

Synchronizing Telecommunications Networks

Basic Concepts



Table of Contents

I. Introduction	4
II. The Need for Synchronization	5
Background.....	5
Impact of Slips on Services.....	6
SDH and SONET Synchronization Needs	7
Synchronization-Caused Error Bursts	8
Synchronization Performance Objectives —Public Network.....	9
Synchronization Performance Objectives —Private Network.....	9
III. Synchronization Architecture.....	10
Major Methods for Synchronization	10
Plesiochronous	10
Hierarchical Source-Receiver	10
Mutual Synchronization	10
Pulse Stuffing.....	10
Pointers	11
Telecommunications Synchronization	11
Source Clocks: Primary Reference Source	12
Receiver Clocks	12
Clock Standards	13
IV. Synchronization Performance	16
Primary Reference Source Contribution	16
Facility Performance	16
Receiver Clock Contribution.....	17
Ideal Operation	17
Stressed Operation — Network Clocks	18
Stressed Operation — CPE Clocks	19
Holdover Operation	20
Interface Standards.....	20
V. Introduction to Synchronization Planning	20
Basic Concepts	20
Planning Issues	22
VI. Conclusion	22
VII. References	23

I. Introduction

With the rapid deployment of digital switching systems and transmission facilities and the introduction of SDH and SONET, the importance of synchronization in telecommunications has dramatically increased. New services and applications are also placing increased demands on the performance and operation of the synchronization network. Stringent synchronization performance and planning is required, not only to avoid unacceptable performance, but to mitigate latent, costly, hard-to-find problems, and to reduce subtle interdependencies among networks of various administrations.

This application note provides an introduction to network synchronization. Section II provides the background for synchronization and illustrates the need for network synchronization. Several synchronization-related impairments, such as slips, misframes, and error bursts, are introduced. The impact of these impairments on the performance of various services and applications is discussed.

Section III describes the various synchronization architectures that are used to maintain acceptable synchronization performance. The primary reference source and receiver clocks are introduced. The functionality of these clocks is presented, along with the relative importance of each function to network synchronization performance and planning. Section III concludes with a discussion on clock requirements of ETSI, ANSI and ITU.

Synchronization performance is considered in Section IV. The contribution to synchronization performance of primary reference sources, synchronization transport facilities, and receiver clocks is presented. It is shown that receiver clocks typically operate at a different frequency than the primary reference source to which they are locked. This frequency offset of the receiver clock has the greatest impact on network synchronization performance.

Section V covers the basic concepts of network synchronization planning. The most common planning issues are also discussed.

References are cited throughout this application note and included in brackets []. A complete listing of references is on page 23.

II. The Need for Synchronization

Background

Synchronization is the means of keeping all digital equipment in a communications network operating at the same average rate. For digital transmission, information is coded into discrete pulses. When these pulses are transmitted through a network of digital communication links and nodes, all entities must be synchronized. Synchronization must exist at three levels: bit, time slot, and frame.

Bit synchronization refers to the requirement that the transmit and receive ends of the connection operate at the same clock rate, so that bits are not misread. The receiver may derive its timing from the incoming line to achieve bit synchronization. Bit synchronization involves timing issues such as transmission line jitter and ones density. These issues are addressed by placing requirements on the clock and the transport system.

Time slot synchronization aligns the transmitter and receiver so that time slots can be identified for retrieval of data. This is done by using a fixed frame format to separate the bytes. The main synchronization issues at the time slot level are reframe time and framing loss detection.

Frame synchronization refers to the need of the transmitter and receiver to be phase aligned so that the beginning of a frame can be identified. The frame in a DS1 or E1 signal is a group of bits consisting of twenty four or thirty bytes, or time slots, respectively, and a single framing pulse. The frame time is 125 microseconds. The time slots are associated with particular circuit users.

A network clock located at the source node controls the rate at which the bits, frames, and time slots are transmitted from the node. A second network clock is located at the receiving node, controlling the rate that the information is being read. The objective of network timing is to keep the source and receive clocks in step, so that the receiving node can properly interpret the digital signal. Differences in timing at nodes within a network will cause the receiving node to either drop or reread information sent to it. This is referred to as a slip.

For example, if the equipment that is sending information is operating with a clock rate that is faster than the receiving equipment's rate, the receiver cannot keep up with the flow of information. When the receiver cannot follow the sender, the receiver will periodically drop some of the information sent to it. The loss of information is referred to as a slip of deletion.

Similarly, if the receiver is operating with a clock rate faster than the sender, the receiver will duplicate information, so that it can continue to operate at

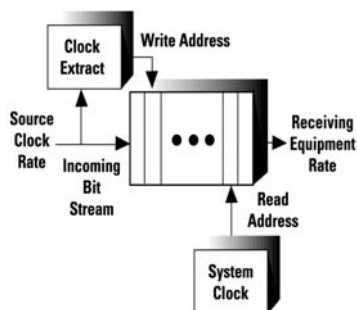


Figure 1. Slip Buffer

IMPACT OF SLIPS	
Voice	Occasional audible clicks
Faxes	Distorted lines
Voiceband Data	Corrupted data
Video	Frame freeze
Encrypted Data	Loss of communications
SONET/SDH	Pressure on pointer budgets. Impairment at PDH boundary.

Poor synchronization affects quality of service. The impact ranges from annoying for voice services to disastrous for encrypted services.

its speed and still communicate with the sender. This duplication of information is called a slip of repetition.

In DS1 and E1 communications, buffers are used to control slips (see Figure 1). The data is clocked into the receiving equipment's buffer at a rate determined by the source end's clock rate. Data is read from the buffer using the receiving equipment's clock. Buffers of varying sizes are used. Typically, the buffer will hold more than one frame of data. In this case, the receiving equipment will drop or repeat an entire frame of data when it slips. This is called a controlled slip.

