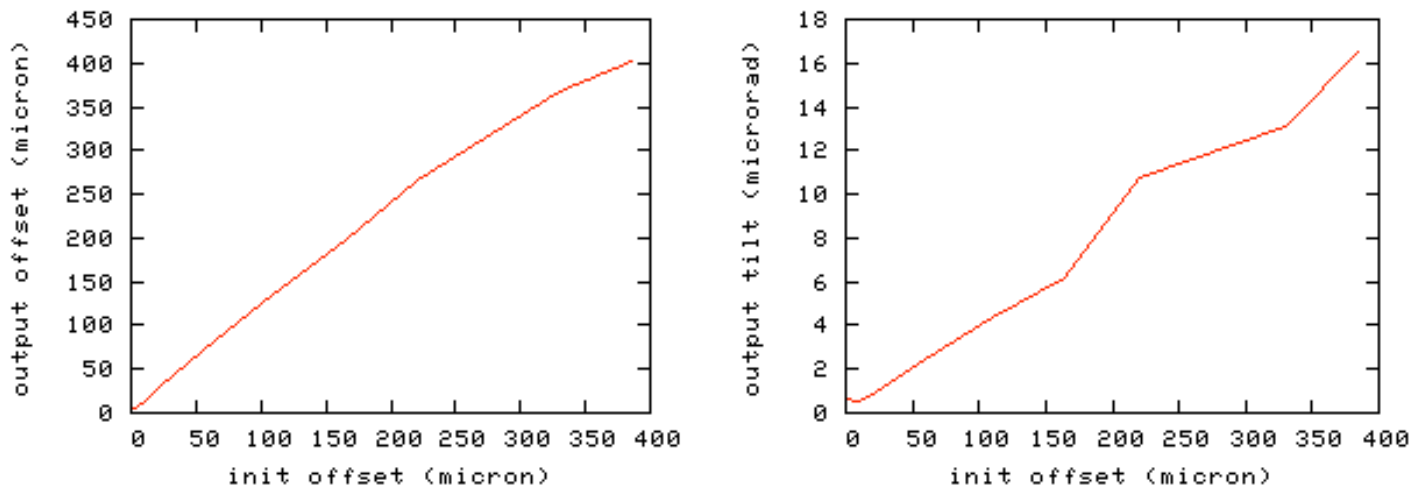


Figure 4.7-3: Output radiation misalignment from initial electron beam offset, 20-nm whole-bunch case.

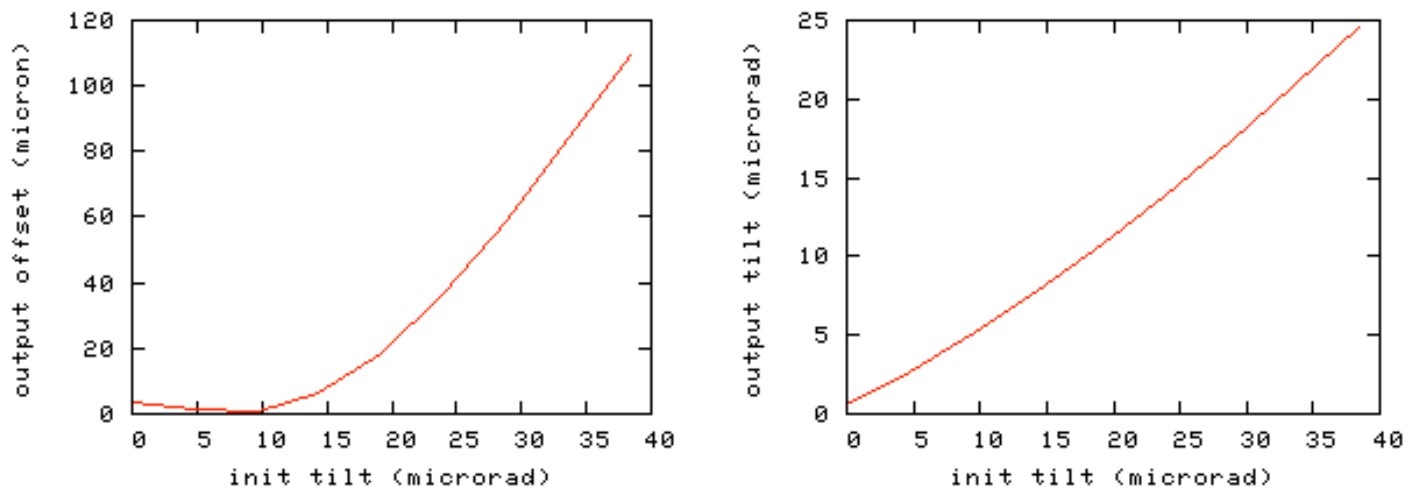


Figure 4.7-4: Output radiation misalignment from initial electron beam tilt, 20-nm whole-bunch case.

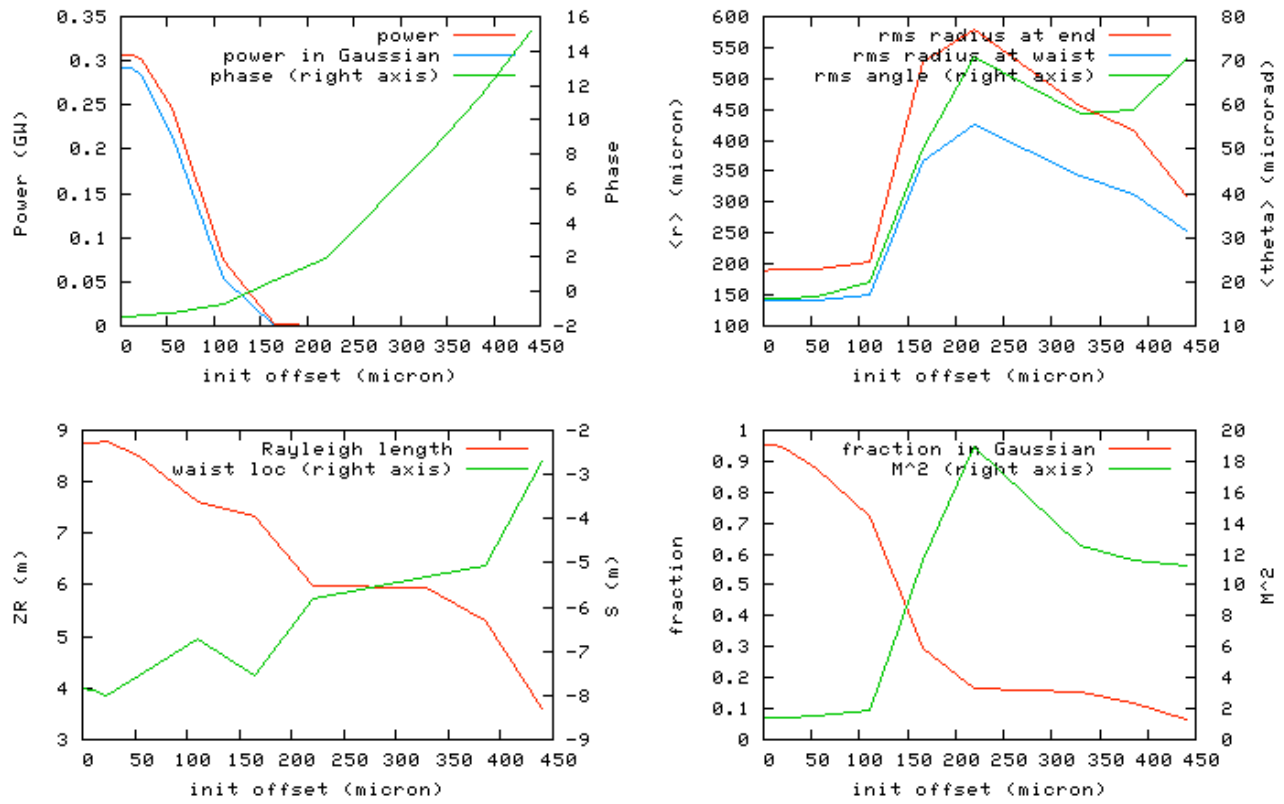


Figure 4.7-5: Sensitivity of output to electron beam offset, 10-nm whole-bunch case.

