

INTERNATIONAL
STANDARD

ISO/IEC
11573

First edition
1994-12-15

**Information technology —
Telecommunications and information
exchange between systems —
Synchronization methods and technical
requirements for Private Integrated
Services Networks**

*Technologies de l'information — Télécommunications et échange
d'information entre systèmes — Méthodes de synchronisation et
exigences techniques pour les réseaux privés avec intégration de services*



Reference number
ISO/IEC 11573:1994(E)

Contents

Section 1 : General

1.1 Scope	1
1.2 Definitions	1
1.3 Abbreviations and acronyms	3
1.4 Impact of slips	4

Section 2 : Technical requirements, Synchronization methods

2.1 Technical requirements	5
2.1.1 Jitter and wander at the input	5
2.1.1.1 C0 and T0 interfaces (144 kbits/s)	5
2.1.1.2 C1 and T1 interfaces (1,544 Mbits/s)	5
2.1.1.3 C2 and T2 interfaces (2,048 Mbits/s)	6
2.1.2 Jitter and wander at the output	6
2.1.2.1 C0 and T0 interfaces (144 kbits/s)	6
2.1.2.2 C1 and T1 interfaces (1,544 Mbits/s)	6
2.1.2.3 C2 and T2 interfaces (2,048 Mbits/s)	7
2.1.3 Frequency deviation at the input	7
2.1.4 Accuracy	7
2.1.5 Lock range	7
2.1.6 Phase discontinuity of slave clocks	7
2.2 Synchronization methods for PISNs	8
2.2.1 High level concepts	8
2.2.2 Reference Clock Switching Criteria	8
2.2.3 Reference Restoral	9
2.2.4 Timing Reference Interfaces and Alarms	9
2.2.5 Buffers	9
2.2.6 Controls	9
2.2.7 Slip performance objectives	9
2.2.8 Strategies	10

Section 3 : Description of the synchronization methods

3.1 Plesiochronous operation	11
3.2 Synchronization from one input	11
3.3 Automatic switch over with signalling	12
3.4 Automatic switch over without signalling	12

Annexes

A Choice of clock references

A.1 Choice of reference from public nodes	13
A.2 Choice of references between private nodes	14
A.3 Avoidance of Timing Loops	15

B Synchronization configuration

B.1 Master Slave configurations (synchronization)	16
B.2 master—master configuration (split timing)	17

C Basis of strategies	
C.1 Slip rate	18
C.2 Allocation of the controlled slips	18
C.3 Unavailability of the links	19
C.4 Nodal solutions	20
C.5 Description of the five options	21
D Synchronized Private Network Examples	
D.1 Example with a small private network	22
D.2 Example with a big private network	22
D.3 Example with two different public clock sources	23
D.4 Example with a transit node	23
E Slave Clock Performance Measurement Guidelines	
E.1 Slave Clocks considerations	24
E.1.1 Ideal Operation	24
E.1.2 Stressed Operation	25
E.1.3 Holdover Operation	25
E.2 Test Configuration Guidelines	25
E.2.1 Reference Clock	25
E.2.2 Digital Reference Simulation	26
E.2.3 Output Timing Signal Recovery	26
E.3 Test categories	26
E.3.1 Ideal Testing	26
E.3.2 Stress Testing	26
E.3.3 Holdover Testing	27
F Signalling for management of synchronization	
F.1 Presentation	28
F.1.1 Configuration parameters	28
F.1.2 Reactions of the node	28
F.1.3 Reference clock switching and restoral	28
F.2 Description of the states	28
F.2.1 Initial states	28
F.2.2 Slave states	29
F.2.3 Autonomous state	29
F.2.4 Wait states	29
F.3 Description of the events	29
F.3.1 Failure of links	29
F.3.2 Signalling information	29
F.3.3 Time out	29
F.4 SDL representation of the state machine	29
G Bibliography	34

Foreword

ISO (the International Organization for Standardization) and IEC (the International Electrotechnical Commission) form the specialized system for worldwide standardization. National bodies that are members of ISO or IEC participate in the development of International Standards through technical committees established by the respective organization to deal with particular fields of technical activity. ISO and IEC technical committees collaborate in fields of mutual interest. Other international organizations, governmental and non-governmental, in liaison with ISO and IEC, also take part in the work.

In the field of information technology, ISO and IEC have established a joint technical committee, ISO/IEC JTC 1. Draft International Standards adopted by the joint technical committee are circulated to national bodies for voting. Publication as an International Standard requires approval by at least 75 % of the national bodies casting a vote.

International Standard ISO/IEC 11573 was prepared by Joint Technical Committee ISO/IEC JTC 1, *Information technology*, Subcommittee SC 6, *Telecommunications and information exchange between systems*.

During the preparation of this International Standard, information was gathered on patents upon which application of the standard might depend. Relevant patents were identified as belonging to ALCATEL Business Systems. However, ISO and IEC cannot give authoritative or comprehensive information about evidence, validity or scope of patent and like rights. The patent-holder has stated that licences will be granted under reasonable terms and conditions and communications on this subject should be addressed to

ALCATEL Business Systems
Business Products Division
54, avenue Jean Jaurès
92707 Colombes Cedex
France
Tel. : (1) 47.85.55.55
Telex: 615.531 F
Telefax: (1) 47.85.54.20

Introduction

When synchronous digital signals are being transported over a communications link, the receiving end must operate at the same average frequency as the transmitting end to prevent loss of information. This is referred to as link synchronization. When digital signals traverse a network of digital communications links, switching nodes, multiplexers, and transmission interfaces, the task of keeping all the entities operating at the same average frequency is referred to as network synchronization.

The design of a PISN requires specification of the timing sources and receivers for the synchronization network. Proper design requires that timing loops in the synchronization network be avoided. A timing loop occurs when a clock is using as its reference frequency a signal that is itself traceable to the output of that clock. The formation of such a closed timing loop leads to frequency instability and is not permitted. While it is relatively straightforward to ensure against timing loops in the primary synchronization reference network, care should be taken that timing loops do not occur during failure or error conditions when various timing references are rearranged.

When a PISN is not connected to the public digital network, synchronization can be achieved by having all PISN equipment derive timing from a single source. This source should be the highest quality clock available. Alternatively, if timing is derived from more than one class I clock, or public clock traceable source, the network is said to be operating *plesiochronously*.

If a PISN is connected to the public network at one or more nodes, the private network designer can coordinate with the public network provider to derive class I clock, or public clock traceable timing from the public digital network. More information is available in Annex A.

This page intentionally left blank

IECNORM.COM : Click to view the full PDF of ISO/IEC 11573:1994

Information technology — Telecommunications and information exchange between systems — Synchronization methods and technical requirements for Private Integrated Services Networks

Section 1 : General

1.1 Scope

This International Standard contains requirements necessary for the synchronization of PISNs. Timing within a digital private network needs to be controlled carefully to ensure that the rate of occurrence of slips between PINXs within the PISN, and the public switched networks is sufficiently low not to affect unduly the performance of voice transmissions, or the accuracy or throughput (if errored data require re-transmission) of non-voice services.

Requirements are also based upon the interconnection of digital private telecommunication networks via digital facilities in the public (switched or not) telecommunication networks.

This International Standard is one of a series of technical standards on telecommunications networks. This International Standard with its companion standards fills a recognized need in the telecommunications industry brought about by the increasing use of digital equipment and facilities in private networks. It is useful to anyone engaged in the manufacture of digital customer premises equipment (CPE) for private network applications, and to those purchasing, operating or applying digital CPE to digital facilities for Private Integrated Services Networks (PISN).

This International Standard establishes technical criteria necessary in the design of a synchronization plan for a PISN. Compliance with these requirements would be expected to result in a quality PISN synchronization design.

1.2 Definitions

For the purposes of this International Standard, the following definitions apply:

1.2.1 Accuracy

A measure of the maximum departure from the nominal clock rate over a 24 h period, made anytime in the lifetime of the clock, during a defined period of time, within the declared environmental conditions. Frequency deviation may be constrained to the specific accuracy by clock operation in the free running or hold over modes, as defined below.

1.2.2 Asynchronous signals

Signals having not the same nominal rate.

1.2.3 Clock free running mode

In such a mode, the PINX works with its own clock source which is not locked to an external reference and is not using storage techniques to maintain its accuracy.

1.2.4 Clock hold over mode

An operating condition of a clock in which it is not locked to an external reference clock, but uses storage techniques to maintain during a limited period of time its accuracy with respects to the last known reference clock.

1.2.5 Controlled Slip

It consists of the repetition or deletion of an integer number of octets caused by the elastic buffer mechanism used at the interface of a non-synchronous bit stream (a plesiochronous or asynchronous one). Slips and controlled slips shall be considered synonymous in this International Standard.

1.2.6 Jitter

Short-term non-cumulative variations of the significant instants of a digital signal from their ideal positions in time.

1.2.7 Lock range

Maximum frequency offset from the nominal, to which a given clock is able to synchronize.

1.2.8 Master

The term "master" refers to the clock source providing the timing to the PINX.

1.2.9 Maximum time interval error (MTIE)

The maximum time interval error (TIE) for all possible measurement intervals within the measurement period. Figure 1 illustrates the definition of MTIE.

1.2.10 Phase Locked Loop (PLL)

A feedback-controlled system that locks a local clock to an incoming reference clock in both frequency and phase.

1.2.11 Plesiochronous

The essential characteristic of time-scales or signals such as their corresponding significant instants occur at nominally the same rate, any variation in rate being constrained within specified limits.

1.2.12 Primary Reference Clock

Equipment that provides a timing signal, with a long term accuracy equal or better than $\pm 10^{-11}$.

1.2.13 Pull in range

Maximum frequency offset from its own clock, to which a given clock is able to synchronize.

1.2.14 Reference Clock

Timing signal used for synchronization, without any assumption on its accuracy.

1.2.15 Slave

The term "slave" refers to the PINX receiving timing from another source.

1.2.16 Slip

Refer to controlled slip

1.2.17 Split Timing

An arrangement where equipment employs separate transmit and receive clocks on a transmission link having no particular relationship to one another.