The basic objective of network synchronization is to limit the occurrence of controlled slips. Slips can occur for two basic reasons. The first is the lack of frequency synchronization among the clocks in the connection, resulting in differences in clock rates. The second is phase movement either on the communications link (such as jitter and wander) or between the source and receiver clock. The latter, phase movement between the source and receiver clock, will be shown to be the largest contributor to slips in communication networks.

Slips, however, are not the only impairment caused by lack of synchronization. In SDH and SONET networks, poor synchronization can lead to excessive jitter and misframes in the transport of digital signals, as discussed in "SDH and SONET Synchronization Needs" on page 7. In private networks, the poor synchronization of customer premises equipment (CPE) can cause error bursts in the network. (See "Synchronization-Caused Error Bursts" page 8.) Therefore, even though minimizing slip rate remains the foremost objective of synchronization, the control of other synchronization-related impairments needs to be considered in the design of a synchronization network.

Impact of Slips on Services

The impact of one or more slips on services carried on digital networks is dependent on the application [1-6]. The effect of a single slip on various services is described below.

For voice service, studies [1] indicate that slips may cause an occasional audible click. This click is not always heard and is not a serious impairment for speech. Therefore, voice services are tolerant of slips. Slip rates up to several slips per minute are considered acceptable.

A study [2] conducted to determine the effects of controlled slips on Group 3 facsimile transmission found that a single slip caused distortion or missing lines in the facsimile. A slip caused up to eight horizontal scan lines to be missing. This corresponds to a missing 0.08 inches of vertical space. In a standard typed page, a slip would be seen as the top or bottom half of a typed line missing. If slips continued to occur, the affected pages would need to be retransmitted. This retransmission must be initiated by the user and is not automatic.

The impact of a slip on voiceband data is to cause a long burst of error [3]. The duration of this error burst is dependent on the data rate and modem type and ranges from 10 milliseconds to 1.5 seconds. During this errored period, the receiving terminal device connected to the modem receives corrupted data. As a result, the user needs to retransmit the data.

When a slip occurs during a video phone session, the video portion of the call is lost. The callers are required to re-establish the video portion.

The impact of a slip on digital data transmission depends on the protocol used for the transmission. In protocols without retransmission capabilities, there will be missing, repeated, or errored data. Misframes may occur resulting in many frames of data being corrupted while the framing pulse is regained. Retransmitting protocols are able to detect the slip and will initiate a retransmission. Retransmission typically requires one second to initiate and accomplish. Therefore, slips will impact the throughput of the application, typically causing a loss of a second of transmission time.

For digital video transmission (video teleconferencing, for instance), tests [4] indicate that a slip usually causes segments of the picture to be distorted or to freeze for periods of up to 6 seconds. The seriousness and length of the distortion is dependent on the coding and compression equipment used. The impairment is most serious for low bit rate encoding equipment.

Encrypted services are greatly impacted by slips [4]. A slip results in the loss of the encryption key. The loss of the key causes the transmission to be unintelligible until the key can be resent and communications reestablished. Therefore, all communications are halted. More importantly, requiring retransmission of the key adversely affects security. For many secure applications, more than one slip per day is considered unacceptable.

SDH and SONET Synchronization Needs

With the introduction of SDH and SONET, new requirements and demands are being placed on network synchronization. SDH and SONET are high-speed, synchronous transport systems. SDH and SONET network elements require synchronization, since the optical signal they transmit is synchronous. If the SDH/SONET network elements lose synchronization, they will not cause slips, however. This is due to the fact that the payload in SDH and SONET is transmitted asynchronously. SDH and SONET use pointers to identify the beginning of a frame. A mismatch in the sending and receiving rate would cause a change in the pointer (see Figure 2).

A pointer adjustment, however, can cause jitter and wander in the transported signal. Jitter is a fast (≥ 10 Hz) change in the phase of a signal. Wander is slow (< 10 Hz) phase change. Excessive jitter from SDH/SONET can cause misframes (loss of frame synchronization). Excessive wander can cause terminating equipment to slip. Therefore, the goal of network synchronization in an SDH/

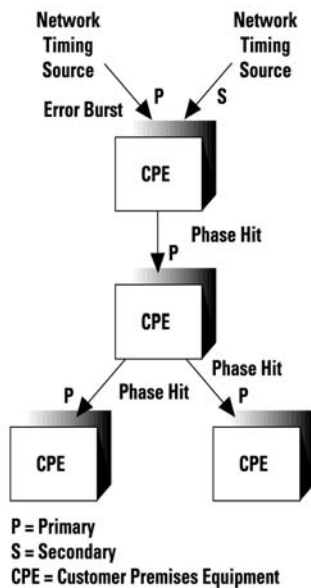


Figure 3. Cascading Errors in Private Networks

SONET network is to limit the number of pointer adjustments made by the SDH/SONET network elements. This is achieved by limiting the short term (<100 second) noise in the synchronization network by using better network clocks throughout the network.