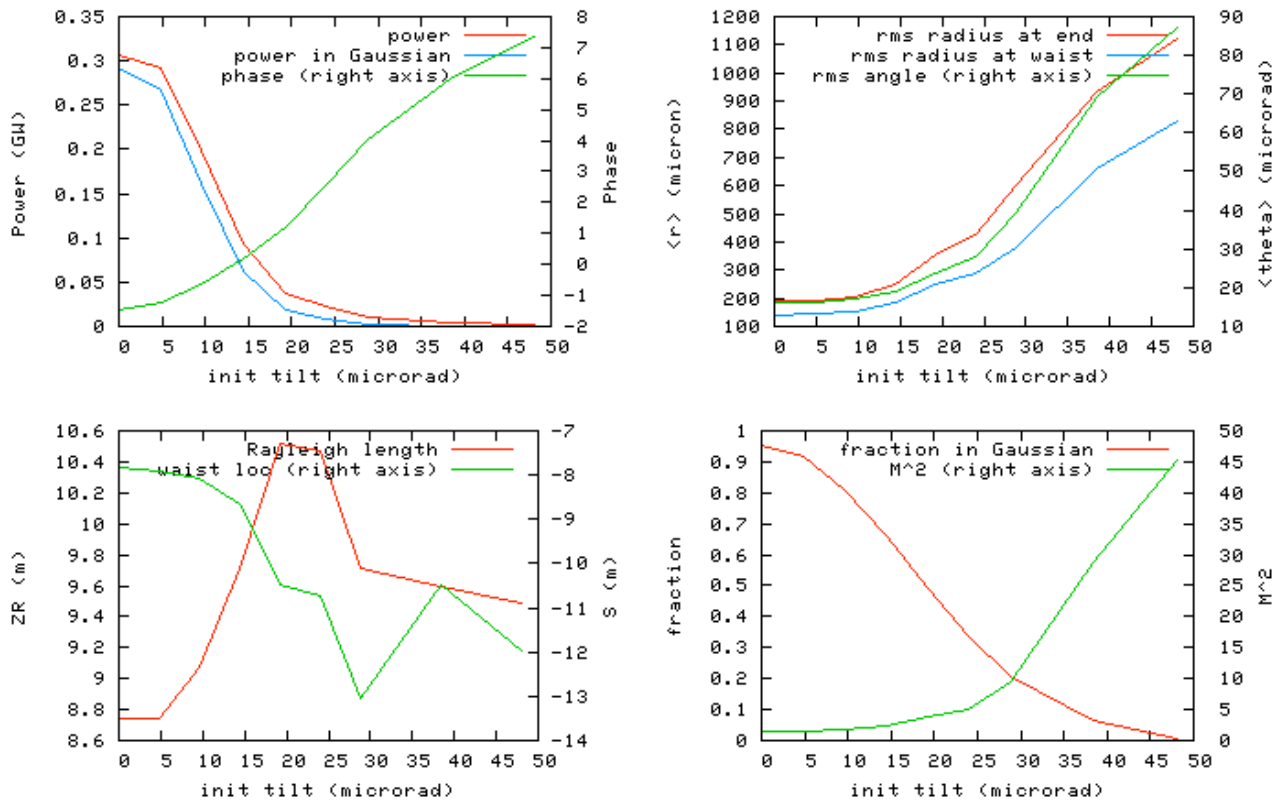


Figure 4.7-6: Sensitivity of output to electron beam tilt, 10-nm whole-bunch case.

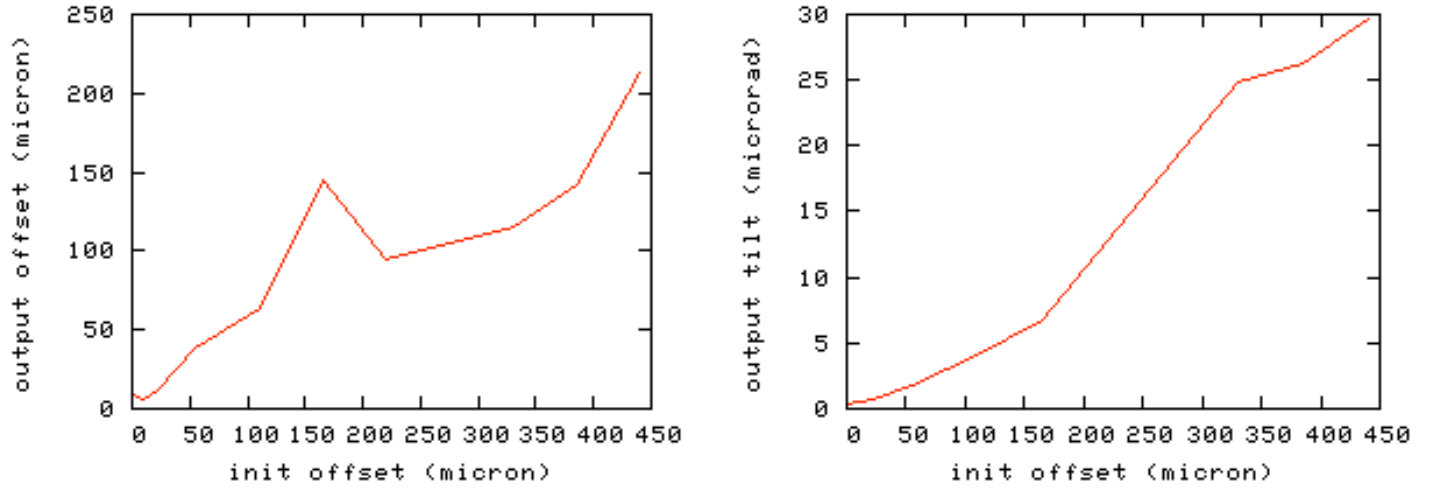


Figure 4.7-7: Output radiation misalignment from an initial electron beam offset, 10-nm whole-bunch case.

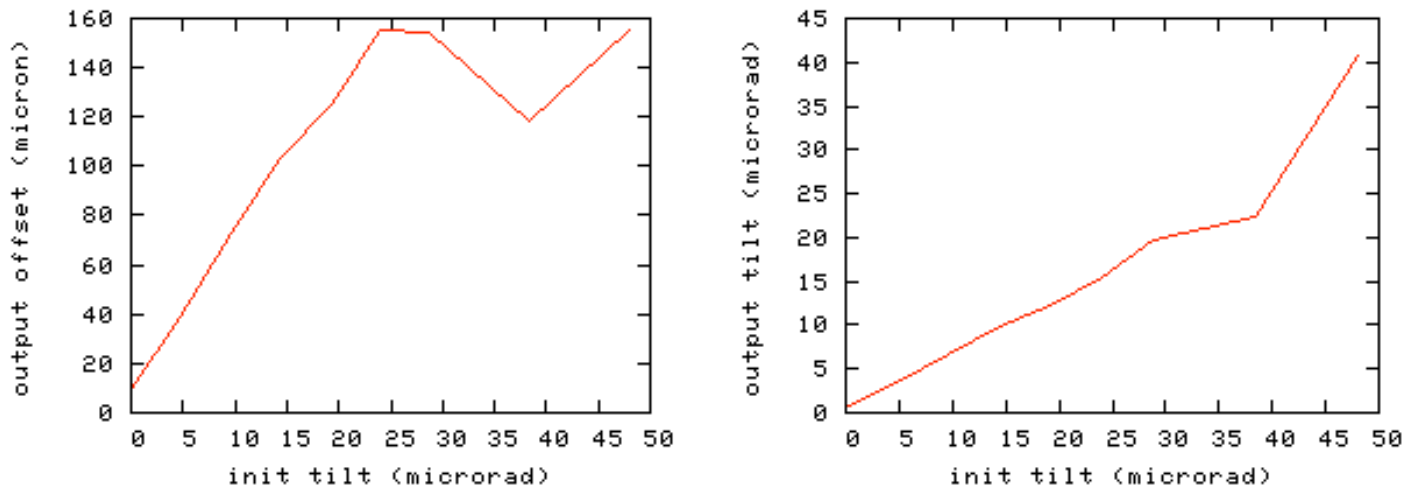


Figure 4.7-8: Output radiation misalignment from an initial electron beam tilt, 10-nm whole-bunch case.