1.2.18 Synchronous

Qualifies signals with corresponding significant instants occurring at the same mean rate; the time difference between these homologous instants is generally limited.

1.2.19 Synchronization

The process of adjusting the corresponding significant instants of signals so that a constant phase relationship exists between them.

1.2.20 Time-Interval Error (TIE)

The variation in time delay of a given timing signal with respect to an ideal timing signal over a particular time period. Figure 1 illustrates the definition of TIE.

1.2.21 Timing loop

An unstable condition in which two or more equipment clocks transfer timing to each other, forming a loop without a designated master timing source.

1.2.22 Time to repair

The time by which, with a stated probability, the link is repaired.

1.2.23 Transparent

A link or group of links is transparent if the signal carried is not re-timed from a clock associated with the link(s). The timing of a signal passing across a transparent link may however be altered due to jitter, wander, filtering, or fault conditions. Figure 2 illustrates the definition of transparent and non transparent links.

1.2.24 Wander

The long-term variations of the significant instants of a digital signal from their ideal positions in time. Long-term implies that these variations are of low frequency.

1.3 Abbreviations and acronyms

AIS :	Alarm Indication signal
BITS :	Building Integrated Timing Supply
CCITT :	International Telegraph and Telephone Consultative Committee
CPE :	Customer Premises Equipment
C0 :	Basic Rate transparent or non transparent links
C1 :	1,544 Mbits/s transparent or non transparent links
C2 :	2,048 Mbits/s transparent or non transparent links
C3 :	Non ISDN transparent or non transparent links
DCS :	Digital Cross-connect System
DSX :	Digital Signal Cross-connect
DTE :	Data Terminal Equipment
ESF :	Extended Super Frame
FM :	Frequency Modulation
GPS :	Global Positioning System
MTIE :	Maximum Time Interval Error (see figure1)
MUX :	Multiplexer
NCTE :	Network Channel Terminating Equipment
NI :	Network Interface
PISN :	Private Integrated Services Network.
PINX :	Private Integrated Services Network Exchange (PABX, Key System, ...).
ppm :	parts per million
PSTN :	Public Switched Telecommunication Network
PLL :	Phase Locked Loop
PRC :	Primary Reference Clock
PM :	Phase Modulation
SES :	Severely Errored Second
SDH :	Synchronous Digital Hierarchy
T0 :	Basic Rate Access to public ISDN
T1 :	1,544 Mbits/s Access to public ISDN
T2 :	2,048 Mbits/s Access to public ISDN
TIE :	Time Interval Error (see figure1)
UI :	Unit Interval (488 ns for T2 and C2, 648 ns for T1 and C1, 5208 ns for T0 and C0)
UTC :	Universal Coordinated Time

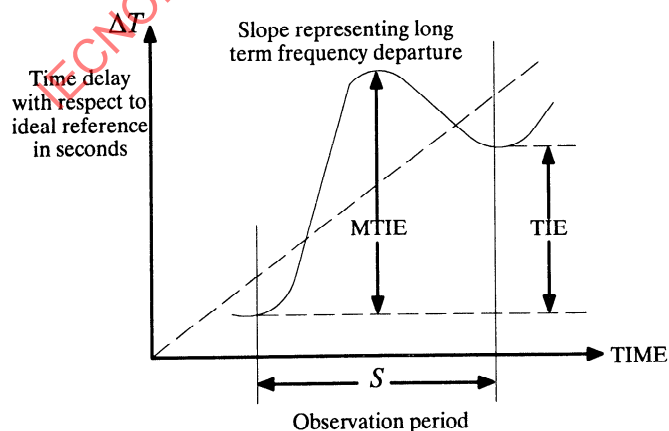


Figure 1 – Time Interval Error (TIE)

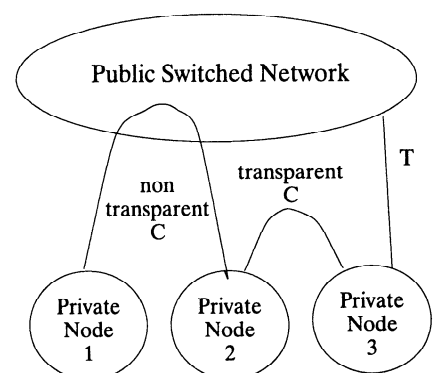


Figure 2 – Links definition

1.4 Impact of slips

When synchronous digital signals are being transported over a link, the receiving end must operate at the same average frequency as the transmitting end to prevent loss of information.

When digital signals traverse a network of digital links, switching nodes, multiplexers and transmission interfaces, the task of keeping all the entities operating at the same average frequency is referred to as synchronization. If the distant transmitter sends at a bit rate higher than the switching system clock, the receive buffer in the switching system eventually will overflow, causing one frame to be lost. If the received bit stream is at a lower bit rate, the buffer will underflow, causing a frame to be repeated. Either occurrence is called a controlled slip.

In a digital PISN, slips can be prevented by forcing all equipment to use a common reference clock. In a digital PISN, when it is not connected to the public digital network, network synchronization is achieved by having all equipments derive timing from a single source that should be the highest quality clock available.

If a PISN is connected to the public network at one or more nodes, the private network shall derive timing from the public digital network timing reference to ensure that the highest quality timing source is used.

The design of a PISN requires specification of the timing sources and receivers to achieve synchronization. Proper design requires that timing loops in the synchronization plan be eliminated.

The impact of slips on service carried on digital networks depends on the application and type of service being provided. Some examples of the effects of which are summarized in the following table.

Table 1 – Impacts of a slip

Service	Potential impact
Encrypted text	Encryption key must be retransmitted
Video	Freeze frame for several seconds Noise burst ("pop") on audio
Digital data	Deletion or repetition of data Possible misframe Reduction of throughput
Facsimile	Connection establishment may be not successful Deletion of 4–8 scan lines or lost of throughput, depending on facsimile system
Voice Band Data	Transmission errors for 0,01 to 2 s Drop call (for some modems)
Voice	Possible "Click"

In addition to slips, synchronization-related impairments caused by transmission effects on equipments such as error bursts and phase discontinuity, can also have an impact on customer service. These degradations can propagate and multiply through the network.

All of the degradations described above can be controlled by appropriate synchronization strategies and clock designs, as described in later clauses.

Section 2 : Technical requirements, Synchronization methods

2.1 Technical requirements

The phase stability of a slave clock can be described by:

- its long–term phase variations (wander and integrated frequency departure);
- its short–term phase variations (jitter);
- phase discontinuities due to transient disturbances.

NOTES

- 1 – The values given in 2.1.1, 2.1.2 and 2.1.3 are taken from ETS 300 012 [3] for C0 and T0, from ETS 300 011 [2] for T2 and C2 and from EIA/TIA–594 [8] for C1 and T1.
- 2 – It has been found necessary to limit the wander value for the T0 and C0 interfaces. The need for the additional parameters for accuracy and lock range derives from the strategies in 2.2.8.
- 3 – Requirements for phase discontinuity are taken from CCITT Recommendation G.812.

2.1.1 Jitter and wander at the input

2.1.1.1 C0 and T0 interfaces (144 kbits/s)

The C0 and T0 inputs shall tolerate at least a sinusoidal input jitter within the mask shown in figure 3 without producing bit errors:

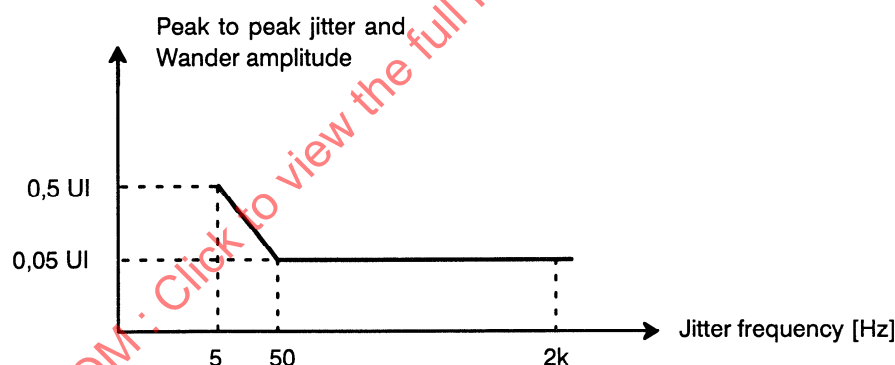


Figure 3 – Tolerable jitter and wander at PINX input for Basic Access

In order to save power, when both B channels are idle, carriers may disable T0 interfaces. Under these conditions, timing information is not available. Synchronization shall be derived from interfaces which are continuously available.

NOTE : The maximum relative wander between two or more interfaces is limited to 4 UI (except for plesiochronous operation).

2.1.1.2 C1 and T1 interfaces (1,544 Mbits/s)

The equipment shall operate with jitter of the received signal which does not exceed the following limits, in both bands simultaneously :

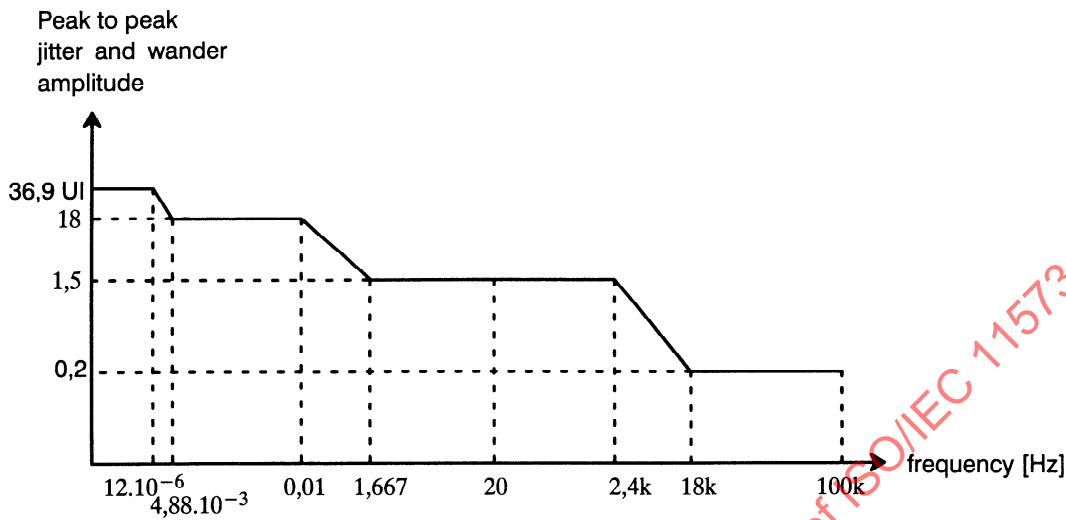
- (1) Band 1 [10 Hz – 40 kHz] : 5,0 UI, peak–to–peak, and
- (2) Band 2 [8 kHz – 40 kHz] : 0,1 UI, peak–to–peak.

