

TWO-WAY TIMING MEASUREMENT WITH UPLINK COMPENSATION

LARRY R. D'ADDARIO

National Radio Astronomy Observatory*
2015 Ivy Road, Charlottesville, Virginia 22903, U.S.A.

INTRODUCTION

In OVLBI, it is necessary to establish a two-way link for transferring timing information to the spacecraft from a stable oscillator on the ground. This is often called a two-way “phase” link, or two-way “Doppler” link, because the implementations chosen for the current generation of spacecraft (Radioastron and VSOP) use only a single carrier in each direction. The phase terminology corresponds to a frequency-domain viewpoint, but in this paper we will take the equivalent time-domain viewpoint because it allows easier elucidation of some points. The Doppler terminology corresponds to the view that only the derivative of time or phase, i.e., frequency, is being transferred. This may be adequate for some applications, but for VLBI we need to take the more general view that time is being transferred.

The purpose of the link is to establish the times of certain events on the spacecraft, such as the zero-crossings of LO waveforms and the sampling times of the signals, as they would be measured by a clock synchronized to one on the ground. There is no such synchronized clock; instead, the spacecraft’s oscillators are locked to a reference signal from the ground, so spacecraft times correspond to times in the transmitted waveform plus the propagation delay. The spacecraft then sends a downlink signal from which we can determine the two-way delay and make appropriate corrections to the timing of signals on the spacecraft. This depends on an assumption of reciprocity, that is, that the uplink and downlink delays are equal, or at least that any non-reciprocal effects are well enough understood to be correctable. Notice that the main purpose of the two-way link is to determine the uplink delay; the downlink delay is determined as a byproduct, but is unimportant because the time of each signal sample was established on the spacecraft before it was transmitted.

In practice the determination of the delay is subject to some errors. A major error is caused by ambiguity in the transmitted waveforms; if they are periodic, then we cannot tell one cycle from another, so if the period is less than the two-way delay then we have a (possibly large) uncertainty. The use of quasi-sinusoidal signals at microwave frequencies gives periods near 10^{-10} sec; with two-way delays of 0.1–0.5 sec, this gives very little information on the total delay. More sophisticated designs could avoid this ambiguity [1]. But even with this limitation, it is possible to measure accurately the *change* in two-way delay that occurs during any continuous period of contact between one earth station and the spacecraft. In this way, whatever error is made remains constant during the contact period and can be corrected for the entire period if it can be determined at one epoch. This was discussed in earlier work by the author [1,2], and the present paper is an extension of that work.

*The National Radio Astronomy Observatory (NRAO) is operated by Associated Universities, Inc., under a cooperative agreement with the National Science Foundation.

Figure 1. Timing diagram showing the relative times of events in the uplink, predicted downlink, and actual downlink at two different epochs during one tracking pass.

EFFECT OF UPLINK COMPENSATION

Consider the timing diagram in Figure 1, which shows uplink and downlink signals as observed at the earth station. Once again, we prefer to think in the time domain, so the uplink and downlink signals are shown as a series of pulses, but these could actually represent the zero-crossing times of quasi-sinusoids. The uplink signal is transponded at the spacecraft and a corresponding downlink signal is generated. We will neglect any delay within the spacecraft. If a particular event in the uplink signal could somehow be labeled so that the corresponding event in the downlink signal could be recognized, then the two-way delay could be measured on the earth station's clock. But since the cycles in each signal are indistinguishable, there is an ambiguity. Fortunately, we know *a priori* the approximate two-way delay, so there is a limited range of uplink "pulses" that might correspond to the present downlink "pulse." We can then pursue the following strategy.

At some fixed time t_0 (probably near the beginning of contact with one earth station), select one uplink event (pulse or zero-crossing) as corresponding to the present downlink event, giving an estimate $\hat{T}_0$ of the total two-way delay. We make this selection so that $\hat{T}_0$ agrees with our *a priori* knowledge of the satellite orbit, since no better information is available, resulting in an error equal to some unknown number of cycles of the uplink waveform. Because there is a fixed relationship between the uplink and downlink waveforms imposed by the satellite transponder, the error can be made to be always the same number of uplink cycles. For example, if there is a fixed number of downlink cycles per uplink cycle set by the transponder frequency ratio R , then when N downlink cycles have elapsed since t_0 we know that the corresponding uplink cycle is the one N/R cycles after the one that corresponded to t_0 . If our initial guess $\hat{T}_0$ of the two-way time was in error by some number

of uplink cycles, our present estimate is still in error by the same number.

