

Slip Stacking in the Fermilab Main Injector

Shekhar Shukla, John Marriner and James Griffin
Fermi National Accelerator Laboratory, Batavia, IL 60563 USA

Abstract

The feasibility of increasing the number of protons available for antiproton production at Fermilab using a kind of momentum stacking called 'Slip Stacking'[1] is investigated.

I. INTRODUCTION

According to the current plan for Run II, the 120 GeV protons used in antiproton production will be obtained by transferring one booster batch into the Main Injector at 8 GeV and accelerating it to 120 GeV. We investigate the feasibility of increasing the antiproton production rate using 'slip stacking' in the Main Injector. This involves stacking two booster batches end to end but with slightly differing momenta, into the Main Injector. The two batches have different periods of revolution and 'slip' relative to each other azimuthally and finally overlap. When they overlap they are captured using a single rf which is the average of the initial frequencies associated with the two batches. The two batches might be moved closer together in momentum if a smaller longitudinal emittance for the final beam is desired. Since the booster and Main Injector acceleration cycles are 66ms and 1.5s respectively, we expect a substantial increase in the pbar production rate, if the process can be completed efficiently.

The following is a list of factors that determine the optimum momentum separation between the two batches, initially and before they are coalesced, and the rf voltages involved.

- 1) A larger momentum separation reduces the time before the batches can be coalesced.
- 2) A larger momentum separation requires a larger horizontal aperture.
- 3) A smaller momentum separation just before the batches are coalesced leads to a smaller longitudinal emittance for the final beam, if the effect of the second rf system is small.
- 4) The rf buckets for the two batches get more distorted as the separatrices move closer together. The losses become fairly high if the separatrices overlap. So the beams should spend as little time with their separatrices close together as possible before they are coalesced.

The procedure used to find a way to obtain a final coalesced beam of small emittance containing a reasonably large fraction of the initial beams, consists of two steps. The first step, described in more detail in section II, consists of finding the approximate heights of the initial buckets and the height of the final bucket after coalescing, which would result in a final beam of small longitudinal emittance with small losses during coalescing. We ignore the distortion of the rf buckets due to multiple frequencies in this step. In the next step, which is described in section III, we use these approximate heights of the buckets as starting points and use rf simulations that include the distortion of the buckets, to find an acceptable strategy of rf manipulations to achieve our

goals and to estimate the final emittance of the beam and the losses during coalescing when this strategy is used.

In Section IV we estimate the effects of beam loading and consider ways to overcome its adverse effects. The effects are serious due to the high intensity of the proton beam. Compensating for the effects is complicated by the simultaneous presence of two beams and rf systems.

II. OPTIMUM RF BUCKET HEIGHTS

The optimum bucket heights before and after coalescing were found assuming gaussian particle distributions for the initial beams. The harmonic number and rf frequency are 588 and approximately 53 MHz respectively. We use a value of 0.15 eV-s for the longitudinal emittance of each of the initial beams. This is the measured emittance in the Main Ring at injection. The emittance in the Main Injector is expected to be lower due to improved Booster performance. The height of a bucket with area 0.15 eV-s is 6.15 MeV.

For given heights of the initial and final buckets after coalescing, the area of the beam contour in the final bucket containing 95% of the initial beam was found by integrating the part of the initial gaussian distribution within the contour. The process was repeated for various heights of the final bucket. Fig.1 shows the area corresponding to various final bucket heights, for an initial bucket height of 6.2 MeV. The height that gave the minimum area was chosen as optimum final bucket height for the given initial bucket. The process was repeated for various values of the initial bucket height. Fig 2 shows the optimum final bucket heights and the heights of the corresponding beam contours containing 95% of the beam for various initial bucket heights. Fig.3 shows the minimum area containing 95% of the beam for various initial bucket heights.

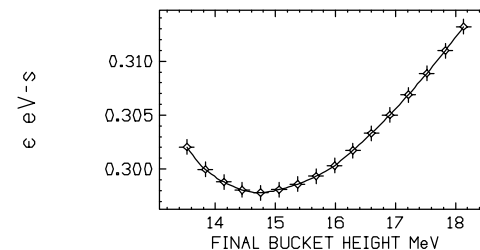


Figure 1: Area containing 95% of beam vs final bucket height for an initial bucket height of 6.2 MeV.

Even if the injected beams are gaussian, the beam distributions before coalescing are not expected to be gaussian if the two rf systems have frequencies that are close together. The distortion of the distributions due to the presence of the second rf was determined using a simulation of a beam in the presence of two rf systems and is described in section III.