For T1 and non transparent C1 interfaces, the equipment shall operate with wander of the received signal of up to 28 UI (18 μ s) peak–to–peak over any 24h period and up to 23 UI (15 μ s) peak–to–peak in any 1h interval.

Wander requirements for transparent C1 interfaces are for further study.

2.1.1.3 C2 and T2 interfaces (2,048 Mbits/s)

The input shall tolerate a sinusoidal input jitter / wander within the mask shown in figure 1 without producing bit errors or losing frame alignment:



NOTE : The value of 36,9 Unit Interval (UI) is the maximum relative wander between two or more interfaces, except for plesiochronous operation.

Figure 4 – Tolerable jitter and wander at PINX input for 2Mbits/s access

2.1.2 Jitter and wander at the output

Wander accumulation within a private network needs to be controlled. Output jitter requirements of this subclause apply when the input jitter meets the requirements of 2.1.1.

2.1.2.1 C0 and T0 interfaces (144 kbits/s)

For further study.

2.1.2.2 C1 and T1 interfaces (1,544 Mbits/s)

Transmit Signal Jitter for T1 and non transparent C1 interfaces

The jitter of the transmitted signal at the equipment output interface shall not exceed the following limits, in both bands simultaneously :

- (1) Band 1 0,5 UI peak-to-peak, and
- (2) Band 2 0,07 UI peak-to-peak.

Transmit Signal Wander for T1 and non transparent C1 interfaces

The wander in the transmitted signal of the equipment shall not exceed the wander of its received signal by more than 2,5 UI. The wander of the transmitted signal shall not exceed 28 UI (18 μ s) peak-to-peak in any 24h period, nor exceed 23 UI (15 μ s) peak-to-peak in any 1h interval under operating conditions defined as having class I clock, or public clock traceability over facilities with typical short-term impairments that do not include events that result in phase transients.

It is recognized that currently, the wander in the transmitted signal may be as large as 7700 UI (5 ms) peak-to-peak in any 24h period and may be as large as 4600 UI (3 ms) peak-to-peak in any 1h interval under normal operating conditions. However, it must be recognized that such wander will result in frame slips within the network.

NOTE: Transparent C1 jitter and wander requirements are for further study.

2.1.2.3 C2 and T2 interfaces (2,048 Mbits/s)

T2, one interface

band 1	$f \in [20\text{Hz} - 100\text{kHz}]$	$< 1,1 \text{ UI}$
band 2	$f \in [400\text{Hz} - 100\text{kHz}]$	$< 0,11 \text{ UI}$

T2, multiple interfaces and non transparent C2 interfaces

band 1	$f \in [4\text{Hz} - 100\text{kHz}]$	$< 1,1 \text{ UI}$
band 2	$f \in [40\text{Hz} - 100\text{kHz}]$	$< 0,11 \text{ UI}$

transparent C2 interfaces

band 1	$f \in [4\text{Hz} - 100\text{kHz}]$	$< 1,6 \text{ UI}$
band 2	$f \in [40\text{Hz} - 100\text{kHz}]$	$< 0,1 \text{ UI}$

2.1.3 Frequency deviation at the input

The interfaces shall tolerate input clock rates within the following ranges around the nominal value:

T0	$\pm 100 \text{ ppm}$
T1	$\pm 32 \text{ ppm}$
T2	$\pm 50 \text{ ppm}$

These values are only relevant for the interfaces, during maintenance and failure conditions, not for the design of the clock unit.

2.1.4 Accuracy

Accuracy is defined in 1.3

Since class I clocks are intended to be used as master clocks in plesiochronous private networks, they only operate in free running mode.

In order to conform with the strategies defined later (2.2.8), clocks shall comply with the following classes :

class I	$\leq \pm 7,10^{-10}$
class II	$\leq \pm 1,10^{-6}$
class III	$\leq \pm 50,10^{-6}$

NOTES:

- 1 In certain network configurations, clocks with higher accuracy are necessary (see Annex D).
- 2 Class III free running clocks are typically not used as timing sources.

2.1.5 Lock range

Slave clocks within accuracy of class II and III shall comply with one of the following lock range classes :

class a	$> \pm 1 \text{ ppm}$
class b	$> \pm 50 \text{ ppm}$

The following combinations of slave clocks are allowed : IIa, IIb, IIIb.

2.1.6 Phase discontinuity of slave clocks

Phase transients are changes in phase relationships. Transients are specified in terms of the maximum transient phase deviation and the maximum equivalent frequency offset during the transient. The MTIE and phase-slope requirements shall also be met under all timing reference degradations, independent of whether a switch of reference has occurred.

(1) In case of internal testing or rearrangement operations within the slave clock, the following conditions shall be met:

- the phase variation over any period of up to 2^{11} UI must not exceed $1/8$ of a UI;
- for periods greater than 2^{11} UI , the phase variation for each interval of 2^{11} UI must not exceed $1/8 \text{ UI}$ up to a total amount of $1 \mu\text{s}$,

Where UI corresponds to the reciprocal of the bit rate of the interface.

(2) Equipment shall operate with transients in the phase of received signals of up to 1 μ s with a maximum rate of change equivalent to 61 ppm frequency offset. Such transients shall be isolated in time.

Additionally, accommodation needs to be made for SDH (SONET) virtual tributary (VT) pointer adjustments with a magnitude of 8 UI (4,7 μ s for VC11 and 3,57 μ s for VC12). Phase slope characteristics of this transient have yet to be defined but typically fall within a 1 second time frame and to be no greater than the equivalent of a frequency offset of 61 ppm. SDH (SONET) quantizes input wander into 8 UI steps. When the upper or lower threshold in a SDH (SONET) pointer processor is reached, the SDH (SONET) pointer processor will generate an outgoing pointer adjustment to either the next downstream pointer processor or far-end desynchronizer.

(3) When equipment receives a phase transient conforming (1), its output shall not exceed (1)

2.2 Synchronization methods for PISNs

Operations include both the provisioning and maintenance of the digital synchronization network. Provisioning means engineering an appropriate configuration for the network and installing any particular equipment necessary to implement the configuration. Maintenance activities involve the detection of synchronization network failures and restoration of timing distribution.

2.2.1 High level concepts

(1) The public network is always to be taken as the reference clock source when available and in operational mode, except where a PINX contains a clock with characteristics in accordance with class I, which needs not synchronize its clock generators to any input.

(2) Synchronization plans need to be hierarchical; timing is passed from better performing sources to clocks of lower or equal performance.

(3) Timing loops need to be eliminated.

(4) Timing sources need to be diverse whenever possible.

(5) The lock range of any node is to be sufficient to cover the accuracy of any potential master.

(6) Cascading of timing references through CPE needs to be minimized.

(7) Non-transparent C type interfaces are treated like T type interfaces from the synchronization point of view. The strategies described here concern only the cases where the private links are transparent.

2.2.2 Reference Clock Switching Criteria

To minimize excessive switching of the timing reference and the accumulation of phase movements, a clock shall not initiate a switch of reference until the timing reference has become degraded. Reference switch over shall occur at or after, but not before, any of the following reference line degradations or events:

T0, C0	loss of incoming signal loss of frame alignment (stable state) link not activated
T1, C1	loss of signal for 50 ms error bursts of a duration of 2,5 s or more at bit error ratios worse than 10^{-3} 17 misframes in a 24h period in response to an external control signal received AIS.
T2, C2	loss of incoming signal loss of frame alignment (stable state) received AIS
T3, C3	loss of incoming signal received AIS

2.2.3 Reference Restoral

Once traceability to the class I has been restored to specification, an automatic switchback may be initiated only if such a switchback does not produce impairments. That is, all digital CPE, regardless of type, shall meet the requirements of 2.1.6 if the clock performs an automatic switchback. To prevent chatter because of repeated automatic switching between two references, a time delay of 10 s or more between reference switches is desirable.

A manual operation is permitted at any time to switch back to the primary reference. Additional references may be added or used to replace existing references by manual procedure (e.g., physically reconnecting new references to the external reference input connectors or changing the source of traffic—carrying reference lines by software reconfiguration).

2.2.4 Timing Reference Interfaces and Alarms

The digital facility or timing reference interface receiver shall provide timing, slip and misframe information necessary to maintain the synchronization system. It is desirable to include information on whether the slip caused a repetition or deletion of a frame. The lack of slip information will increase the probability of an undetected reference degradation occurring.

2.2.5 Buffers

The input signal shall be buffered in the CPE. The buffer shall provide at least 125 μ s (one frame) of storage to allow for controlled slips. In addition, the buffer shall provide a minimum of 18 μ s of hysteresis to absorb jitter, wander, and frame timing differences between incoming signals and to decrease the probability of incurring back-to-back slips.

2.2.6 Controls

The digital CPE synchronization system shall provide the capability of manually overriding automatic reference switch over. This action may be required as part of a diagnostic procedure. A manual capability to disable the automatic switch over of references because of excessive slip rate shall be provided. This is required when only one digital link provides slip information and a potentially ambiguous situation is entered during fault conditions.

2.2.7 Slip performance objectives

The designer of a PISN should consider the following slip rate performance objectives from Annex C.