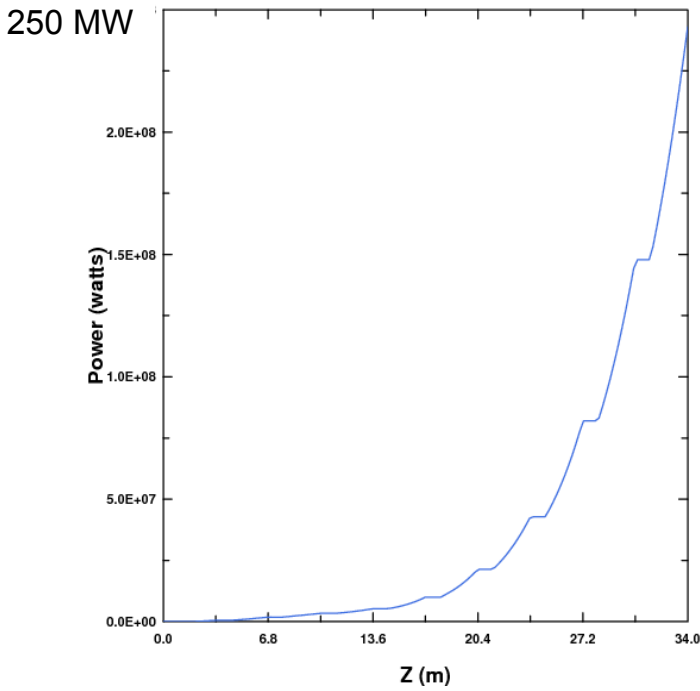


Figure 4.8-1: FEL power at 10-nm wavelength versus z for a long bunch optimized setup which minimizes output power sensitivity to electron beam jitter.

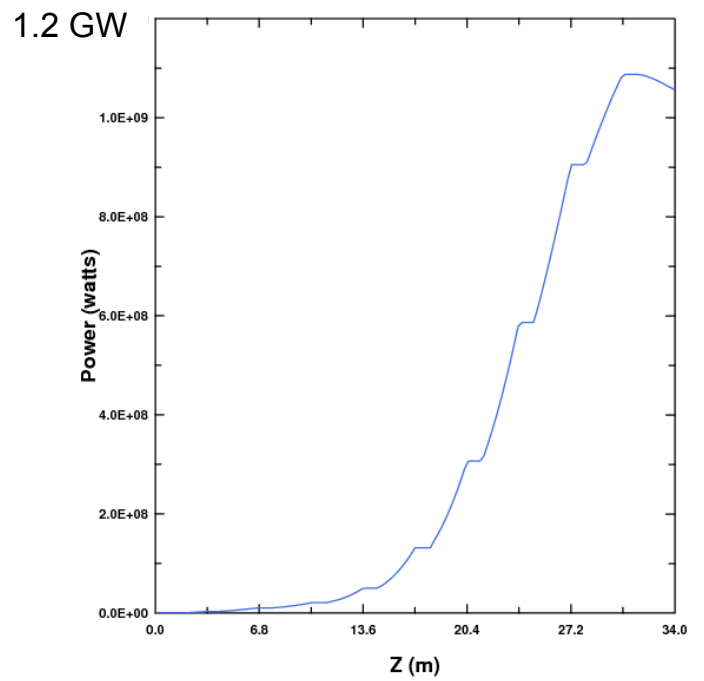


Figure 4.8-2: FEL power at 10-nm wavelength versus z for an optimized medium bunch setup which minimizes output power sensitivity to electron beam jitter.

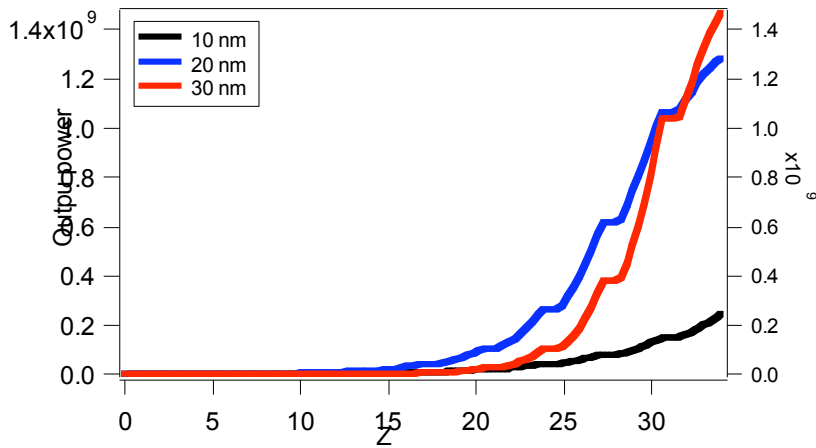
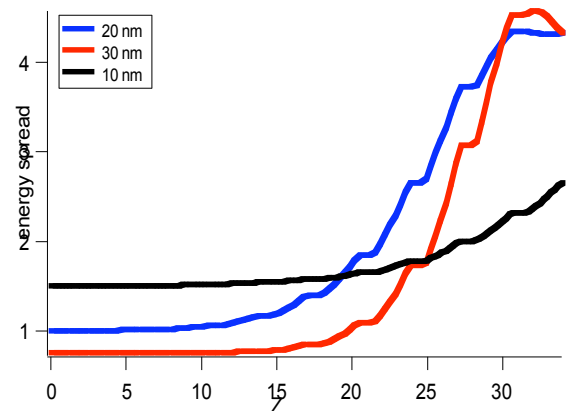
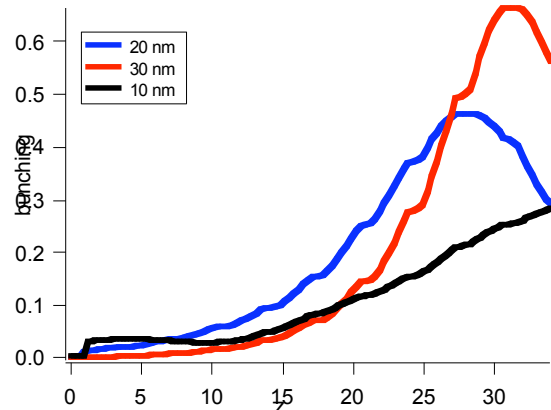


Figure 4.8-3: FEL radiation power, microbunching, and electron beam RMS energy spread at 10-, 20- and 30-nm output wavelength versus z for an optimized, "long bunch" FEL-2 whole bunch configuration that minimizes output spectral bandwidth.