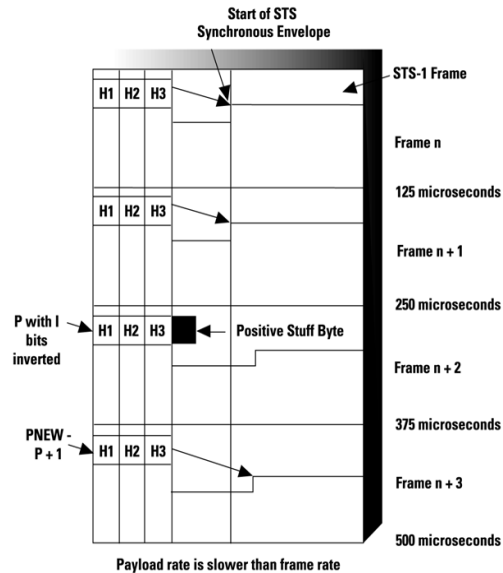


Figure 2. SONET Pointer Adjustment

Synchronization-Caused Error Bursts

In private networks, synchronization can cause additional impairments in the form of error bursts. Consider a private network where Customer Premises Equipment (CPE) clocks are chained. A short impairment on the timing reference of the first CPE clock in the chain will impact all equipment and facilities in the chain (see Figure 3). In response to a short error, most CPE clocks will produce error bursts on all outgoing lines [4]. Thus, a short impairment on the reference of the first clock will cause the first clock to produce an impairment on all of its outgoing lines. The second clock in the chain will see the impairment caused by the first clock and will react in the same manner, producing an outgoing impairment on all of its lines. In this manner, the error burst propagates (and occasionally multiplies) through the CPE network.

Synchronization-caused error bursts are short and transient in nature and are usually indistinguishable from excessively errored transmission lines. Thus, synchronization problems can be mistaken for high line error rates. These performance difficulties can be avoided by using properly designed CPE clocks and by carefully planning the synchronization distribution in the private network. It should be noted that these error burst problems typically do not occur in public networks.

Synchronization Performance Objectives — Public Network

Several synchronization performance objectives have been established by ITU and ANSI to control slip rates, pointer adjustment events, and synchronization-caused error bursts.

For international connections, the slip rate threshold for an “acceptable” connection is set by ITU at one slip every five hours. To achieve satisfactory slip rates on an end-to-end basis, the maximum long-term frequency inaccuracy allowed at the output of a digital system clock is 1×10^{-11} . This requirement was established both by ANSI [7] and by ITU [9, 10]. Short term requirements allow for 1 to 10 microseconds of time error in a day at the output of each network clock [7, 9].

New short-term requirements are being adopted [7]. This serves two purposes. First, it ensures that random variations in timing will not produce slips. Second, it limits the short-term stability of the timing signal which, in turn, limits the number of pointer adjustments and the resulting jitter in SDH/SONET networks. ANSI requires that the band-limited short-term noise at the output of a clock not exceed 100 nanoseconds [7].

Synchronization Performance Objectives — Private Network

There is a specification in the draft stage from ETSI [8] which provides jitter and wander requirements for the synchronization network suitable for SDH and PDH. Limits for different layers of the synchronization network as well as performance of clocks for SDH equipment are established. This document provides standards for those administrations who follow ETSI.

There are few synchronization performance objectives for private networks. The synchronization performance of a private digital network can be more than 1000 times worse than a public switched network [4]. ANSI requires that the first CPE in the synchronization chain in a private network meet 4.8 milliseconds of time error in a day. This corresponds to approximately 40 slips per day per CPE. In addition, ANSI currently has no objectives limiting the number of synchronization-caused error bursts in a private network. These are interim requirements, however. In the next few years, these objectives are expected to change to 18 microseconds of daily timing error and no synchronization-caused error bursts.

The major cause of this poor performance in private networks is the reliance on poor quality stratum 4 CPE clocks. In addition, private networks can have complex and unconstrained architectures, with great amounts of cascading of the timing reference. With Stratum 4 clocks, slips are caused not only by transmission errors but by equipment-induced impairments. In addition, CPE synchronization can be a significant source of errors on transmission facilities in a private network. This is discussed further in Section IV under “Receiver Clock Contribution, Stressed Operation — CPE Clocks” (page 18).

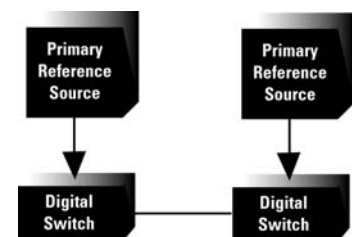


Figure 4. Plesiochronous Operation

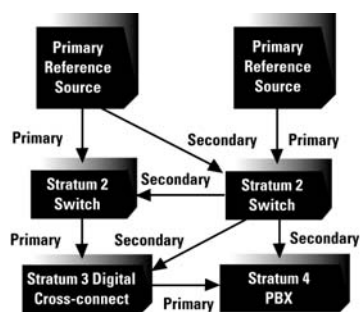


Figure 5. Hierarchical Source-receiver Operation

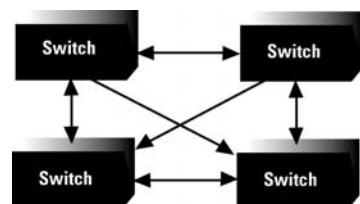


Figure 6. Mutual Synchronization Operation

III. Synchronization Architecture

Major Methods For Synchronization

There are several major methods used to synchronize digital networks: plesiochronous operation, hierarchical source-receiver operation, mutual synchronization, pulse stuffing, and pointers. These are defined below.

Plesiochronous

Each node receives a reference from a different independent timing source (Figure 4). Tolerable slip rates are maintained due to tight timing accuracy of either side of the connection. Standards place a boundary on the accuracy of clocks used to time plesiochronous connections. In networks that use plesiochronous situations, controlling clocks must maintain long-term frequency accuracy to 1×10^{-11} . This mode of operation is typical for connection across administration boundaries.

Hierarchical Source-Receiver

A primary reference source at a master node generates a reference clock that is shared and distributed (Figure 5). The source node sends its reference to receiver nodes. The reference clock is hierarchically distributed throughout the network. The two major components of this network are the receiver clocks used to regenerate the reference clock and the digital paths used to transmit clocks through the network.

Mutual Synchronization

In mutual synchronization, clocking information is shared by all nodes in the network (Figure 6). Each clock sends and receives a timing reference to (from) all other clocks in the network. Network timing is determined by each clock by averaging the synchronization signals it receives from all other clocks in the network. This operation can theoretically provide identical timing signals to each node, but in actual application, with imperfect clocks and imperfect transmission of timing information, the timing fluctuates as it hunts for a common frequency.