Table 2 – Slip performance objectives

Performance Category	Mean slip rate	Proportion of time (see NOTE)
a)	≤ 5 slips in any 24 h period	$> 99,56 \%$
b)	> 5 slips in any 24 h period and ≤ 30 slips in any 1 h period	$< 0,4 \%$
c)	> 30 slips in 1 h period	$< 0,04 \%$

NOTE : Total time greater than 1 year .

To comply with these objectives, the PISN designer shall select for each PINX one of the strategies listed in 2.2.8.

2.2.8 Strategies

Three different types of strategies have been identified :

- Strategy 1 : the clock signal of the private node has class I accuracy. In this case, the private network can operate in a plesiochronous mode with the public network (see clause 3.1).
- Strategy 2 the accuracy of the clock is class II. The private node is synchronized, but only one input is required. It can work asynchronously when the master clock fails (see clause 3.2). In some network configurations, two inputs are required in order to meet the slip performance objectives of table 2.
- Strategy 3 : the accuracy of the private node's clock is class III. If more than one link to a class I clock source is available, and the private node is capable of being synchronized by at least two inputs, with switch over mechanisms between these inputs. the switch over can be
- Strategy 3.1 – with exchange of information (see clause 3.3, and Annex A, B and C)
- Strategy 3.2 – without exchange of information (for some network configurations only, see Annex D)

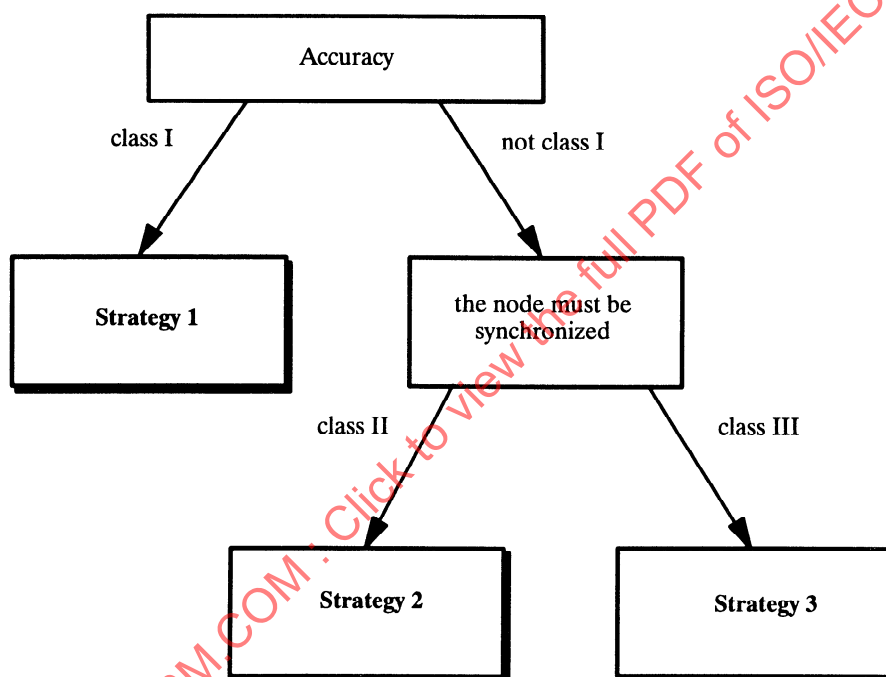


Figure 5 – Synchronization strategies classification

Section 3 : Description of the synchronization methods

3.1 Plesiochronous operation

Strategy 1 in 2.2.8 applies to a node having a clock in accuracy class I. Such a node is not synchronized by any external signal: the node has the ability to work plesiochronously and is in the clock free running mode. The relevant parameters for such a node are:

- jitter and wander at the output (see 2.1.2);
- accuracy in the free running mode (class I in 2.1.4);
- phase stability as described in 2.1.6 in case of internal rearrangements.

A node working with strategy 1 can be a master for any node (with strategies 2 and 3).

3.2 Synchronization from one input

Strategy 2 in 2.2.8 in part II applies to a node having a clock in accuracy class II and only one input for synchronization purposes. The relevant parameters for such a node are:

- jitter and wander at the input (see 2.1.1);
- jitter and wander at the output (see 2.1.2);
- accuracy: the node has the average frequency of the class I source when synchronized, and performs in class II during the failure of the synchronizing input, and during the recovery time;
- during the switch over from the external clock to its own clock (in the free running or in the hold over mode), and during the reverse operation, the phase stability shall conform to 2.1.6. This is not depicted in figure 5. During this time, the node shall stay in performance category (b) for the slip rate (see Annex C).

A node working with strategy 2 can be a master for a node with strategies 2 and 3. It can be enslaved by a public node, or by a private node with strategy 1, or with strategy 2 or with strategy 3.1. Timing loops have to be avoided

The behavior of the node during the failure of the synchronizing link can be described as follows:

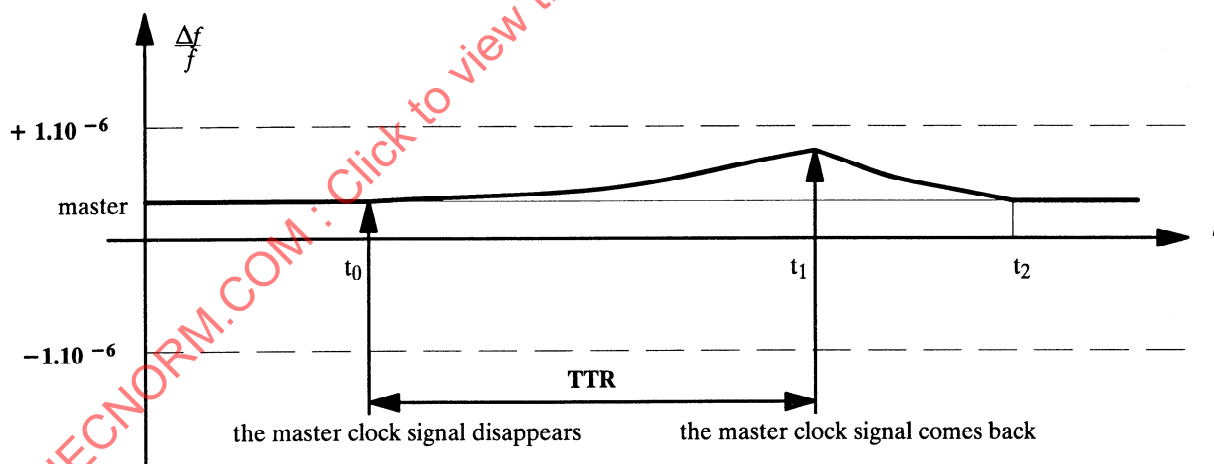


Figure 6 – Example of failure and restoration in case of strategy 2

Until t_0 : the node is synchronized by a digital bit stream of class I.

At t_0 : the bit stream or the master clock signal is no longer valid (see 2.2.2).

From t_0 to t_1 : the node has no master and runs in the free running or in the hold over mode. The duration of this period is called TTR and may be in the range of days.

At t_1 : the master clock signal has been restored and the node begins the re-synchronization phase.

After t_2 : the node is synchronized by a master clock.

NOTES

- 1 No assumption has been made on the way the master clock is determined to have been restored.
- 2 The master clock before t_0 and the master clock after t_1 are not necessarily the same clocks.

3.3 Automatic switch over with signalling

Strategy 3.1 in 2.2.8 corresponds to a node with a clock in accuracy class III, which is capable of taking synchronization from either of two entries. Each node switches automatically from one master to the other, based on the exchange of information between nodes in the private network.

The exchange of information between nodes is used to avoid loops and also to inform the other nodes, even when a failure is hidden by another node.

A number of methods to achieve this are known and may continue to exist. For interoperability, the goal is to have one standardized method. Annex F provides currently available information on a method of synchronizing a PISN which utilizes automatic switch over with signalling.

The relevant parameters for such a node are:

- jitter and wander at the input (see 2.1.1);
 - jitter and wander at the output (see 2.1.2);
 - accuracy: the node has the average frequency of the class I source when synchronized, and performs in class II during rearrangement, and returns to the accuracy of class I when synchronized to the secondary.
 - during the period of rearrangement from primary to secondary, the phase stability shall conform to subclause 2.1.6 .
- This is depicted in figure 7.

The behavior of the node during the failure of the synchronizing link is described as follows:

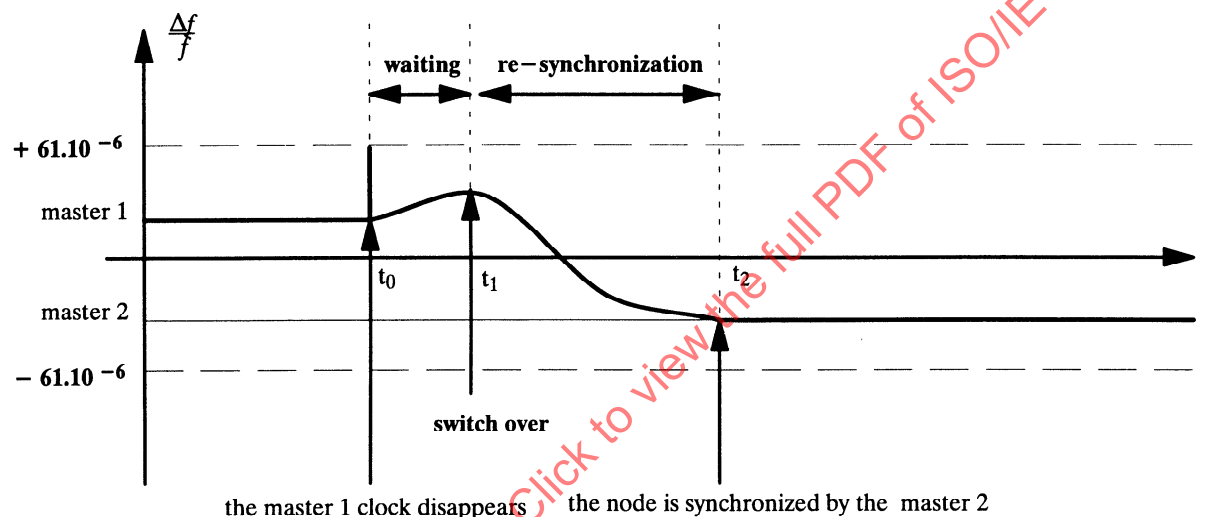


Figure 7 – Example of failure and restoration in case of strategy 3.1

Until t_0 :the node is synchronized by its primary master clock signal.