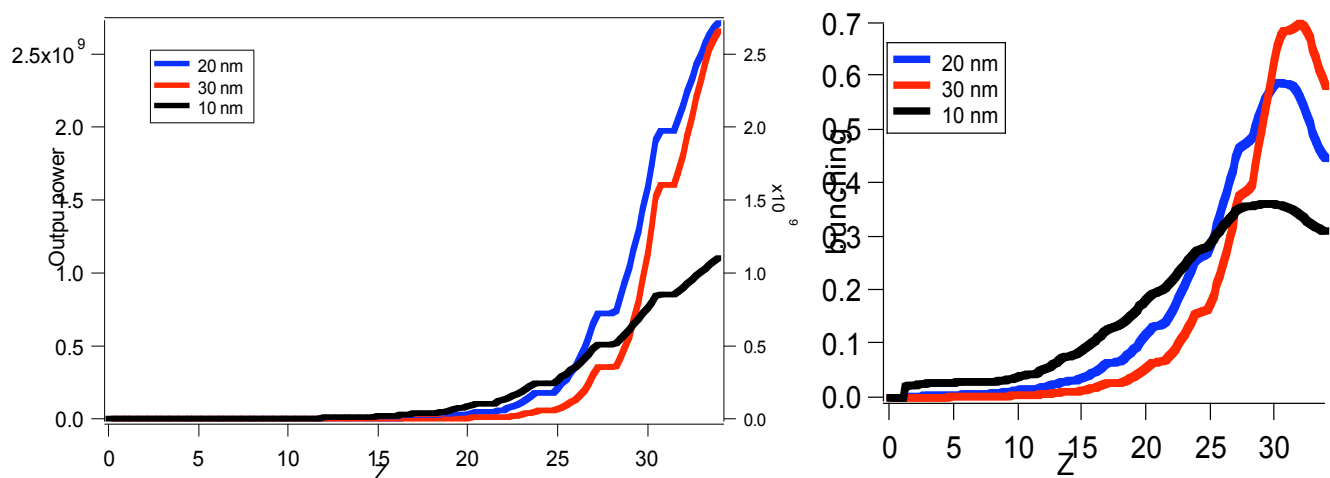


Figure 4.8-4: FEL radiation power, microbunching, and electron beam RMS energy spread at 10-, 20- and 30-nm wavelength versus z for an optimized, “medium bunch” FEL-2 whole bunch configuration which minimizes output bandwidth.

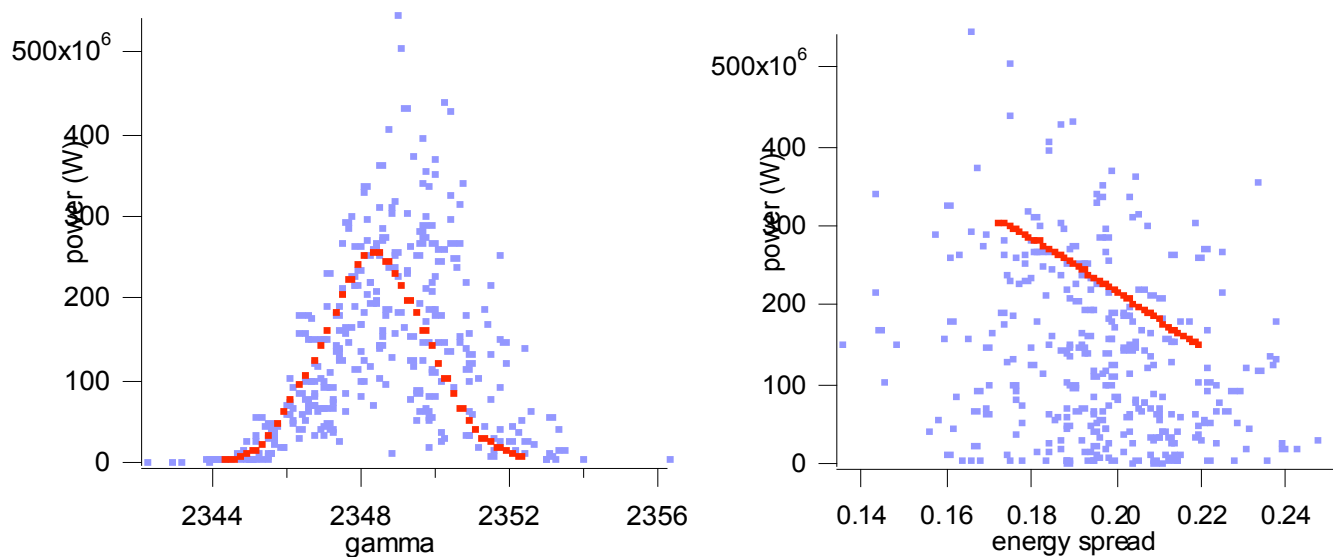


Figure 4.8-5: Time-independent jitter scans for an optimized (for narrow bandwidth) FEL-2 using the whole bunch approach. The electron beam had 500-Amp, “long bunch” parameters. The blue points correspond to individual “shot” output powers for simultaneous variation of input seed laser power, e-beam energy, incoherent energy spread, current, and emittance. The red curves correspond to a simple output power scan when *only* the variable corresponding to the abscissa is varied.

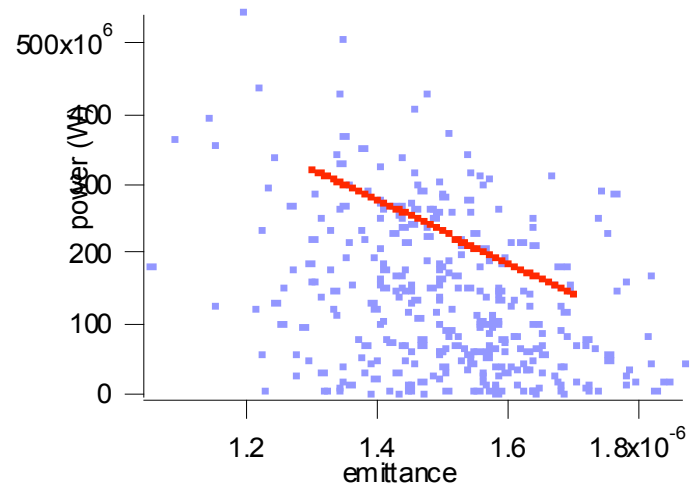
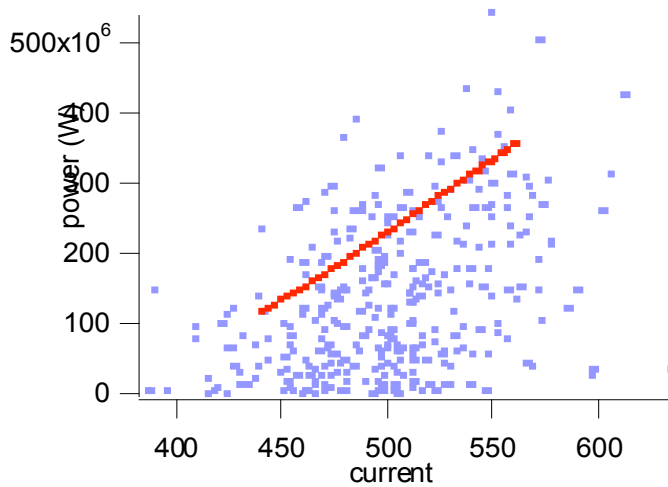


Figure 4.8-6: Time-independent output power jitter scans plotted against the projected current, emittance, and input seed laser power for an optimized (for narrow bandwidth) FEL-2 using the whole bunch approach with “long bunch” e-beam parameters. The red curves correspond to a simple parameter scan when only the abscissa variable is varied.