Pulse Stuffing

This method is used to transmit asynchronous bit streams above the DS1/E1 level. The bit streams to be multiplexed are each stuffed with additional dummy pulses. This raises their rates to that of an independent local clock. The outgoing rate of the multiplexer is higher than the sum of the incoming rates. The dummy pulses carry no information and are coded for identification. At the receiving terminal, the dummy pulses are removed. The resulting gaps in the pulse stream are then removed, restoring the original bit stream.

Pointers

This method is used by SDH and SONET to transmit payloads that are not necessarily synchronous to the SDH/SONET clock. Pointers are used to indicate the beginning of a frame in the payload. Frequency differences between SDH/SONET network elements or between the payload and the SDH/SONET equipment are accommodated by adjusting the pointer value (see Figure 7). Therefore, the payload need not be synchronized to the SDH/SONET equipment. SDH/SONET equipment are usually synchronized so that the number of pointer adjustments are kept to a minimum. This is desirable since each pointer adjustment will cause jitter and wander on the payload.

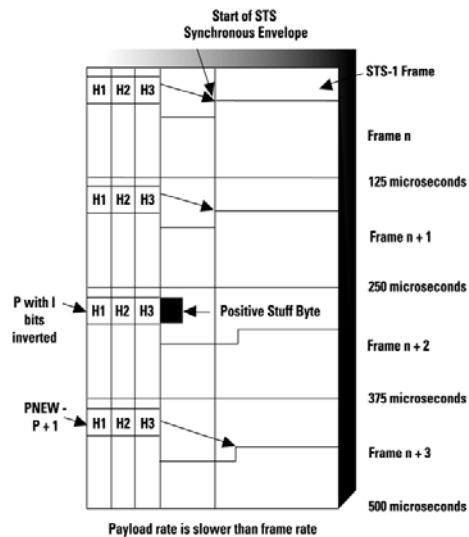


Figure 7. SONET Pointer Adjustment

Telecommunications Synchronization

Most telecommunication administrations use the hierarchical source-receiver method to synchronize its E1/DS1 network. The master clock for a network is one or more Primary Reference Sources. This clock reference is distributed through a network of receiver clocks (Figure 5).

A node with the most stable, robust clock is designated as a source node. The source node transmits a timing reference to one or more receiver nodes. Receiver nodes usually have equal or worse performance than the source node. The receiver node locks onto the timing reference of the source node and then passes the reference to other receiver nodes. Timing is thereby distributed down a hierarchy of nodes.

Receiver nodes are usually designed to accept two or more references. One reference is active. All other alternate references are standby. In the case where the active reference is lost, the receiver node can switch references and lock to an alternate reference. Thus, each receiver node has access to

timing from two or more source nodes. Most networks are engineered so that all receiver clocks are given two or more diverse references. In private networks, this may not be possible due to limited connectivity between nodes.

Clocks are placed into the hierarchy based upon performance levels. ANSI [7] designates performance levels as stratum levels: stratum 1, 2, 3, 4E, and four, in order of best performance to worst. ITU [9] designates four performance levels: primary reference source, transit node, local node, terminal or CPE node. Stratum 1 or primary-reference sources are master nodes for a network. Stratum 2 or transit-node clocks are typically found in toll switching and some digital cross-connect equipment. Local switching, most digital cross-connect systems, and some PBXs and T1 multiplexers have stratum 3 or local-node clocks. Most T1 multiplexers, PBXs, channel banks, and echo cancellers incorporate stratum 4 or CPE clocks.

Source Clocks: Primary Reference Source

A Primary Reference Source (PRS) is a master clock for a network that is able to maintain a frequency accuracy of better than 1×10^{-11} [7]. One class of PRS is a stratum 1 clock. A stratum 1 clock, by definition, is a free running clock [7]. It does not use a timing reference to derive or steer its timing. Stratum 1 clocks usually consist of an ensemble of Cesium atomic standards.

However, a PRS need not be implemented with primary atomic standards [7]. Other examples of PRS are Global Positioning System (GPS) and LORAN-C clocks. These systems use local rubidium or quartz oscillators that are steered by timing information obtained from GPS or LORAN-C. They are not considered Stratum 1 since they are steered, but are classified as primary reference sources. These clocks are able to maintain an accuracy within a few parts in 10^{-13} to a few parts in 10^{-12} .

The slip rate contribution of a PRS is usually negligible. A network which derives timing from two PRS clocks will experience at most five slips per year, caused by the inaccuracy of the clocks. This is negligible compared to the performance of receiver clocks, as will be seen in Section IV. Therefore, it has been the trend of telecommunication network operators to rely more on PRS clocks and to use multiple PRS clocks to time their network.

Receiver Clocks

The major role of a receiver clock is to recover clocking from a reference signal and maintain timing as close to the source node's timing as possible. This requires that the receiver clock performs two basic functions. First, it must reproduce the source clock's timing from a reference signal, even though the reference may be errored. Second, it must maintain adequate timekeeping in the absence of a timing reference.

The usual mode of operation of a receiver clock is extracting timing from the source clock's reference. In this mode, the receiver clock must be able to handle short reference errors that may occur. These errors may be timing instabilities (jitter) or short reference interruptions (error bursts). These errors are usually caused by the facility transporting the reference from the source clock to the receiver clock.

A receiver clock uses low-pass filters to handle short term timing instabilities. For short interruptions, the receiver clocks are designed to have two or more references so that it can switch references under short-term impairments. Most network clocks (ANSI stratum 2, 3, and 4E, ITU transit and local clocks) are designed to cause no more than 1,000 nanoseconds of timekeeping error with each reference switching or other transient event [7, 10]. In addition, network clocks are designed to hold daily time-keeping to within 1 to 10 microseconds in the absence of interruptions.