At t_0 :the primary master clock is no longer valid (see subclause 2.2.2) and the node enters into the clock free running or hold over mode.

From t_0 to t_1 :the node exchanges information with its neighbors in order to take the best decision. This period of time, called "waiting" in the figure 7 has a duration in the order of seconds.

At t_1 :the node has decided to switch to the secondary master clock signal.

From t_1 to t_2 :the node synchronizes its clock to the second master clock signal.

From t_0 to t_2 :the node must stay in the performance category (b) and must conform to 2.1.6.

3.4 Automatic switch over without signalling

Strategy 3.2 corresponds to a node with a clock in accuracy class III, which is synchronized by at least two reference clocks.

The relevant parameters for such a node are:

- jitter and wander at the input (see 2.1.1);
- jitter and wander at the output (see 2.1.2);
- accuracy: the node has the average frequency of the class I source when synchronized, and performs in class II during rearrangement, and returns to the accuracy of class I when synchronized to the secondary.

For some specific configurations, failures may be hidden by a node, and timing loops could occur; consequently in those cases, this strategy shall not be used.

Annex A (informative)

Choice of clock references

A.1 Choice of reference from public nodes

There are fundamentally two architectures that may be used to pass timing across the interface between the public network and private digital networks. The first is for the private digital network to accept a primary reference clock at one location and to then provide timing references to all other private digital network locations over interconnecting facilities (see figure A.1). The second is for the private digital network to accept a reference clock at each interface (see figure A.2).

In method one (figure A.1), the private network owner has control of the synchronization of his network, i.e., location 1 is providing clock references to locations 2, 3, 4 and 5. From an administrative viewpoint, this looks good since there is only one set of references between the public and the private digital network.

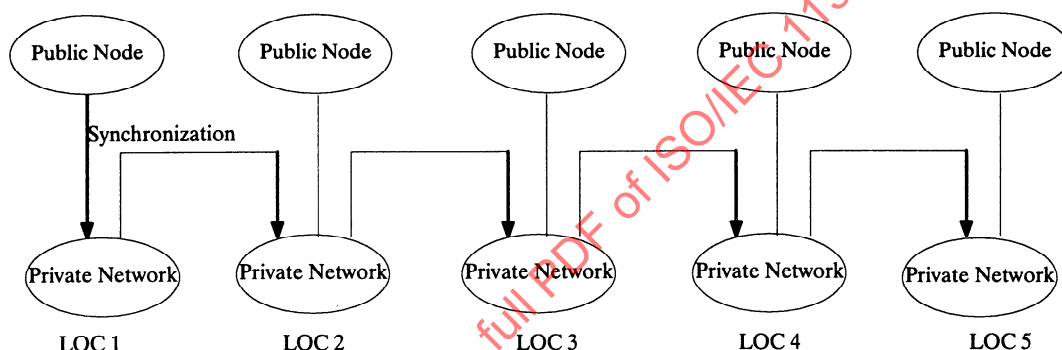


Figure A.1 – Private Digital Network Synchronization Reference at One Interface with a carrier

However, from a technical and maintenance viewpoint, there are limitations. For example, a loss of the references at Location 1 would cause all private digital network locations to slip against the public network(s). This creates a trouble that is at best difficult to detect. Trouble reports may be generated at any office and maintenance forces may spend time looking for a trouble that may be in another office, state or company. Another weakness is the existence of multiple clocks and timing links in the private digital network timing chain. This daisy chain type of timing has two inherent weaknesses. One, the probability for failure is increased because of the additional links and clocks. Second, each clock and facility has the capability of adding jitter and wander, e.g., a signal transmitted from location 5 may not meet the interface requirements for jitter or wander. Synchronous facilities of the future may cause other difficulties.

In the second method (figure A.2), primary reference clocks are provided to the private digital network at each interface with the public network. In this arrangement, the loss of a primary reference would cause a minimum of troubles. For example, a loss of reference to Location 1 would cause slips only between Office A and Location 1, and between Locations 1 and 2.

Additionally, the slips against the public network would occur at the same interface as the source of the trouble, making trouble location and subsequent repairs easier. The timing path(s) are shorter than method 1, thereby increasing reliability and reducing the level of jitter or wander.

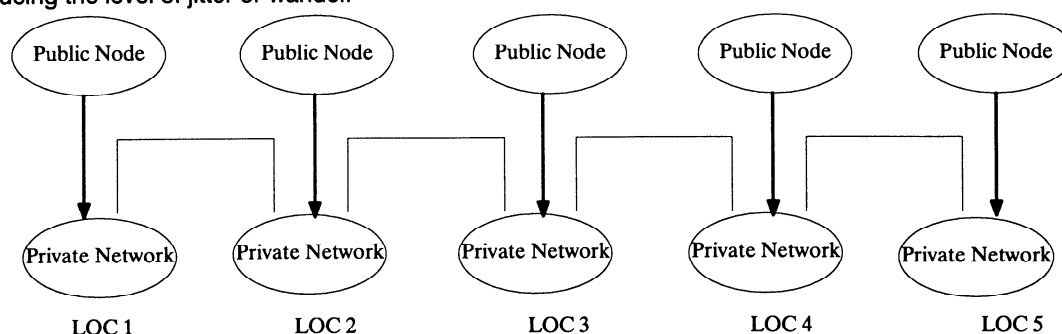


Figure A.2 – Private Digital Network Synchronization Reference at Each Interface with a carrier

A.2 Choice of references between private nodes

The clock references correspond to the location where clock signals are to be received.

Two particular entries shall be chosen among all the inputs : a primary main input called p, and a secondary called s. It is possible to use more than two entries to synchronize a node of the network, but to comply with 2.2.7, a mechanism with only two entries is sufficient.

The method for choosing timing inputs is based on giving a coefficient to any potential clock input to a node ; two and one coefficient to each node. The coefficient of a node corresponds to the optimum path to the public or class I clock. Thus it represents the lowest coefficient of its incoming links.

Each termination of a link (input) is given a coefficient which represents the minimum number of links between that termination and the public network or the master node.

The coefficient of the link corresponds to the coefficient of the node plus 1 unit ; however, the value of the coefficient's increment may be weighted (i.e. one unit or more) depending upon reliability or quality criteria of the links or nodes.

Choice of p (primary source)

The choice of p is then made by taking the input to the node with the lowest coefficient.

If two or more entries have the lowest coefficient any of the inputs with the lowest coefficient may be used.

Choice of s (secondary source)

The choice of the secondary input is then made by taking the input with the lowest coefficient except the input chosen for p.

To avoid timing loops, where the nodes do not utilize switch over mechanisms with signalling, secondary (ies) must be chosen carefully.

The following figure gives an example of the principle. To not complicate the figure, the case where several links are provided between two nodes is not represented.

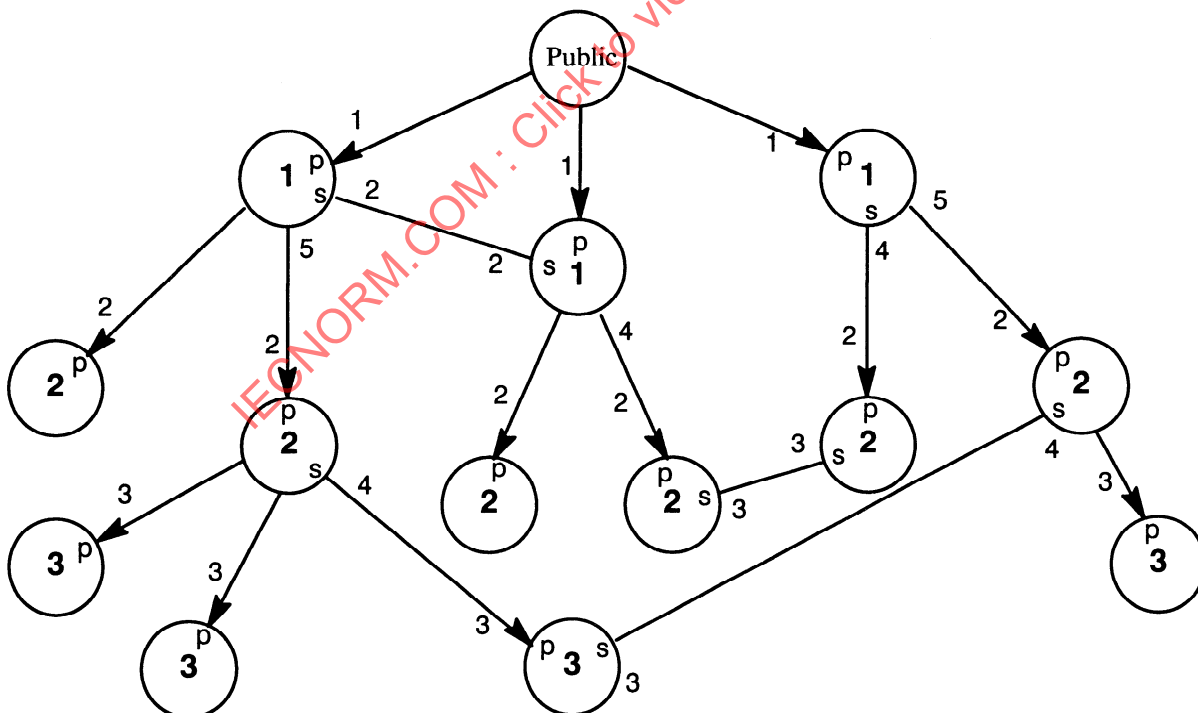


Figure A.3 – Choice of the clock references

A.3 Avoidance of Timing Loops

Improper use of secondary timing references in the synchronization network can possibly create timing loops in the network. That is, a timed clock receives timing from itself via a chain of timed clocks. Timing loops are to be avoided in digital networks. When a timing loop is formed, equipment clocks involved in the timing loop become unstable, and network performance can degrade beyond that which is obtained when all clocks are operating in the free run mode. The potential for loops exists when either primary or secondary reference signals are passed between clocks of the same class and certain failure conditions exist. Figure A.4 (a) is a typical example of a possible timing loop in a private digital network. If references P1 and P2 fail, a timing loop would be formed when clocks 1 and 2 switch to their secondary references. A more appropriate design is shown in Figure A.4 (b). The alternative being the use of signalling exchanges.