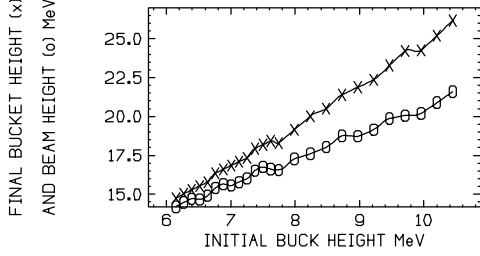


Figure 2: Optimum heights of the final bucket and the beam for various heights of the initial bucket.

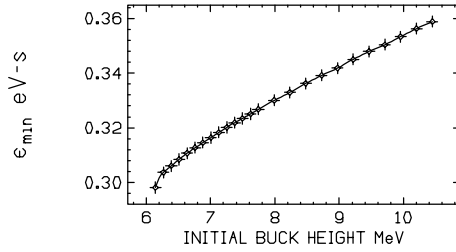


Figure 3: Minimum area containing 95% of beam for various heights of the initial bucket.

III. ACCELERATION AND COALESCING

The fractional difference in periods of revolution for the two batches is given by

$$\frac{\Delta\tau}{\tau} = \eta \frac{\Delta p}{p}, \quad (1)$$

where, $\frac{\Delta p}{p}$ is the fractional momentum difference and η is the slip factor. The slip factor is given by

$$\eta \equiv \frac{1}{2} - \frac{1}{2} \frac{\gamma_t}{\gamma} \quad (2)$$

For the MI, $\gamma_t = 21.8$, and at injection, using $\gamma = 9.55$, $\eta = 8.86 \times 10^{-3}$. The duration of a booster acceleration cycle, $T = 66.7$ ms. At injection, the length of a booster batch $l = 1.57 \mu s$, and the period of revolution in MI, $\tau = 11.14 \mu s$. If the two batches are injected 46 MeV apart and allowed to slip, they would overlap completely after half a Booster cycle, i.e., 33 ms. Simulations show that for a bucket height of 10 MeV, the distortion of the particle distributions due to the presence of the second rf is negligible. However, to obtain a small longitudinal emittance of the final beam, the two beams have to be accelerated towards each other before they are coalesced. The bucket height has to be reduced so that the two beams can be brought close together. Were it not for the effect of the second rf one could accelerate the beams and reduce the

bucket height very slowly to minimize particle loss. The presence of the second rf encourages faster rf manipulations once the beams are close to each other.

Optimization of the rf curves is complicated. After a few trials, the variation of the rf voltage, frequency and synchronous phase angle depicted in Fig.4 was accepted as satisfactory. The two beams are captured with a single rf while they are still accelerating. The efficiency of acceleration and coalescing for a final longitudinal emittance of 0.34 eV-s

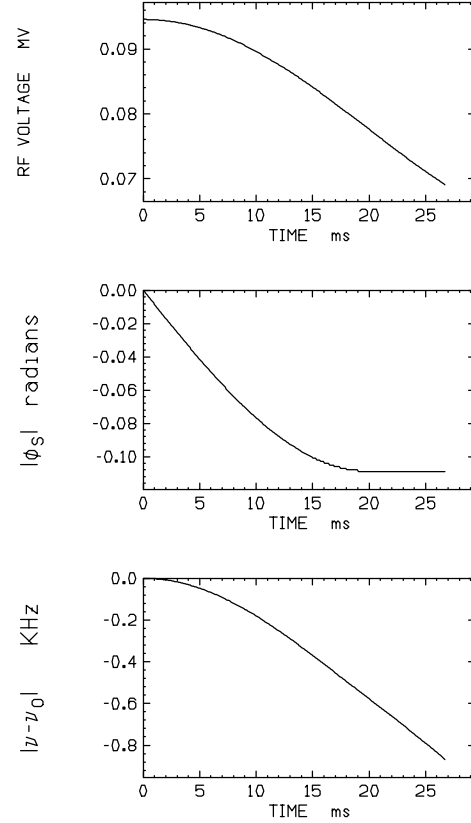


Figure 4: Variation of the rf voltage, synchronous phase and rf frequency used in the tracking simulation.

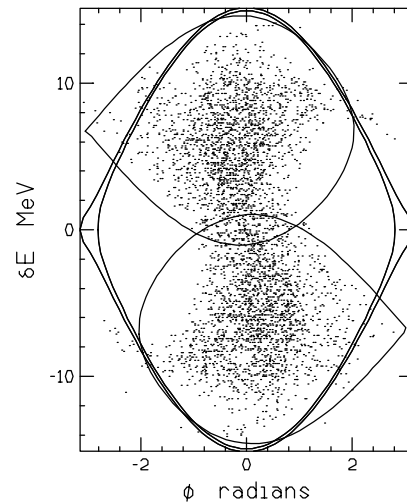


Figure 5: Beam distributions just before coalescing.

is 95%. Fig 5 shows the beam distributions just before coalescing as well as the shapes of the initial and final buckets. The dashed curve inside the final bucket is a contour containing 0.34 eV-s of area.