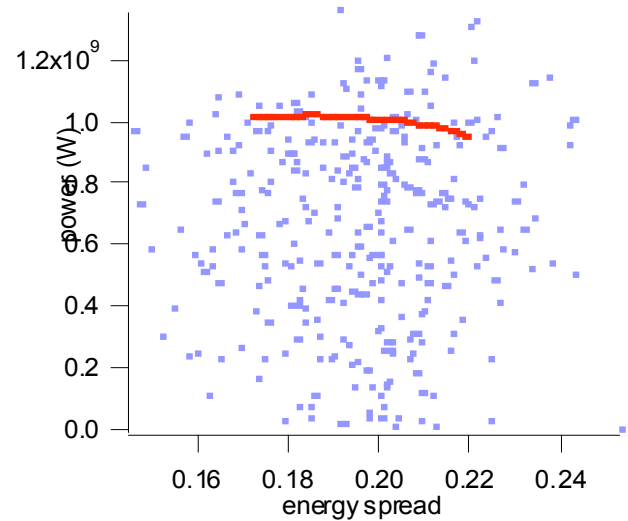
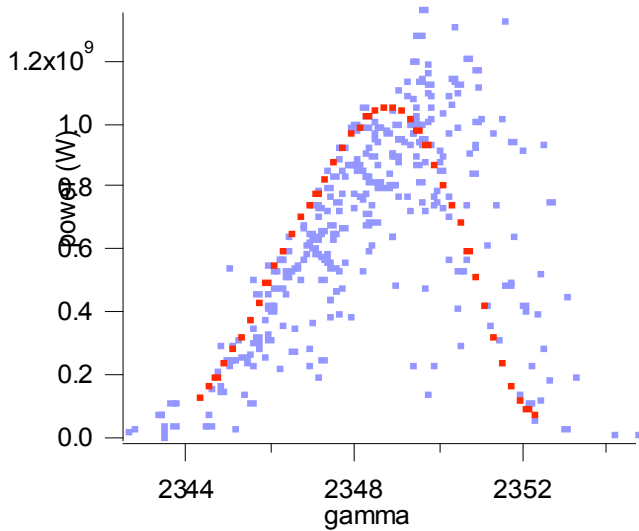
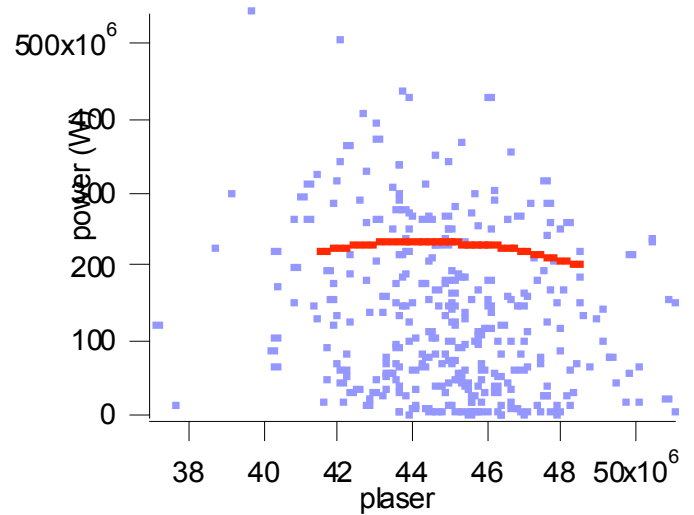


Figure 4.8-7: Time-independent jitter scans for an optimized (for narrow bandwidth) FEL-2 using the whole bunch approach plotted against projections of beam energy and energy spread. The electron beam had 800-Amp, “medium bunch” parameters. See the captions to Figs. 4.8-5 and 4.8-6 for further explanation.

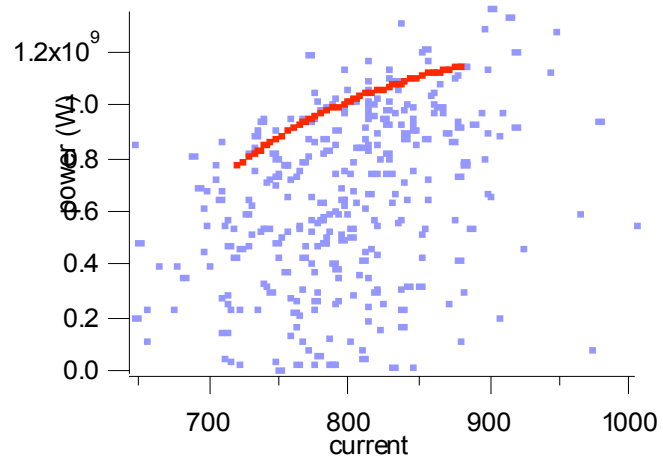
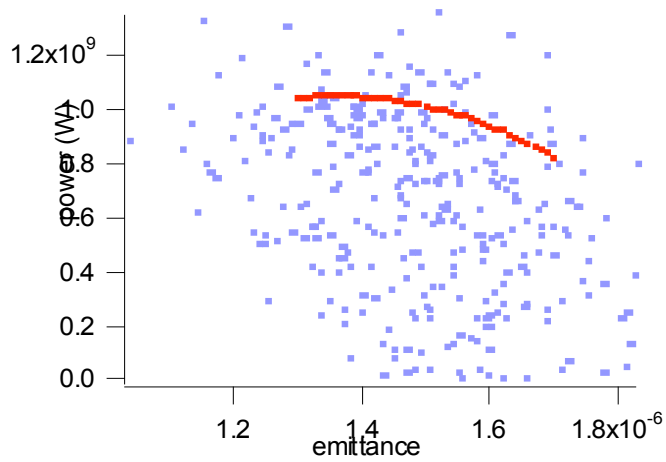
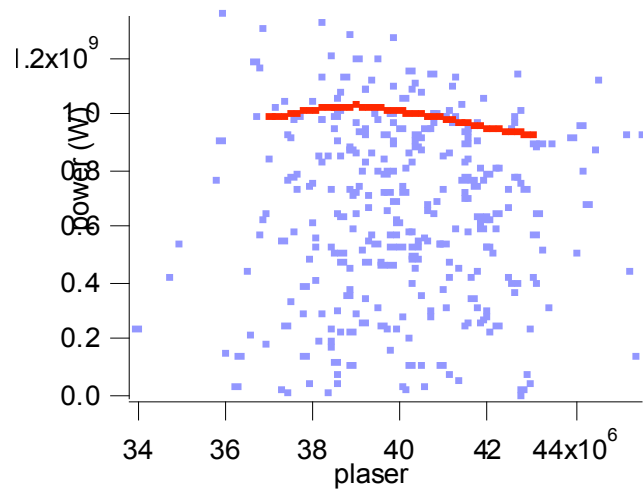


Figure 4.8-8: Same as Fig. 4.8-6 except for medium electron bunch parameters.



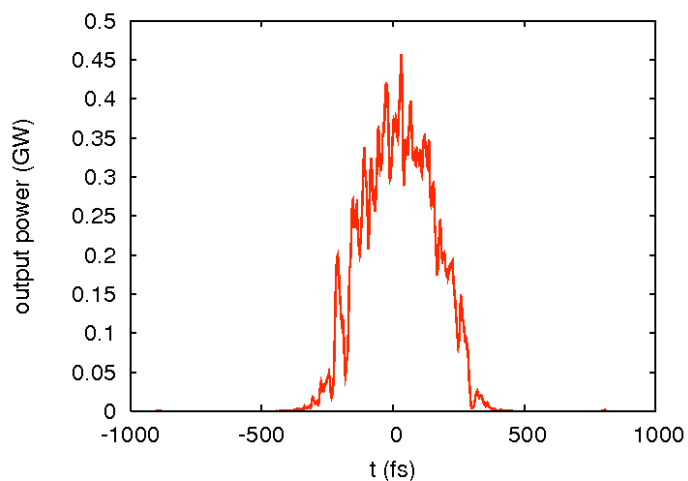


Figure 4.9-1: Output power at 10 nm versus time for the “L2” distribution using the whole-bunch configuration.