Stratum 4 (CPE) clocks do not have any requirements for their timing recovery mode of operation. In response to short interruptions, a Stratum 4 clock will typically cause 10-1000 microseconds of time-keeping error. In addition, an error burst will accompany this phase jump. Therefore, CPEs are very intolerant of facility errors. (See Section IV "Receiver Clock Contribution, Stressed Operation — CPE Clocks," page 18, for typical stratum 4 performance).

The second mode of operation is a receiver clock running with a loss of all of its timing references. Holdover is the capability to remember the last known source frequency and maintain frequency accuracy after all timing reference is lost. All clocks, other than CPE, are required to have holdover capability. CPEs are allowed to enter free-run mode when it loses all timing reference. Free-run mode refers to a mode of operation where the clock's timing is controlled by the local oscillator and no memory of an external reference is used to correct the oscillator frequency.

Clock Standards

ITU and ANSI classify receiver clocks into levels based on performance. ITU designates clocks as transit, local, and CPE/terminal clocks. ANSI designate clocks as stratum 2, 3, 4E, and 4, in decreasing order of performance. In order to meet a certain level of performance, a clock must meet requirements for several functions. These are: rearrangement timekeeping, holdover, free-run accuracy, hardware duplication, and external timing capabilities. These functions are summarized in Tables 1 and 2.

Table 1. ITU Clock Standards

FUNCTION	TRANSIT NODE	LOCAL NODE	TERMINAL (CPE)
Accuracy	No requirement	No requirement	5×10^{-5}
Holdover			
Init. Freq Offset	5×10^{-10}	1×10^{-8}	Not required
Long Term	1×10^{-9}	2×10^{-8}	
Time Interval Error	1 microsec.	1 microsec.	No requirement
Phase Change Slope	61 ppm	61 ppm	No requirement

Table 2. ANSI Clock Standards

FUNCTION	STRATUM 2	STRATUM 3	STRATUM 4E	STRATUM 4
Accuracy	1.6×10^{-8}	4.6×10^{-6}	3.2×10^{-5}	3.2×10^{-5}
Holdover	1×10^{-10}	3.7×10^{-7}	Not required	
Time Interval Error	1 microsec.	1 microsec.	1 microsec.	Not required
Phase Change Slope	61 ppm	61 ppm	61 ppm	Not required
Duplication	Required	Required	Not required	Not required
External Inputs	Required	Required	Not required	Not required

Rearrangement time-keeping capability is the most important requirement in receiver clocks. This is because receiver clocks can often experience short interruptions of its timing reference. The short interruption will cause the clock to undergo a rearrangement. A rearrangement is defined as a clock switching its reference or bridging a short duration error. Clock hardware side switching is also considered a rearrangement. Under rearrangement conditions, all clocks, except stratum 4 CPE clocks, must cause no more than 1 microsecond of timing error with respect to its timing source. In addition, when the clock causes the timing error, it cannot adjust the phase quickly. The phase must change with a slope of less than 61 ppm. The phase change slope requirement is necessary so that down-stream clocks can remain locked to the clock undergoing the rearrangement.

Holdover requirements vary dramatically between network clocks. A stratum 2 and a transit node are allocated a frequency inaccuracy of 1×10^{-10} and 1×10^{-9} after the first 24 hours of reference outage [7, 10]. These strict specifications are required since these clocks are typically used to control the timing in toll offices which have tens of thousands of circuits. This specification ensures that no circuit experiences more than a single slip in the first 24 hours of holdover. In contrast, since stratum 3 and local clocks typically are deployed in small offices and impact fewer circuits, they are allowed up to 255 and 14 slips, respectively, on every circuit during the first 24 hours.

Stratum 4 CPE clocks are not required to have holdover. A stratum 4 clock without holdover will immediately enter free run condition whenever the timing reference is lost.

The free run condition refers to a clock's stability when it is operating on its own internal oscillators without being steered or corrected by a history of an external reference. For clocks with holdover, the free run mode of operation is observed only in an extended reference outage (weeks to months) and is extremely rare. Thus, the free run accuracy specification is the least critical of the clock specifications. This point is highlighted further by the fact that ITU does not specify free run accuracy. For stratum 4 CPE clocks, its free run will determine its slip performance during even a short loss of reference.

Additional requirements are that ANSI stratum 2 and 3 clocks must have duplicated hardware and external clock inputs. Duplicated hardware ensures that the equipment continues to operate during a hardware failure of the clock. An external clock input refers to a dedicated-for-timing clock input. This is used to feed timing directly into a clock. This input is useful for flexible synchronization planning, where the timing reference for a clock may not terminate on the digital system.

IV. Synchronization Performance

The synchronization performance in a hierarchical source-receiver network is characterized by three components: the accuracy of the master clock, the performance of the facilities distributing the reference, and the performance of the receiver clocks obtaining a reference over the facility. It will be shown that synchronization inaccuracy of the master clock usually contributes a small portion of the timing inaccuracies in a synchronization network. Synchronization performance is dominated by a combination of the facility and receiver clock performance. In actual networks, a receiver clock, locked to the master clock, will operate with a long-term frequency that is different than the master clock. The frequency inaccuracy of a receiver clock is typically 10-100 times the inaccuracy of the master clock. Therefore, receiver clocks contribute the largest portion of timing errors and slips in a network.

Primary Reference Source Contribution

The slip rate contribution of a PRS is usually negligible. Cesium, GPS, and LORAN-C will typically have long term accuracies on the order of a few parts in 10^{-13} to a few parts in 10^{-12} . This results in slip rates ranging from a slip every five years to three slips per year. This is a small fraction of the five slips per day goal for an end-to-end connection and can usually be considered negligible.