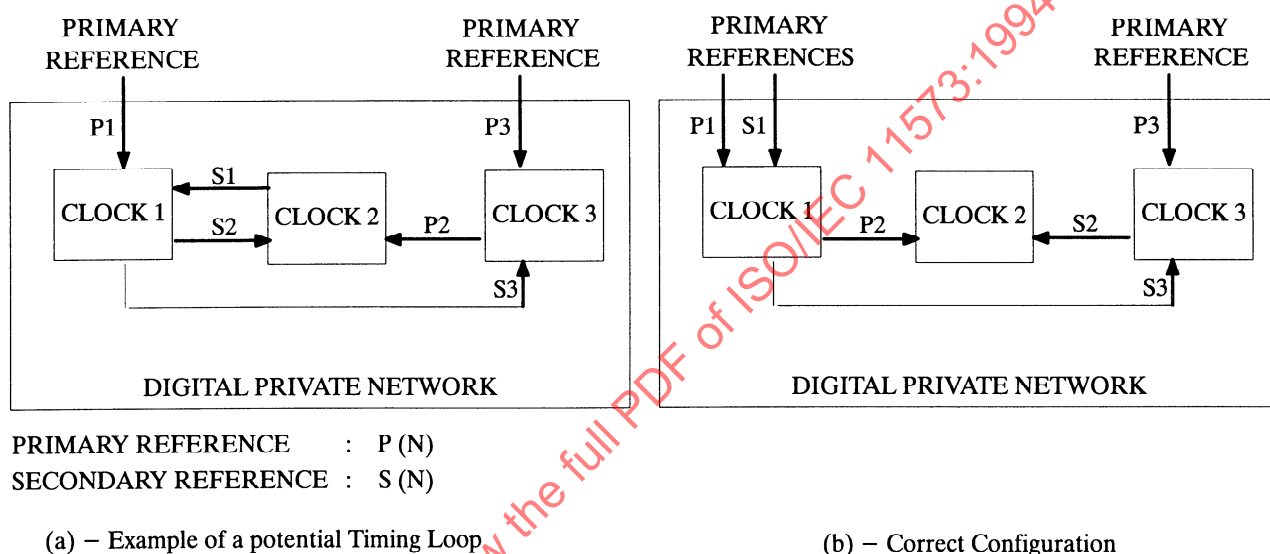


Figure A.4 – Avoidance of Timing Loops

Annex B (informative)

Synchronization configurations

B.1 Master Slave configurations (synchronization)

In the master–slave configurations, the private nodes work synchronously. Figure B.1 represents the simplest master–slave configuration:

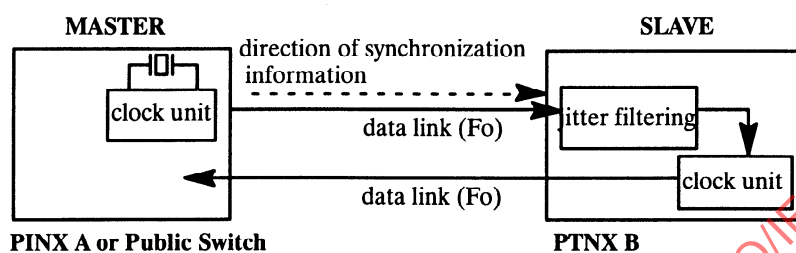


Figure B.1 – Master Slave Configuration

When 3 or more nodes are connected together, the master–slave configuration can be applied in a cascade configuration :

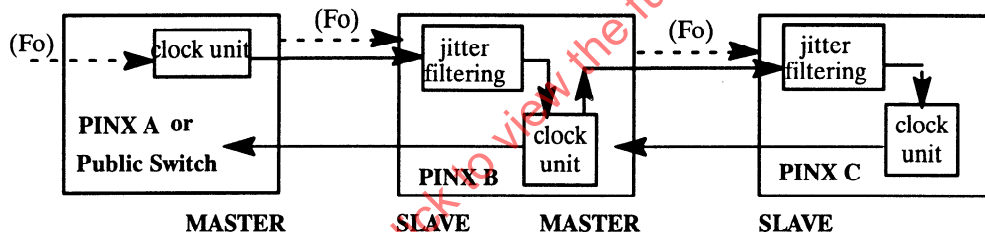


Figure B.2 – Cascade Master – Slave Configuration

When several PINXs are connected to another one, the configuration may be a multiple master configuration :

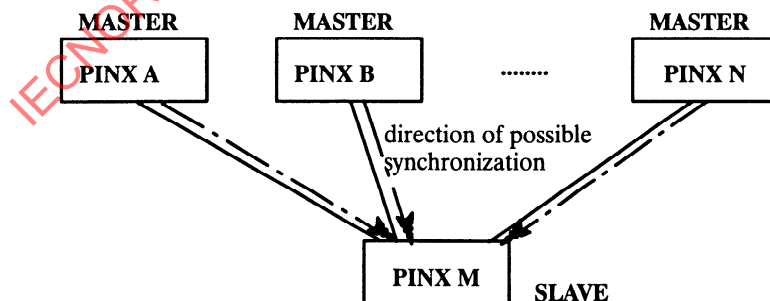


Figure B.3 – Multiple Master – Slave Configuration

B.2 master–master configuration (split timing)

The inputs (x) and (y) are reference clocks. When (x) and (y) are providing a clock signal, two cases can occur:

- (1) x and y come from the same clock, and are synchronous. In such a case, PINX A and PINX B are synchronous.
- (2) x and y are plesiochronous clock sources (from two different networks for instance). In that case, the two PINXs work plesiochronously.

Without x and y (in case of failure for instance), the two PINXs can work

- (1) plesiochronously if they both have clocks in accuracy class I
- (1) asynchronously otherwise

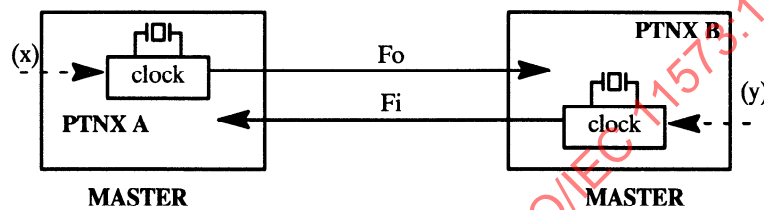


Figure B.4 – Master – Master Configuration

Annex C

(informative)

Basis of strategies

C.1 Slip rate

CCITT Recommendation G.822 [1] specifies the objectives of slip rates for 64 kbits/s international digital connections. The slip rate is reproduced in the following table:

Table C.1 – G.822 slip rate objectives

CCITT Performance Category	Mean slip rate	Proportion of time (NOTE)
a)	≤ 5 slips in any 24 h period	$> 98,9 \%$
b)	> 5 slips in any 24 h period and ≤ 30 slips in any 1 h period	$< 1,0 \%$
c)	> 30 slips in 1 h period	$< 0,1 \%$

NOTE – Total time greater than 1 year.

C.2 Allocation of the controlled slips

CCITT Recommendation G.822 [1] proposes the following allocation for the various portions of the Hypothetical Reference Connection:

Table C.2 – Allocation of the controlled slips

Section	Allocation of the objectives	Part of the total time (NOTE)	
		(b)	(c)
international transit	8,0 %	0,08 %	0,008 %
national transit	6,0 %	0,06 %	0,006 %
local section	40,0 %	0,4 %	0,04 %

NOTE – (b) and (c) refer to performance categories from table C.1

No allowance has been made by CCITT for private networks. For the slip performances of a PISN, an additional allocation of 40% is used as an objective for each private network.

Table C.3 – Allocation of the controlled slips within a private network

Performance Category	Mean slip rate	Proportion of time (see NOTE)
a)	≤ 5 slips in any 24 h period	$> 99,56 \%$
b)	> 5 slips in any 24 h period and ≤ 30 slips in any 1 h period	$< 0,4 \%$
c)	> 30 slips in 1 h period	$< 0,04 \%$

NOTE : Total time greater than 1 year .

C.3 Unavailability of the links

Surveys of link performance have given the following information regarding the unavailability of the links:

- 1) the failure rate of a link (λ) is in the range of $1,10^{-4}$ to $5,10^{-4}$
- 2) the time to repair (τ) is about 24 h.

From these values, $\lambda\tau = 0,004$ is chosen as the basis for the analysis of PISN types in the following subclauses. Other sources of slips (e.g. phase transient, node failures) are PISN dependant and have not been included in this analysis.

C.3.1 Unavailability of the public clock source after (r-1) nodes

– let us call r the number of links for the shortest path from a node to the public ISDN; it means that $(r-1)$ nodes are on the path to the public for that shortest path.