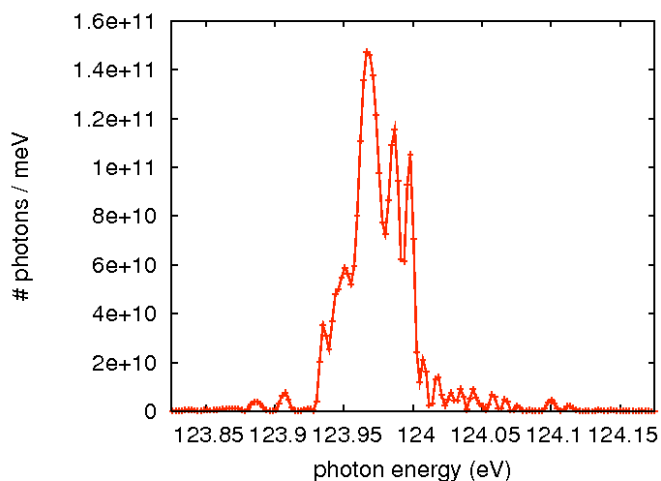


Figure 4.9-2: Spectrum at 10-nm wavelength for the “L2” distribution using the whole-bunch configuration.

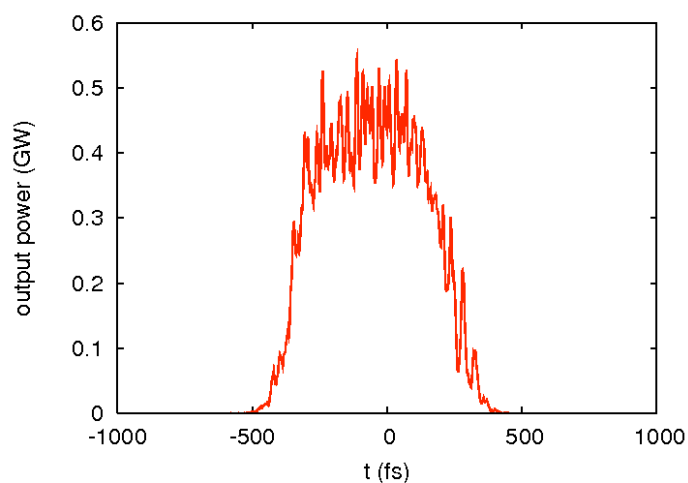


Figure 4.9-3: Output power at 10 nm versus time for the “L4” distribution using the whole-bunch configuration.

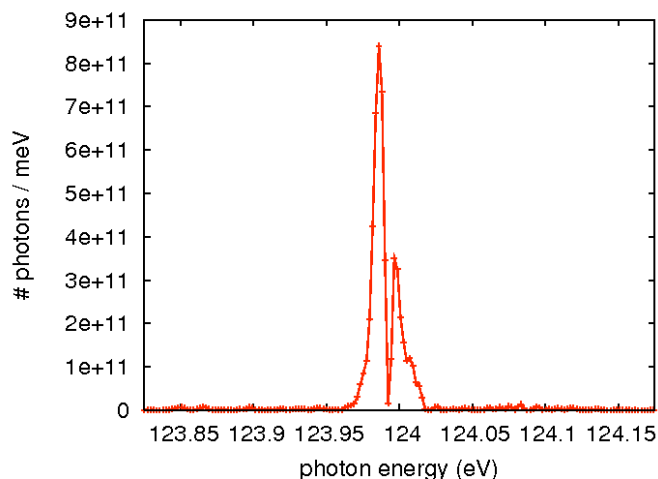


Figure 4.9-4: Spectrum at 10-nm wavelength for the “L4” distribution using the whole-bunch configuration.

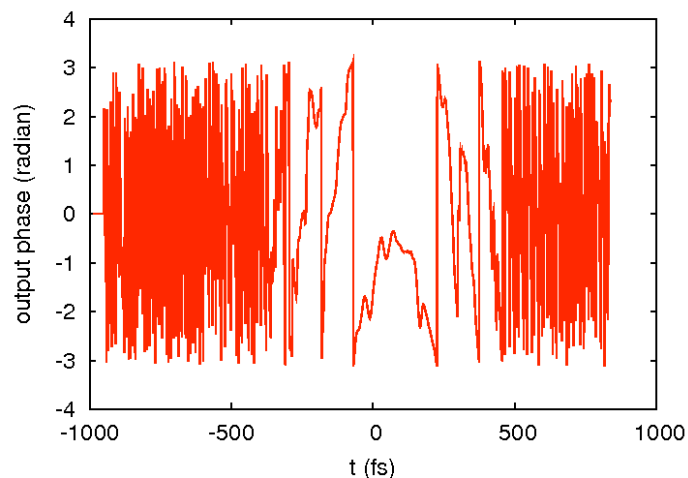


Figure 4.9-5: Far field on-axis eikonal phase profile at 10 nm for the “L2” distribution using the whole-bunch configuration.

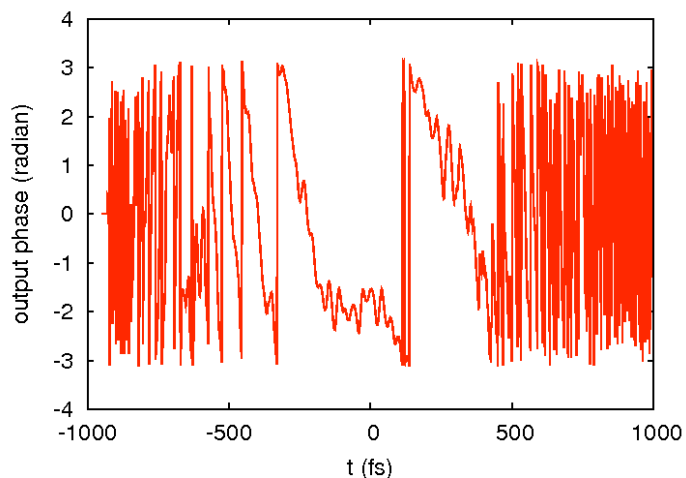


Figure 4.9-6: Far field on-axis eikonal phase profile at 10 nm for the “L4” distribution using the whole-bunch configuration.

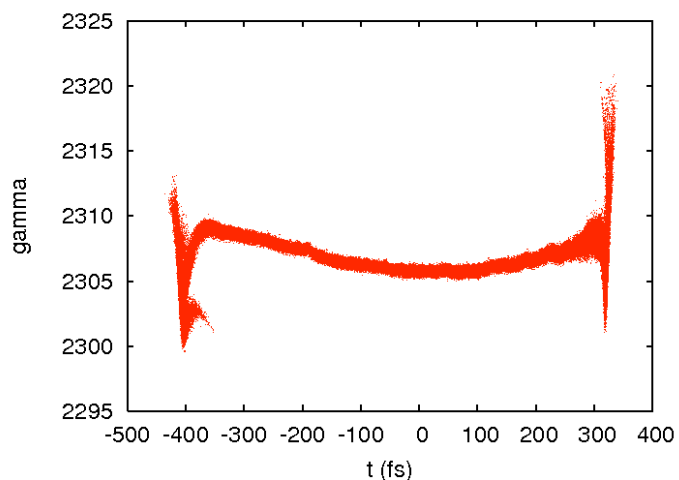


Figure 4.9-7: Phase space profile for the “M2” medium-duration bunch distribution at entrance to the undulator.

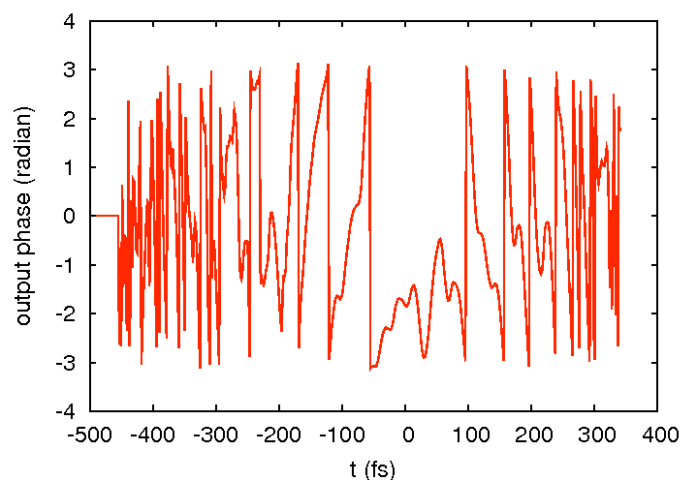


Figure 4.9-8: Far field on-axis eikonal phase profile at 10 nm for the “M2” distribution using the whole-bunch configuration.