Facility Performance

There are two major factors in determining a facility's performance for transporting timing reference. They are errors and timing instabilities (jitter and wander).

A facility used for timing reference can have a significant number of disruption events. The number of error burst events can range from an average of 1 to 100 events per day depending on facility type, mileage and other factors. For example, the ITU objective for end-to-end severely errored seconds (SES) performance is 175 per day [11]. An SES is a second of transmission when at least 320 CRS-6 errored events occur. This is roughly equivalent to having a bit error rate of 1×10^{-3} for the duration of the second. Performance objectives in ANSI are between 40 and 50 SES per day, depending on mileage.

These constant degradations will adversely affect the distribution of timing reference. As previously discussed, a receiving clock will react to each error. The clock is allowed to move up to 1 microsecond in response to each error on its timing reference. The accumulation of facility errors and the resulting phase error in the receiving clock will greatly impact the slip rate in a network and can lead to tens of microseconds of phase movement per day if the network is planned poorly.

Timing instabilities on a reference depend on the technology used by the facility to transport the reference. If the reference is carried asynchronously (e.g., by DS3 transmission), the reference will have jitter typically less than 600 nanoseconds in magnitude and insignificant amounts of wander. These levels are usually not a concern.

References passed over satellite will have excessive wander. This is caused by small movements of the satellite from its geostationary position. The magnitude of the wander is typically 1.8 milliseconds per day. This makes satellite transmission unsuitable for use as a timing reference.

References passed as payload through SDH/SONET can have significant amounts of wander. A DS-1 or E1 signal, mapped and transported through SDH/SONET, can experience tens of microseconds of wander per day [12]. Therefore, timing is never passed as payload through SDH/SONET. In networks that use SDH/SONET transport, the optical carrier is used to transport timing since it does not experience pointer adjustments and the resulting jitter and wander.

Receiver Clock Contribution

A receiver clock is a clock whose timing output is controlled by the timing signal received from a source clock of equal or higher quality. As stated above, receiver clocks must reproduce the source clock's timing from a reference signal, even though the reference may be errored, and it must maintain adequate time keeping in the absence of all timing references. The receiver clock performance can be characterized by its operation in three scenarios [13]:

- Ideal Operation
- Stressed Operation
- Holdover Operation

Ideal operation describes the short term behavior of the clock and is important to control pointer adjustments in SDH and SONET networks. Stressed operation is the typical mode of operation of a receiver clock, where a receiver clock is expected to receive timing from a source clock over a facility that has short term impairments. Finally, holdover operation characterizes the clock's performance in the rare case when all timing references to the clock are lost.

Ideal Operation

In ideal operation, the receiver clock experiences no interruptions of the input timing reference. Even though this is not typical of real network operation, understanding a clock's performance under ideal operation gives bounds for the clock's performance. It is also important to limit the short term noise of a clock. A clock's short term noise will impact the occurrence

of pointer adjustments in SDH/SONET networks, and the resulting SDH/SONET payload jitter and wander.

Under ideal conditions, the receiver clock should operate in strict phase lock with the incoming reference. For short observation intervals less than the time constant of the Phase Locked Loop (PLL), the stability of the clock is determined by the short term stability of the local oscillator as well as quantization effects and PLL noise. In the absence of reference interruptions, the stability of the output timing signal behaves as white noise phase modulation. The high frequency noise is bounded and uncorrelated (white) for large observation periods relative to the tracking time of the PLL.

Stressed Operation — Network Clocks

This category of operation reflects the performance of a receiver clock under actual network conditions where short interruptions of the timing reference can be expected. As described in Section IV, “Facility Performance” page 16, these interruptions are of short duration in which the timing reference time is not available. The number of interruptions can range from 1 to 100 per day.

All interruptions will affect the receiver clock. During the interruption the timing reference cannot be used. When reference is restored or if the interruption persists and clock switches references, there is some error regarding the actual time difference between the local receiver clock and the newly restored reference. The timing error that occurs due to each interruption depends on the clock design, but should be less than 1 microsecond [7, 10]. This random timing error will accumulate as a random walk, resulting in a white noise frequency modulation of the receiving clock’s timing signal.

In addition to the white noise frequency modulation, interruption events can result in a frequency offset between the receiver clock and the source clock. This is due to a bias in the phase build-out in the receiver clock, when reference is restored. The amount of bias is dependent on the clock design. The magnitude of this bias plays a crucial role in the long term synchronization performance of the receiver clock.

This bias will accumulate through a chain of receiver clocks. The end result is that there will be a frequency offset between all clocks in a synchronization chain. The magnitude of the frequency offset grows with the number of clocks in the chain. Therefore, in actual network conditions, receiver clocks will operate with a slightly different long term frequency than the primary reference clock. The magnitude of this frequency offset is a function of the performance capabilities of the receiver clock (its timing error bias during rearrangements) and the number of short interruptions (SES) on the facility carrying the reference.

It is this long term frequency offset, caused by short term facility impairments and receiver clock bias, that is the major cause of slips in a network. The long term frequency offset can vary from a few parts in 10^{-12} to a few parts in 10^{-10} , depending on the network configuration and on clock and facility performance. This frequency offset is several orders of magnitude worse than the frequency difference between two primary reference sources. For this reason, there is a growing tendency among network operators to install multiple primary reference sources in their network and to limit the amount of cascading of timing reference takes through the network.

Stressed Operation — CPE Clocks

Under stressed conditions, stratum 4 CPE clocks perform very differently than other network clocks. This is due to the fact that most CPE clocks do not incorporate a phase build-out routine to limit the time-keeping error that occurs during the short interruption. Most CPE clocks perform poorly in response to a short error on its timing reference.