– the rate of unavailability of the public clock source approximates to: $U_r = r \times 0,004$

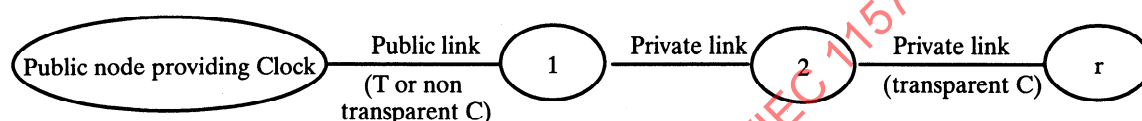


Figure C.1 – Unavailability – Serial configuration

C.3.2 Unavailability of the clock with n links, seen from one node

let us call : n the number of links used for synchronization purpose in one node;

λ_i the failure rate of the link i ;

τ_i the time to repair the link i ;

Assuming strictly independent links, the proportion of time during which the n links are broken is given by :

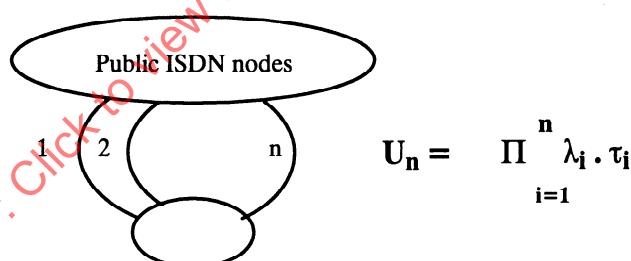


Figure C.2 – Unavailability – Parallel configuration

– with $\lambda_i \cdot \tau_i = 0,5\%$, the proportion of time during which several links of a node are broken at the same time is given in the table

Table C.4 – Unavailability of the reference clock – Parallel configuration

Number of links broken at the same time	1	2	3
Proportion of time (with $p_i = 1$)	0,4 %	0,0016 %	0,00001 %

C.3.3 Conclusions

- the proportion of time for the unavailability of one link providing the clock source from the public is in the range of performance category b) (table C3)
- the proportion of time for the unavailability of 2 independant links is much lower than performance category c) (table C3)

C.4 Nodal solutions

According to the previous calculations, we have the following options for the strategy of synchronization:

Table C.5 – Options for the strategies of synchronization

	Before failure	Primary Failure (F_P)	Secondary Failure (F_S)	Options
performance categories	a)	a)	a) b) c)	Option 1 (2) Option 4 (2) Option 3
		b)	b) c)	Option 2 (3) Option 5 (3)
		c)	c)	not compliant
Calculated Proportion of time	$N = (100 - F_P - F_S)\%$	$F_P = q \times 0,4\%$ (1)	$F_S = p_i \times p_j \times 0,004\%$	

NOTES

- 1 $q = \text{Max}(p_i)$; p_i and p_j are the numbers of links for the two broken paths.
- 2 these options are better than the requirements.
- 3 may have some problems after the first failure.
- 4 referring to table C.3, F_P must be less than 0,4% and F_S must be less than 0,04%

Table C.6 – The options and the objectives

Option	Proportion of time in performance category a) Objective : $\geq 99,56\%$	Proportion of time in performance category b) Objective : $< 0,4\%$	Proportion of time in performance category c) Objective : $< 0,04\%$
Option 1	100%	0%	0%
Option 2	$(100 - F_P - F_S) \%$ (NOTE)	$(F_P + F_S) \%$ Note	0%
Option 3	$(100 - F_S) \%$	0%	S
Option 4	$(100 - F_S) \%$	$F_S \%$	0%
Option 5	$(100 - F_P - F_S) \%$ Note	0%	$(F_P + F_S) \%$

NOTE : these areas are outside the objectives.

It has to be noticed that the percentage allowed does not fit very well with our calculation of failure's probability:

- the first failure occurs with a higher probability than the value allowed (unavailability $q \times 0,4\%$ and requirement 0,4%);
- the second failure's probability is lower than the value allowed.

Therefore the use of performance category b) after one failure is unlikely to be satisfactory, but the use of performance category c) after 2 failures is adequate.

The preferred option is number 3. The options 1 and 4 have a better performance during the second failure, but this is not necessary. On the other hand, the options 2 and 5 may have some problem if the percentage of time for the first failure is over 0,4%.

The different options are described and commented in the next paragraphs. In the different cases, several implementations are possible since the parameters we can use are:

- the accuracy of the clock during the failure, which has a direct consequence on the number of slips in the non synchronized state
- the number of links providing a clock source.

C.5 Description of the five options

C.5.1 Option 1 : a) – a) – a)

In this option, the node of the private network does not generate more than 5 slips every 24 h, even in case of failures on one or two different branches bringing the clock from the public network

First implementation: the clock of the node is not synchronized by the public network, but it must be in accuracy class I.

Second implementation: the node is able to find an alternate path to the public clock after the first and the second failure. This is only possible if the node has at least three clock paths to the public ISDN, and some information (via signalling between nodes) on the actual quality of these routes.

This option, with at least 3 entries for synchronization is necessary for only 0,00004% of the failures.

C.5.2 Option 2 : a) – b) – b)

In the normal operation, the node is synchronized to be in performance category a). Then, after the first failure, it goes into performance category b), which means some degradation. This is due to the fact that the node is no longer driven by a primary reference clock, and runs with its own clock, in free running or in hold over mode.

To operate in performance category b), a clock with accuracy class II is required.

C.5.3 Option 3 : a) – a) – c)

The node stays in performance category a) after the first failure (by finding an operational alternative path to a primary reference clock, if any) and then, after the second failure, it enters performance category c).

To operate in performance category c), a clock with accuracy class III is enough.

C.5.4 Option 4 : a) – a) – b)

First implementation : after the first failure, the node finds an alternate path to a primary reference clock. Information about the quality of the alternate path is needed. After the second failure, the node stays in performance category b), which is not required.

Second implementation : the operator can guarantee that the proportion of time for one failure is less than 0,04%. It does not seem to be the case anywhere in Europe.

NOTE : this option is significantly better than the requirements.

C.5.5 Option 5 : a) – b) – c)

The only possible scenario for this option is the following :

- the node starts in performance category a): it is synchronized through its main input to a primary reference clock;
 - when a failure occurs on that input (less than 0,4% of the time), it switches over to a second input.
- The node can take such a decision because the node behind the second input guarantees performance category b): this specific node is still synchronized by the public, or in hold over mode with an accuracy class II;
- when this second input has a failure, the node is in performance category c).

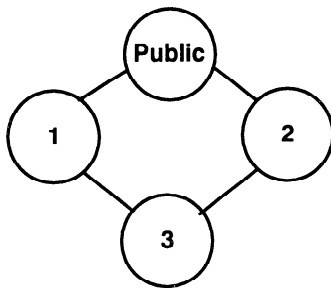
It has to be noted that with this option, the previous node must guarantee performance category b). It means that this previous node uses options 2 or 4.

Annex D (informative)

Synchronized Private Network Examples

D.1 Example with a small private network

The following example is a very simple one. Node 3 has no visibility to the public network.



Node	Strategy 1	Strategy 2	Strategy 3.1	Strategy 3.2
1	yes	yes (1)	yes (1)	no (2)
2	yes	yes (1)	yes (1)	no (2)
3	yes	yes	yes (1)	no (2)

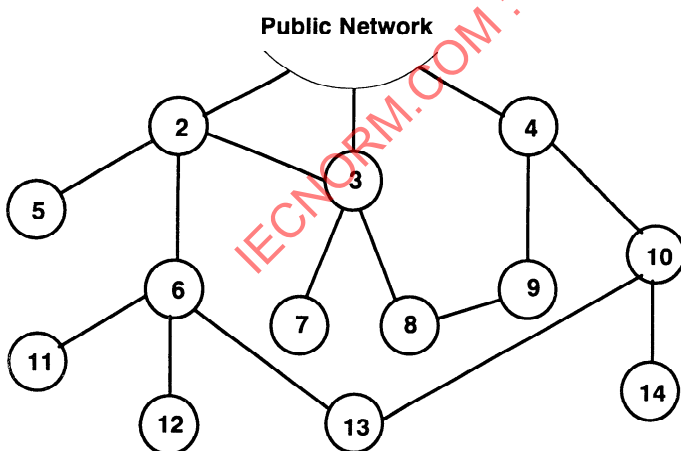
NOTES

- 1 (1) shows recommended solutions.
- 2 Except if the appropriate masters use strategy 1 or strategy 2.

Figure D.1 – Example 1

- Strategy 1 : can be used for each node. Due to the cost of class I clocks it is not a recommended solution.
- Strategy 2 : can also be used for the 3 nodes, and can be recommended for the nodes 1 and 2, but may be excessive for node 3 which does not drive any other node.
- Strategy 3.1 : is useable and recommended for the 3 nodes.
- Strategy 3.2 : cannot be used in this network by more than one node.

D.2 Example with a big private network



Node	Strategy 1	Strategy 2	Str 3.1	Str 3.2
1				
2	yes	yes (1)	yes (1)	no
3	yes	yes (1)	yes (1)	no
4	yes	yes (1)	yes (1)	no
5	yes	yes	yes	yes (1)
6	yes	yes	yes (1)	no
7	yes	yes	yes (1)	yes
8	yes	yes	yes (1)	no (2)
9	yes	yes	yes (1)	no (2)
10	yes	yes	yes (1)	no
11	yes	yes	yes	yes (1)
12	yes	yes	yes	yes (1)
13	yes	yes	yes (1)	no (2)
14	yes	yes	yes	yes (1)

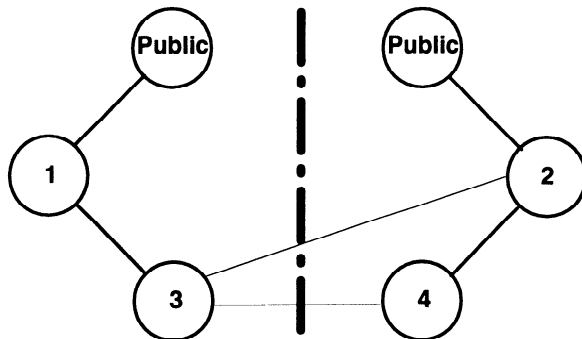
NOTES

- 1 (1) shows recommended solutions.
- 2 Except if the appropriate masters use strategy 1 or strategy 2.

Figure D.2 – Example 2

D.3 Example with two different public clock sources

The following example describes the situation of a single private network connected to two different public clock sources (within the same country or not). The table below shows the strategies that can be used for each node.



Node	Strategy 1	Strategy 2	Str 3.1	Str 3.2
1	yes	yes (1)	yes (1)	no
2	yes	yes (1)	yes (1)	no
3	yes	yes	yes (1)	no (2)
4	yes	yes	yes (1)	no

NOTES

- 1 (1) shows recommended solutions.
- 2 Except if the appropriate masters use strategy 1 or strategy 2.