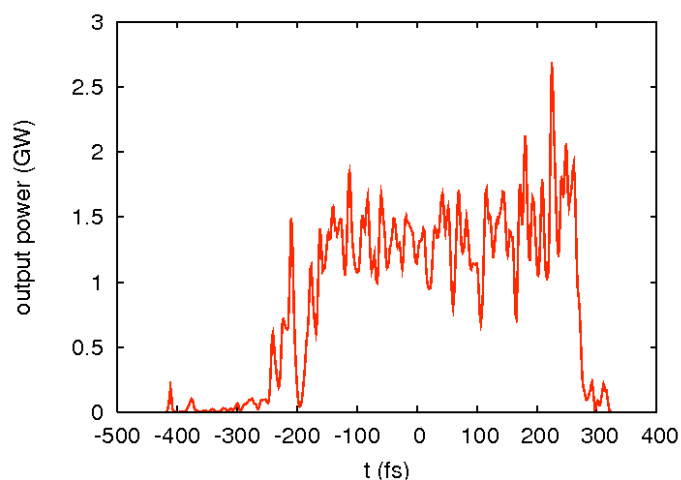


Figure 4.9-9: Output power profile at 10 nm wavelength for the “M2” distribution for FEL-2 using the whole-bunch configuration.

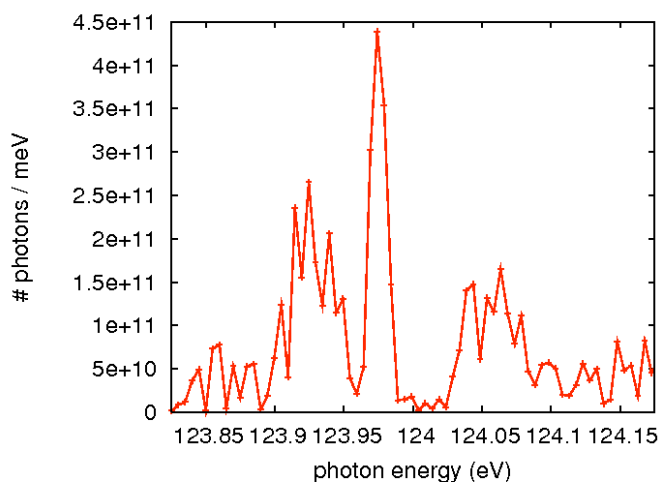


Figure 4.9-10: Spectrum at 10-nm wavelength for the “M2” distribution using the whole-bunch configuration.

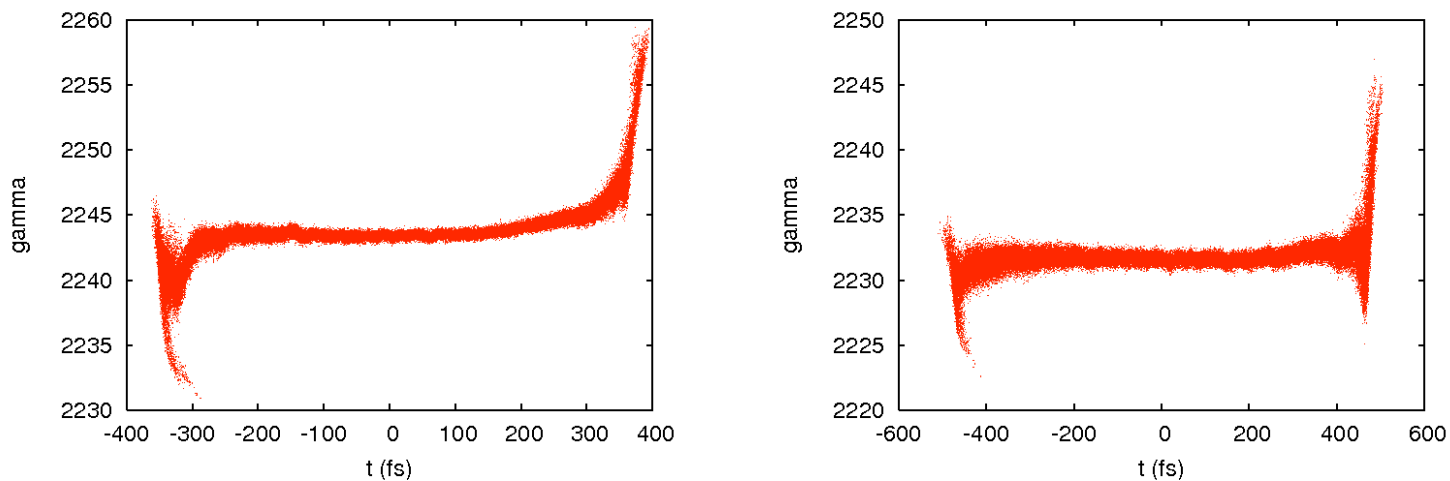


Figure 4.9-11: Phase space profiles for the “M4” (left) and “M6” (right) medium-duration bunch distributions.

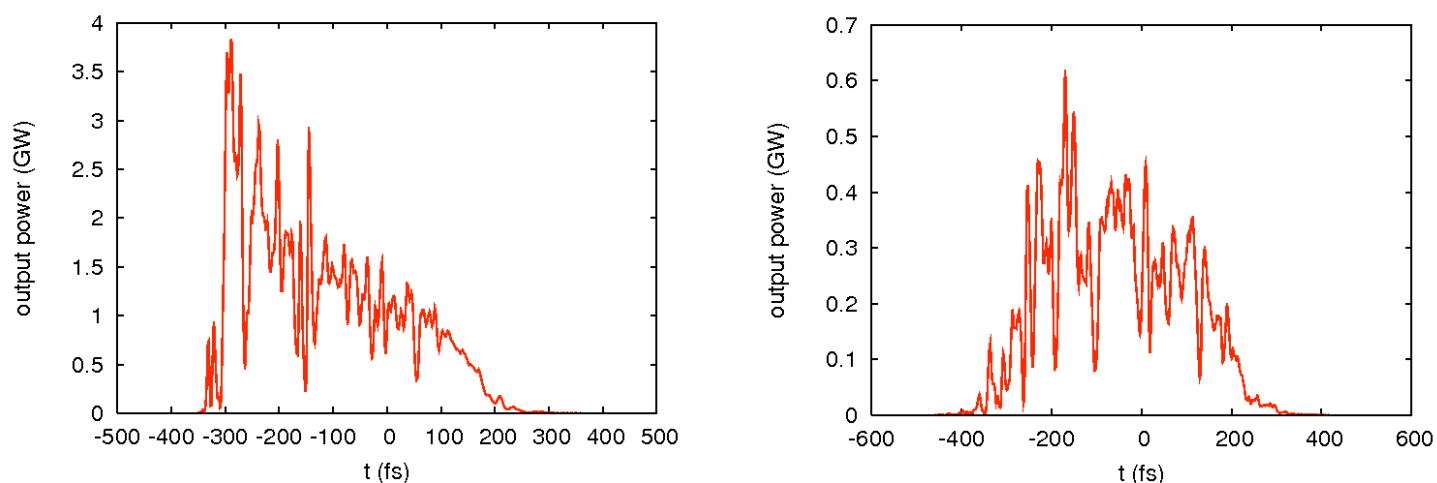


Figure 4.9-12: Output power profiles at 10 nm wavelength for the “M4” (left) and “M6” (right) distributions for FEL-2 using the whole-bunch configuration.