When a stratum 4 clock experiences a short interruption, it will declare the reference unusable and will switch its reference to a backup timing source. This back-up source may be either another timing reference or its internal oscillator. During this switch of reference, the clock will typically produce a large, fast phase hit of 10 to 1000 microseconds. The magnitude of this hit is often large enough to cause multiple slips. This phase hit occurs on all outgoing lines of the CPE.

Downstream clocks are unable to remain locked to a reference with such a phase hit. To the downstream device, the phase hit is indistinguishable from a facility error. As a result, the downstream clock will switch its reference, cause another phase hit, and the error event propagates. Therefore, one error on a facility at the top of the synchronization chain can cause all lines and nodes in the synchronization chain to have errors (see Figure 8).

The performance of private networks using stratum 4 clocks is typically poor. It can be 1000 times worse in performance than is seen in public networks, operating at an effective long-term frequency accuracy of 1×10^{-9} to 1×10^{-7} . Slip performance of dozens of slips per day per CPE is not unusual. In addition, the phase hits caused by poor CPE synchronization appear as transmission errors. CPE synchronization can cause up to hundreds of transmission errors per day. Excessive transmission errors in private networks is a common symptom of poor synchronization performance [14, 15].

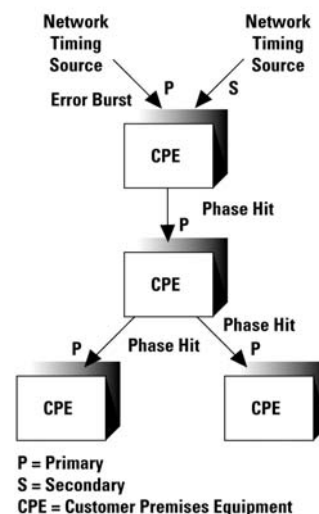


Figure 8. Cascading errors in private networks

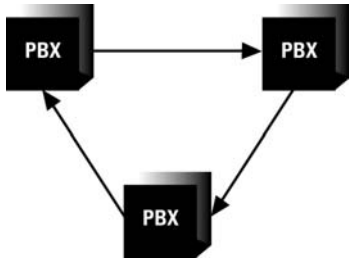


Figure 9. Timing Loop

Holdover Operation

A receiver clock will operate in holdover in the rare cases that it loses all its timing references for a significant period of time. There are two major contributors to holdover performance: initial frequency offset and frequency drift. Initial frequency offset is caused by the settability of the local oscillator frequency and the noise on the timing reference when the clock first enters holdover. Frequency drift occurs due to aging of the quartz oscillators. ITU clock requirements bound both contributors to holdover performance separately. ANSI holdover requirements apply to the aggregate performance.

Interface Standards

Current clock standards do not ensure acceptable operation under stress conditions. ANSI and ITU interface requirements are designed to apply to ideal operation only. Under ideal operation, daily time-keeping error is to be held to 1 to 10 microseconds and long term frequency offset should be less than 1×10^{-11} . However, since stress operation performance is undefined, actual network performance is allowed to be poorer than 1×10^{-11} long term.

V. Introduction to Synchronization Planning

The role of synchronization planning is to determine the distribution of synchronization in a network and to select the clocks and facilities to be used to time the network. This involves the selection and location of master clock(s) for a network, the distribution of primary and secondary timing throughout a network, and an analysis of the network to ensure that acceptable performance levels are achieved and maintained.

Basic Concepts

To achieve the best performance and robustness from a synchronization network, several rules and procedures must be followed. Some of the most important are avoiding timing loops, maintaining a hierarchy, following the BITS concept, using the best facilities for synchronization reference transport, and minimizing the cascading of the timing reference.

Timing loops occur when a clock uses a timing reference that is traceable to itself (Figure 9). When such loops occur, the reference frequency becomes unstable. The clocks in a timing loop will swiftly begin to operate at the accuracy of the clock's pull-in range. This will result in the clock exhibiting performance many times worse than it does in free-run or holdover mode. Therefore, it is important that the flow of timing references in a network be designed such that timing loops cannot form under any circumstance. No combination of primary and/or secondary references should result in a timing loop. Timing loops can always be avoided in a properly planned network.

Maintaining a hierarchy is important to achieve the best possible performance in a network. Under ideal or stress conditions, passing timing from better to worse clocks will maximize performance. Synchronization will still be maintained in normal operation if timing is passed from a worse clock to a better clock. Only performance may suffer slightly, since a better clock is more immune to short term network impairments and will accumulate less timing error. It is only in the case where an upstream clock enters holdover or free run that non-hierarchy causes major problems. In this case, the poorer performing upstream clock in holdover may have a frequency accuracy worse than the downstream clock can lock to. The downstream clock would not remain locked and will also go into holdover. This results in multiple clocks being in holdover and excessive slips in the network.

Most administrations follow the Building Integrated Timing Supply (BITS) or SSU concept for synchronization distribution (Figure 10). In the BITS or SSU method, the best clock in an office is designated to receive timing from references outside the office. All other clocks in the office are timed from this clock. In many cases the BITS or SSU is a timing signal generator, whose sole purpose is for synchronization. Other administrations rely on clocks in switches or cross-connect systems for the BITS or SSU. The BITS or SSU clock should be the clock that is best performing in stress and holdover and is the most robust. With the BITS or SSU concept, the performance of the office will be dictated by the BITS/SSU clock, since only the BITS/SSU clock is subject to stress on its timing reference.