However, the error will not be a constant amount of *time* unless the uplink period is constant, so it might seem reasonable to transmit a constant-frequency uplink signal. But there are good reasons not to do this. The distance to the spacecraft is changing rapidly, so the Doppler effect will cause rapid variation of the on-board timing. Although this could in principle be measured and later compensated, the variation is so large that second order effects and non-ideality in the spacecraft hardware make it impractical. (For a detailed explanation, see [3].) Therefore, a practical approach involves transmitting an uplink waveform that cancels the *a priori* uplink delay so that the on-board timing is nearly stable. The uplink is then only quasi-periodic; although the two-way time error remains a constant number of cycles, it is not a constant amount of time. We will show later that this effect is very small.

HARDWARE IMPLEMENTATIONS

We will soon look at a hardware implementation that achieves this strategy. But before doing so we will look at other possible implementations and discuss their shortcomings.

Figure 2. Block diagram of a straightforward two-way timing implementation at an earth station.

Consider the scheme in Figure 2. It involves separate synthesis of an uplink signal and an estimate of the downlink signal, each based on the best *a priori* estimate of the spacecraft orbit. Of course, the uplink signal being generated at time t is based on the expected position of the satellite at a slightly later time, and the downlink estimate is based on its assumed position slightly earlier. The actual downlink signal is then compared

against the synthesized estimate, and the time difference τ is measured.¹ The total two-way delay is then estimated as

$$\tilde{T}(t) = \hat{T}(t) + \tau(t) \quad (1)$$

where $\hat{T}(t)$ is the *a priori* estimate.² Notice that the error in $\hat{T}(t)$ may not be constant; as time goes on (from t_0), the satellite orbit may evolve relative to the predicted orbit. One might think that this evolution would be reflected in the measured $\tau(t)$ and that the total $T(t)$ could have constant error. This would indeed be the case if it were not for a subtle effect: the uplink compensation and the downlink prediction that apply to the signal being received at time t may be based on two different satellite positions. This is because the uplink compensation actually applied to that signal was the one transmitted at $t - T(t)$, whereas the one corresponding to the downlink prediction was transmitted at $t - \hat{T}(t)$. Since our *a priori* knowledge of the satellite position is not perfect ($T \neq \hat{T}$), an error results.³

How big is this effect? If the satellite range is in error by 1000 m, then the two-way delay is in error by $6.7 \mu\text{sec}$, so the uplink compensation disagrees with the downlink prediction by this amount. This is many cycles of the observing frequencies of interest, as well as many sampling times at the bandwidths of interest. It is then significant for VLBI signal processing, but perhaps not for orbit determination since the satellite will typically move only a few cm during the error time.

What if no uplink compensation were used? Then the residual $\tau(t)$ would include the total deviation of the two-way delay from the (single) prediction, and the above formula would produce an ambiguity error in $T(t)$ that remains constant. So the time-variation of the error can be attributed to the attempt to compensate the uplink.

Is there a way to avoid this problem? Yes; consider the design of Figure 3. Here we use only one prediction, from which we synthesize the uplink compensation. Rather than synthesize a predicted downlink signal separately, we use a delayed version of the uplink signal, scaled in frequency by the transponder ratio. In effect, we build at the earth station a simulation of the round-trip path. This would be equivalent to Figure 2 except for the fact that the delay is kept *constant*; it is left permanently at the value $\hat{T}_0$ that was our best estimate of the two-way delay at one selected time during the contact with one earth station. As the two-way delay changes during the rest of the contact, the uplink portion continues to be compensated as accurately as possible, but the downlink prediction will deviate more than before from the actual value. This will be reflected in the measured residual $\tau(t)$, which will now vary more rapidly. So finally our estimate of the two-way delay is

$$\tilde{T}(t) = \hat{T}_0 + \tau(t) = T(t) - \Delta\tau, \quad (2)$$

where $T(t)$ is the true two-way delay and $\Delta\tau$ is the error in our estimate. In either of these designs, the purpose of synthesizing the predicted downlink signal is to minimize the size of the residual, making it easier to measure. In the scheme of Figure 3 this is done less accurately (intentionally) than in the scheme of Figure 2. In return for allowing larger residuals, we avoid one type of timing error.