IV. BEAM LOADING

Because of the high beam current, the beam loading voltage in the rf cavities is a serious concern. If the quality factor, Q , is high and the bunch length is short, the cavity voltage $V(t)$ following the passage of a bunch of charge q is given by

$$V(t) = \frac{q\omega_r R}{Q} e^{-(\alpha+i)\omega_r t} \quad (3)$$

where R is the cavity shunt impedance, ω_r is the cavity resonant frequency, and $\alpha=1/2Q$.

In the case that the bunches are spaced by $\tau=2\pi/\omega_r$, the voltage after the passage of n bunches is easily found to be

$$V(n\tau) = \frac{q\omega_r R}{Q} \frac{1 - e^{-n\pi\alpha}}{1 - e^{-\pi\alpha}} \quad (4)$$

We can use eq.4 to estimate the beam loading voltage. As an example, we consider the case where there are two batches of 84 bunches each in the Main Injector and that the last 42 bunches of the first batch and the first 42 bunches of the last batch overlap and are exactly in phase. We ignore the difference in revolution frequencies of the two batches and the difference between the resonance frequency of the cavity and the revolution frequencies. Under these circumstances, one can use a generalization of eq.4 to estimate the beam loading voltage as shown in fig. 6. The calculation is for a total of 18 cavities with $R/Q=100 \Omega$ and $Q=5000$. The voltage increases when the beam passes through the cavities. During the time that the two beams overlap the voltage increases at twice the rate. When the beam is absent the voltage decays at a rate determined by the time constant α .

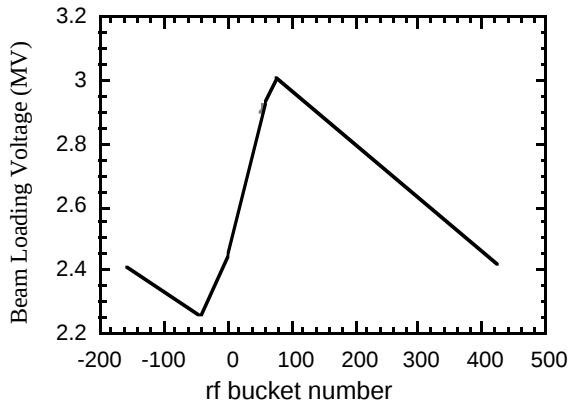


Figure 6. Beam loading voltage

Approximately 0.4 ms later the bunches are out of phase and the beam voltage becomes very small.

This estimate of the beam loading voltage indicates that, if uncompensated, the beam loading voltage (3 MV) would dwarf the rf voltage (100 kV). We propose to control the beam loading voltage by:

1. Tuning all cavities to the nominal 8 GeV frequency.
2. Using a small number of cavities (2 or perhaps 4) to produce the required rf voltage and de-Qing the remaining cavities. One simple technique that appears to be moderately effective is to turn off the screen voltage to reduce the tube plate resistance. This technique is estimated to de-Q the cavities by a factor of 3.
3. Using feed-forward on all the cavities. A resistive gap measures the wall current. This current, after being properly scaled, can be applied to the cavity drivers. Based on current Main Ring experience it is expected to achieve a factor of 10 reduction in the effective beam current.
4. Using feedback on all the cavities. A signal proportional to the gap voltage is amplified, inverted, and applied to the driver amplifier. This technique is expected to achieve a factor of 100 reduction (based on previous experience in the Main Ring and results achieved elsewhere).

If all these efforts are successful, beam loading should be reduced to a negligible value. Experiments to measure the suppression of the beam loading voltage and calculations of the tolerance of the slip stacking process to large beam loading voltages should determine whether slip stacking is feasible in the presence of relatively high impedance rf cavities.

V. CONCLUSIONS

Slip Stacking appears to be promising for enhancing the antiproton production rate at Fermilab. The reduction in beam loading voltage required can be estimated with a computer simulation and should be done soon. The proposed methods of reducing the beam loading voltage and the effects of the reduced voltage on the beam need to be studied experimentally, perhaps using the Main Ring. Whether Slip Stacking turns out to be useful will probably be determined by the ability of the Main Injector to accelerate the increased number of protons efficiently.

VI. REFERENCES

- [1] D. Boussard and Y. Mizumachi, "Production of Beams with High Line-Density by Azimuthal Combination of Bunches in a Synchrotron", CERN-SPS/ARF/79-11; 1979 IEEE Particle Accelerator Conference (1979).
C. Ankenbrandt, "Slip-Stacking: A New Method of Momentum Stacking" Fermilab FN-352 (1981).