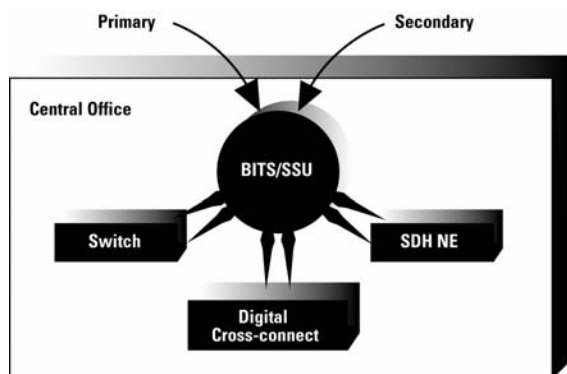


Figure 10. BITS/SSU Configuration

Using the best facilities to transport synchronization reference is required to minimize slips. The best facility may be defined as the reference with the fewest impairments. This refers to a reference that has the least average number of SES and is free from excessive timing instabilities (jitter and wander). References that are payloads on SDH/SONET should not be used for timing, since they are subjected to pointer processing, which adds excessive wander and jitter onto the reference. Similarly, references that are

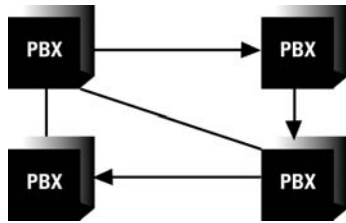


Figure 11. Excessive Cascading

transmitted by ATM Constant Bit Rate services will exhibit large amounts of wander and should not be used for timing.

Cascading of timing references through a network should be minimized (Figure 11). Timing performance will always degrade as timing is passed from clock to clock. The more clocks and facilities in a synchronization chain, the greater the accumulated degradation will become, and the larger the frequency offset grows. Each facility will add impairments to which the clocks in the chain must react. Therefore, for best performance, synchronization chains should be kept short.

Planning Issues

Not all synchronization planning concepts can be simultaneously adhered to. This is especially the case in private networks. Private networks' limited connectivity often results in lack of secondary references and long synchronization chains. In addition, network architecture can make non-hierarchical situations unavoidable. Lack of external timing options in most CPE makes use of a BITS configuration infeasible. In addition, most private networks rely on the poor performance of stratum 4 CPE clocks. With all these factors, designing an adequately-performing private-network synchronization plan can be difficult.

In carrier networks, the introduction of SDH and SONET facilities can impact the amount of cascading in a network. SDH/SONET network elements retime the facility. As facilities become SDH or SONET, chains of SDH/SONET clocks can appear between network offices. In addition, since most SDH/SONET clocks are poorer than stratum 3 in performance, hierarchy issues appear. Therefore, with the introduction of SDH or SONET, the synchronization plan should be reviewed to ensure adequate performance and robustness.

VI. Conclusion

An overview of network synchronization has been presented. It has been shown that synchronization performance has a strong impact on digital data services, encrypted services, and on new technologies, such as SDH and SONET. The major contributor to synchronization performance in real network operation is the frequency offset that a receiver clock exhibits relative to the primary reference source to which it is locked. This performance degradation can be controlled by the introduction of several primary reference sources, by the use of robust clocks, and by proper synchronization planning.

References

- [1] AT&T, "Effects of Synchronization Slips," ITU-T Contribution COM SpD-TD, No. 32, Geneva, November, 1969.
- [2] J. E. Abate, and H. Drucker, "The Effect of Slips on Facsimile Transmission," ICC'88, 1988 IEEE.
- [3] H. Drucker, and A.C. Morton, "The Effect of Slips on Data Modems," ICC'87, CH2424-0/87/0000-0409, 1987 IEEE.
- [4] J. E. Abate, et al, "AT&T's New Approach to the Synchronization of Telecommunication Networks," IEEE Communications Magazine, Vol. 27, No. 4, April 1989.
- [5] K. Inagaki, et al, "International Connection of Plesiochronous Networks Via TDMA Satellite Link," International Conference on Communications, 1982 IEEE, 0536-1486/82/0000-0221.
- [6] M. Decina and Umberto de Julio, "Performance of Integrated Digital Networks: International Standards", International Conference on Communications, 1982 IEEE, 0536-1486/82/0000-0063.
- [7] American National Standard for Telecommunications, "Synchronization Interface Standards for Digital Networks," ANSI T1.101-1994.
- [8] European Telecommunication Standards, "The Control of Jitter and Wander Within Synchronization Networks," Draft ETS DE/TM-3017.
- [9] ITU-T Recommendation G.811, "Timing requirements at the output of primary reference clocks suitable for plesiochronous operation of international digital links."
- [10] ITU-T Recommendation G.824, "The control of jitter and wander within digital networks which are based on the 1544 kbit/s hierarchies."
- [11] ITU-T Recommendation G.821, "Error performance of an international digital connection forming part of an integrated services digital network."
- [12] G. Garner, "Total Phase Accumulation in a Network of VT Islands for Various Levels of Clock Noise," Contribution to ANSI T1X1.3, Number 94-094, September, 1994.
- [13] ITU-T COM XVIII D.1378, "Standard Clock Testing Methodology," 1987.
- [14] "When the Timing is Right, Networks Run Like Clockwork," AT&T DataBriefs, Vol. 2, No. 4, November 1992.
- [15] "AT&T ACCUNET Synchronization Planning Service Makes Phantom Problems Disappear," AT&T DataBriefs, Vol. 2, No. 4, November 1992.



Symmetricom
2300 Orchard Parkway
San Jose, CA 95131, USA
tel: 408-433-0910
fax: 408-428-7897
e-mail: info@symmetricom.com
<http://www.symmetricom.com>

Symmetricom Limited
2 The Billings
Walnut Tree Close
Guildford, Surrey
GU1 4UL, England
tel: 44-1483-510300
fax: 44-1483-510319

For more information:

**Synchronizing Telecommunications Networks:
Synchronizing SDH/SONET—Application Note**

**Synchronizing Telecommunications Networks:
Fundamentals of Synchronization Planning—
Application Note**

APP/STN-BC/D/0200/2M