It should be apparent that the scheme of Figure 3 is equivalent to that illustrated in Figure 1, since the design guarantees that the number of elapsed cycles of the predicted

¹For quasi-periodic signals, this is actually a phase measurement, $\tau = \phi/(2\pi f)$ for phase ϕ at frequency f . We can nevertheless allow τ to be much larger than one cycle by making a continuous series of measurements and adding 2π to ϕ whenever a cycle is completed.

²Here and later, the notation $\hat{x}$ means an estimate of x based on real-time measurements, whereas $\hat{x}$ is an *a priori* estimate.

³The reader may wish to pause here and consider this point carefully. It should be clearly understood before continuing.

Figure 3. Block diagram of an improved two-way timing implementation.

downlink is always R times the number of elapsed cycles of the uplink. This is not true of the scheme of Figure 2.

However, it is not necessary to implement explicitly the $\times R$ and delay boxes in Figure 3. In fact, a convenient implementation might use two independent synthesizers and have a block diagram very much like Figure 2. It would differ only in the programming of the downlink synthesizer, which would be driven by the estimated spacecraft position at $t + \hat{T}/2 - \hat{T}_0$ rather than that at $t - \hat{T}/2$.

Recall that there is still a variation in the timing error due to the variation in uplink frequency during a period of contact. The timing error at the reference epoch t_0 can be written

$$\Delta\tau_0 = \frac{\Delta n}{f_0(1 + \dot{R}_0/c)} \quad (3)$$

where Δn is the error in uplink cycles, f_0 is the nominal uplink frequency, and R is the range; the denominator is then the Doppler-corrected uplink frequency. Similarly, the error at some other time is

$$\Delta\tau = \frac{\Delta n}{f_0(1 + \dot{R}/c)} \quad (4)$$

and the change in error is

$$\Delta\tau - \Delta\tau_0 = \frac{\Delta n}{f_0} \left(\frac{\dot{R}_0 - \dot{R}}{c} \right). \quad (5)$$

For Radioastron and VSOP, the largest $|\dot{R}_0 - \dot{R}|$ is about 10^4 m/s, and this occurs only rarely during a low perigee pass. If the *a priori* range error is 1000 m, then $\Delta n/f_0$ is 6.7 μsec , so

$$\max |\Delta\tau - \Delta\tau_0| = 2.2 \times 10^{-10} \text{ sec}. \quad (6)$$

Although this is several cycles at some desired observing frequencies (4.9 cycles at 22 GHz), it is quite negligible if the variation is spread over many coherence times of the stable reference oscillators. It is also small with respect to typical sampling periods (around 10^{-7} sec).

EFFECT OF EARTH STATION MOTION

The designs of Figures 2 and 3 neglect an important fact: the earth station moves during the two-way propagation time. This has no effect on the measurement discussed so far, namely that of the two-way delay. But, as mentioned in the Introduction, our real objective is to determine the uplink delay only, so as to be able to correct the on-board timing. The earth station motion makes the path non-reciprocal, and we now investigate how accurately this can be accounted for.

The two-way delay can be written $T(t) = T_u + T_d$ where T_u is the uplink delay and T_d is the downlink delay. Assuming that the earth station motion results only from the rotation of the earth, it moves a distance $D = \omega_E r_E T \cos l$ during two-way time T where ω_E is the angular velocity of the earth, r_E is the radius of the earth, and l is the station latitude. The component of this motion along the range direction is then $R_u - R_d = c(T_u - T_d) \approx D \cos \delta \sin h$ where δ, h are the declination and hour angle, respectively, of the satellite at the transponding time. The worst-case error in knowledge of this component is then $2D\sigma/(R_u + R_d)$ where σ is the transverse (sky plane) error in knowledge of the satellite position. The last expression is independent of range and is about $3 \times 10^{-6}\sigma$, or about 3 mm when the orbit is known to 1000 m; then $T_u - T_d \equiv \Delta T$ can be known to about 10^{-11} sec.

We thus have

$$2T_u(t) = T(t) + \Delta T(t) = \tilde{T}(t) + \Delta\tilde{T}(t) + \Delta\tau + \Delta\tau_e \quad (7)$$

where $\Delta\tau_e$ is the error in our correction for the earth station motion.