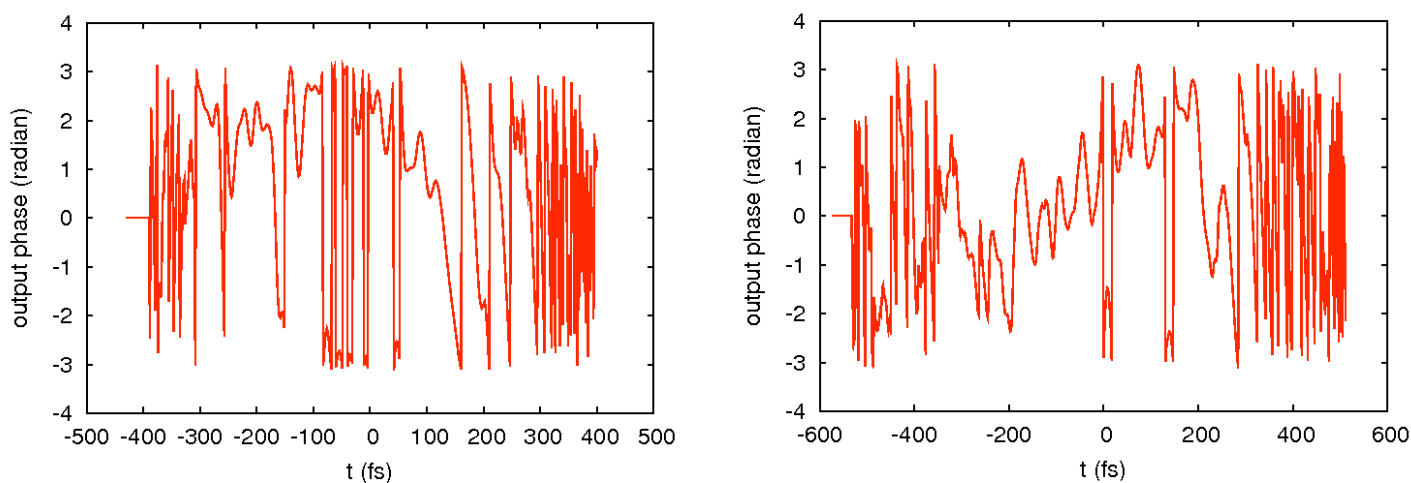


Figure 4.9-13: Output eikonal phase profiles at 10 nm wavelength for the “M4” (left) and “M6” (right) distributions for FEL-2 using the whole-bunch configuration.

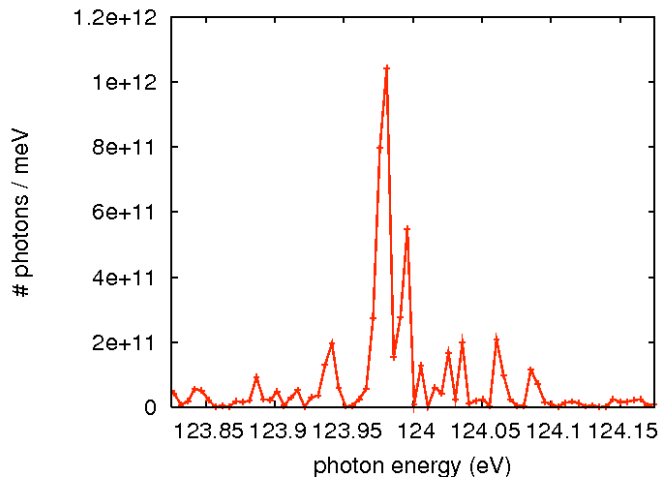


Figure 4.9-14: Spectrum at 10-nm wavelength for the “M2” distribution using the whole-bunch configuration.

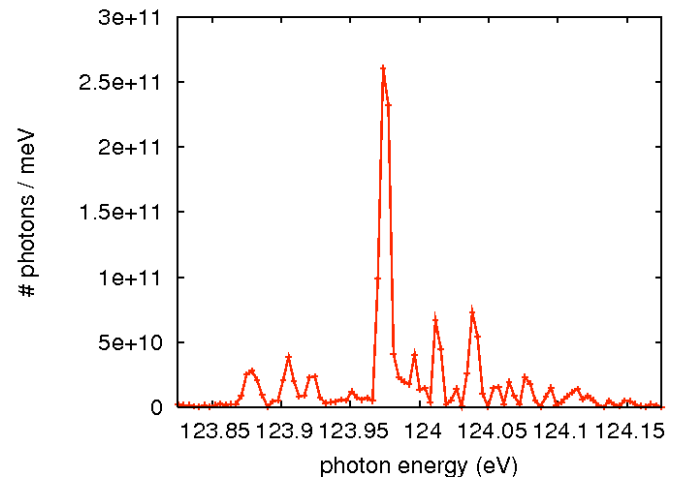


Figure 4.9-15: Spectrum at 10-nm wavelength for the “M2” distribution using the whole-bunch configuration.

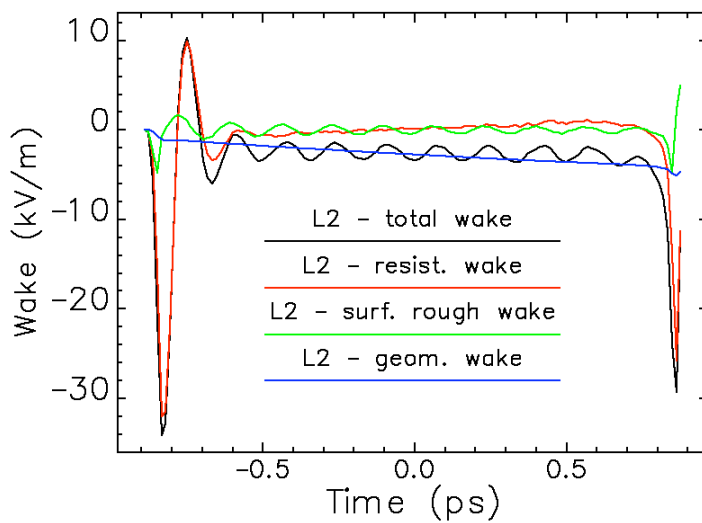


Figure 4.10-1: Time-dependent longitudinal wake results for the L2 S2E distribution. In addition to the total wake (black line), the individual components of the resistive, surface roughness, and geometric cross-section interruption wakes are also plotted.

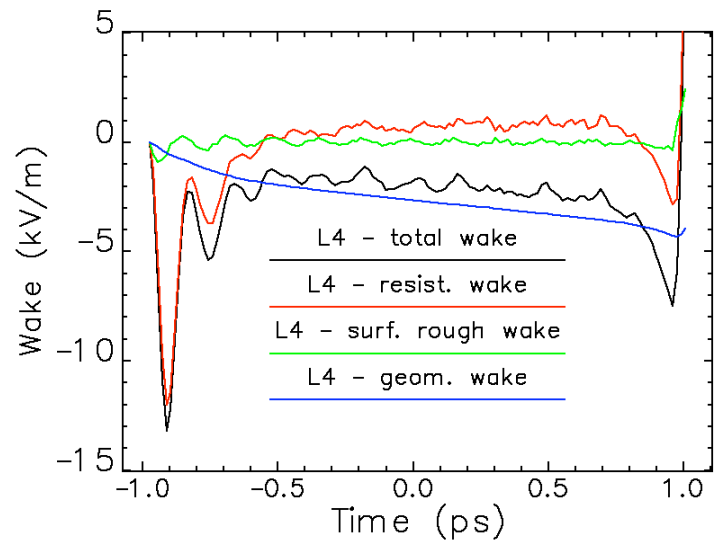


Figure 4.10-2: Time-dependent longitudinal wake results for the L4 S2E distribution. The different lines and colors have the same meaning as in Fig. 4.10-1.