UPLINK TIME RESIDUAL

Although we now have [eqn (7)] a measure of the uplink delay whose errors are either constant or small, this is not sufficient to accomplish our purpose. We need to establish the time of certain events on the satellite. The event that arrives on the downlink at our initializing time t_0 can be assumed to have occurred at $t_0 - [\hat{T}_0 + \Delta T(t_0)]/2$. If the uplink compensation were perfect, then the times of later events could be found by counting the periodic downlink signal. If there were no uplink compensation, then we could again count the downlink signal, but we would then "correct" the time by adding $\hat{T}_u(t) - \hat{T}_u(t_0)$. The actual case is in between these, and is more difficult. Let $\Delta T_u(t)$ be the uplink residual, i.e., the difference between the actual uplink delay and that assumed for the uplink compensation; then

$$\Delta T_u(t) = T_u(t) - \hat{T}_u(t - T(t) + \hat{T}(t)) \quad (8)$$

where argument t indicates that the quantity applies to the event received on the downlink at time t . (In the following, the argument is t if not explicitly written.) Our best estimate of this residual is then

$$\Delta\tilde{T}_u = \tilde{T}_u - \hat{T}_u(t - \tilde{T} + \hat{T}). \quad (9)$$

We can then use this estimate to correct the event times obtained by counting the downlink. The error we make in the estimate (9) is given by subtracting (8) from (9) and substituting from (7):

$$\Delta\tilde{T}_u - \Delta T_u = (\Delta\tau + \Delta\tau_e)/2 + \hat{T}_u(t + \hat{T} - T) - \hat{T}_u(t + \hat{T} - T - \Delta\tau - \Delta\tau_e). \quad (10)$$

The last two terms are the *a priori* uplink time for the actual time of transmission and the estimated time of transmission, respectively. Although we still do not know accurately the time that the uplink corresponding to the present downlink was transmitted, we know it better than we did *a priori*, and we have made a first-order correction. If the *a priori* range error is 1000 m and the maximum satellite range rate is 10^4 m/s, then the remaining error is $< 2.2 \times 10^{-10}$ sec.

To calculate the estimate in (9), the earth station must remember the uplink prediction that was being used $\tilde{T}$ seconds ago, even though $\tilde{T}$ is not known until now. This can be done in several ways; the uplink prediction could be re-computed for that time, or the predictions could be put in a temporary buffer as they are transmitted.

SUMMARY

We have shown that the use of a Doppler-compensated uplink in the two-way timing transfer for present-generation OVLBI satellites may lead to an error in knowledge of the times of events on the satellite. This error can be avoided by a strategy that differs from a straightforward one in two ways: (a) the downlink phase detector should not measure the residual with respect to a predicted downlink signal, but rather with respect to the uplink signal after a *fixed* delay (Figure 3). (b) The uplink residual at the satellite should be estimated by using the measured change in two-way time to estimate the time of transmission of the uplink that corresponds to the present downlink [equation (9)].

An analysis of the errors in this strategy has shown that there are several:

- (i) The absolute one-way time is only as good as the *a priori* range, but this error can be kept constant to better than 2.2×10^{-10} sec.
- (ii) The earth station motion during the two-way time is not known perfectly with respect to the satellite, resulting in an error of up to about 10^{-11} sec.
- (iii) Because of (i), the time of transmission of the uplink corresponding to the present downlink is not perfectly known, and use of the present two-way data allows only a first-order correction, leaving a residual of up to 2.2×10^{-10} sec.

Thus, the time error can be kept constant to within a few hundred picoseconds. More straightforward methods allow variations of up to several microseconds.

The numerical results quoted here are based on an *a priori* range error of 1000 m and a maximum range rate of 10^4 m/sec.

REFERENCES

- [1] D'Addario, L. R., 1991, "Time Synchronization in Orbiting VLBI," accepted by **IEEE Trans. on Instrumentation & Meas.**, scheduled for June 1991 (preprints available from the author).
- [2] D'Addario, L. R., 1990, "OVLBI Timekeeping With the VLBA Formatter," OVLBI Earth Station Memo No. 3, 1 Aug 1990 (available from NRAO, Charlottesville, VA 22903).
- [3] Linfield, R. and J. Springett, 1991, "The Need for Doppler Compensation in the Phase Uplink for VSOP and Radioastron," OVLBI Earth Station Memo No. 9, 10 Jan 1991 (available from NRAO or JPL).