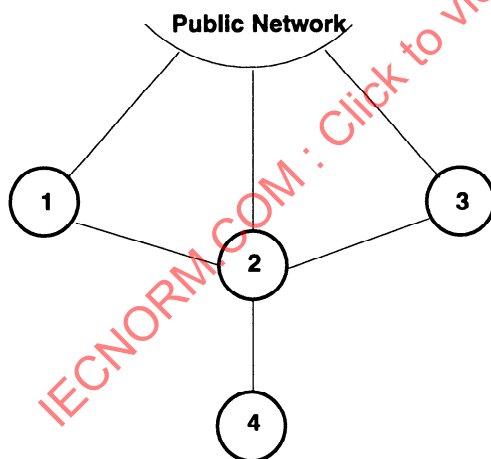
Figure D.3 – Example 3

This example shows that two islands of the private network (nodes 1 and 3 on one side, and nodes 2 and 4 on the opposite side) can work plesiochronously.

But the split of the islands can be different. It is the responsibility of the designer of the private network to negotiate with the appropriate public network operators.

D.4 Example with a transit node

This example describes the situation where transit functions are provided and the PINX clock may not be locked to the connected PINX's clock. Two links are slipping independently. In effect, PINX 2 in this example needs to have a clock twice as accurate as those performing end only functions.



Assume that the nodes 1, 2 and 3 use strategy 2 (i.e. they only use the single input from the public network as master).

Take the case where a call is made from node 1 via node 2 to node 3 to access the public network. In the case where the node 2 is free running, both links 1–2 and 2–3 slip. Since nodes 1 and 3 are locked to the same source, both slip at the same rate. To achieve better than 30 slips per hour overall, there must be no more than 15 per hour on each link. To meet this the clock must have an accuracy of $0.5 \cdot 10^{-6}$.

If in a large network, a single call path could encounter two non-synchronized exchanges, even better performance might be required, since then there would be 3 or more links slipping.

If these cases are rare, and result for instance from more than one single failure, they may entirely fall into performance category c).

Figure D.4 – Example 4

Annex E (informative)

Slave Clock Performance Measurement Guidelines

E.1 Slave Clocks considerations

Slave clocks are required in a digital synchronization network to provide a time keeping function at each node. When a network is synchronized, the relative time error between slave clocks in the network is maintained within tight bounds. Time error variations in slave clocks will be manifested in variations in slip buffer fills where one node receives a synchronous payload from a remote node. If relative time variations are constrained properly, buffer slips will be kept at an acceptable level.

Slave clocks can be viewed as providing two basic functions:

- (1) Receive from the incoming reference a good estimate of the original master node timing;
- (2) in the absence of reference, attempt to maintain good time keeping with respect to the master clock.

The first function requires that a slave clock attempt to reproduce the original master node timing from an impaired reference signal. Slave clocks inherently function as low pass timing filters. The short-term slave clock can be viewed as the superposition of the short-term stability of the local oscillator on the long term timing of the master clock.

In addition, slave clocks need to bridge short interruptions in the incoming reference. Synchronization reference is carried via digital transport facilities. It is normal to expect some level of impairment on these facilities. Slave clock performance is adversely affected when the reference-carrying facilities experience disruptions. These disruptions need not be major outage events to adversely affect slave clock performance. For example, when a slave clock sees an error burst condition on the incoming reference line, it may consider the phase data extracted from the line as suspect. The slave clock could suspend updating the control loop during the suspect interval. After the suspect interval is over, the slave clock typically performs a process termed phase build-out. This process attempts to restore the phase error in the loop in such a way that no residual error results from the disruption. However, there is inevitably some small residual error.

In reality, reference performance can include a significant number of disruption events. Error burst events such as Severely Errored Seconds (SES) are known to occur on links in the order of 10 to 100 events per day. Given that slave clocks in the network are adversely affected by some of these disruption events, it is wrong to assume that a slave clock is normally operating in phase lock with the incoming timing reference. In fact, slave clocks in the network are constantly degraded to some extent by the disruptions that typically occur on reference carrying facilities.

During long outages, a slave clock attempts to maintain its time keeping performance. For class III clocks, the oscillator is permitted to rapidly return to a free running condition. Class II clocks are required to have a holdover capability. Generally, the behavior of the crystal oscillator can be learned, and a predictor can be applied to compensate for predictable behavior such as drift. In practice, prediction is not typically used, and holdover is achieved using an estimate of the slave oscillator's frequency offset compared to the incoming reference to compensate for the initial offset.

The slave clock model is best understood by considering three categories of slave clock operation:

- (1) ideal operation;
- (2) stressed operation;
- (3) holdover operation.

E.1.1 Ideal Operation

Ideal operation represents an idealized operating condition which would not be typical of real network operation. In ideal operation, the slave clock experiences no interruptions of the input timing reference. Under such conditions, the slave clock would be expected to operate in 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 output timing signal is determined by the short-term stability of the local slave clock time base. In the absence of reference interruptions, the stability of the output timing signal behaves asymptotically as a white noise PM process as the observation period is increased to be within the tracking bandwidth of the PLL. The output of the slave clock can be viewed as a superposition of the high frequency noise of the local oscillator riding on the low frequency portion of the input reference signal. In phase-locked operation, the high frequency noise is bounded, and is uncorrelated (white) for large observation periods relative to the bandwidth of the phase-locked loop.

Under ideal conditions, the only nonzero parameter of the model is the white noise PM component.

E.1.2 Stressed Operation

This category of operation reflects the performance of a slave clock under actual network conditions. In the presence of interruptions, the stability of the output timing signal behaves as a white noise FM process as the observation period is increased to be within the tracking bandwidth of the PLL.

The presence of white noise FM can be justified based on the simple fact that, in general, network slave clocks extract time interval rather than absolute time from the time reference. An interruption is by nature a short period during which the reference time interval is not available. When reference is restored, there is some ambiguity regarding the actual phase difference between the local slave clock and the reference. Depending on the sophistication of the slave clock phase build-out, there can be various levels of residual phase error which occur for each interruption. There is a random component which is independent from one interruption event to the next which results in a random walk in phase; i.e., a white noise FM noise source.

In addition to the white noise FM component, interruption events can actually result in a frequency offset between the slave clock and its reference. This frequency offset results from a bias in the phase build-out when reference is restored. This is a critical point. The implication of this effect is that, in actual network environments, there is some accumulation of frequency offset through a chain of slave clocks. Thus, slave clocks controlled by the same primary reference clock are actually operating plesiochronously to some degree.

E.1.3 Holdover Operation

This category of operation accounts for the infrequent times when a slave clock loses reference for a significant period of time. Assume that an interruption of reference exceeding 10 seconds in duration would mark the onset of a holdover event. In holdover, the key components of the slave clock model are the frequency drift and the initial frequency offset. The drift term accounts for the environmental effects (e.g. temperature and power supply voltage), and aging associated with quartz oscillators. The initial frequency offset is associated with the intrinsic settability of the local oscillator frequency.

The measurement methodology proposed in this annex is structured to take into account the behavior of slave clocks in real network environments. This annex presents a model for characterizing actual slave clock performance. A key element of this model is that it reflects the stress conditions in real networks under which slave clocks would be expected to perform acceptably. This annex also presents a standard methodology for measuring slave clock performance. The measurement strategy is to be able to derive the values of the model parameters for the given clock under test. Once slave clock performance can be described by a set of parameters values, it is relatively easy to develop recommended slave clock performance specifications.

E.2 Test Configuration Guidelines

The objective of the test procedure is to be able to estimate the parameters in the slave clock testing arrangement is shown in figure E.1. The components and their interconnections are described next.

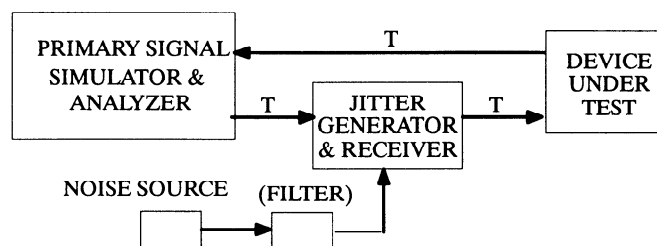


Figure E.1 – Standard test configuration

E.2.1 Reference Clock

The test configuration is designed to provide the slave clock under test with a digital reference timed from a stable reference oscillator. In slave clock testing, the relative phase – time compared to the reference input is critical. In holdover testing, the longer – term stability and drift of the reference oscillator is important. Thus, the absolute accuracy of the reference input is not critical. It is important that short – term instability of the reference oscillator be small to ensure low measurement noise and a low background tracking error in the control loop of slave clock being tested. For testing, a cesium reference clock has been employed as a reference. The background tracking noise for typical loop bandwidths is much less than 1 ns using the cesium reference.

E.2.2 Digital Reference Simulation

The testing arrangement is designed to provide an impaired digital reference (see 2.2.2) to the slave clock for stress testing of the slave clock. To accomplish this, a digital signal simulator and analyzer that has the capability to be externally synchronized is employed. The jitter produced from these synthesizers would be expected to be less than 1 ns rms.

The primary signal simulator is programmed to produce the desired interruption events to stress the slave clock. The digital signal is next bridged through a jitter generator and receiver. The jitter generator is used to insert background jitter to the digital signal. It is important to simulate a realistic level of background jitter for several reasons. Primarily, when interruptions occur, the background jitter can be a major source of phase build-out error as the synchronization unit attempts to bridge the interruption. Secondly, the jitter transfer characteristics of the slave clock can be evaluated.

The jitter generation unit is provided with an external jitter modulation input. The jitter signal used is band limited white noise. The main reason for low pass filtering the jitter is to avoid producing bit errors from high frequency alignment jitter. The jitter power needs to be set to reflect the input jitter levels defined in 2.1.1. It is important that sinusoidal jitter be avoided as a test jitter input, because it is not representative of actual network conditions.

E.2.3 Output Timing Signal Recovery

To test a slave clock, reference input is provided from the output of the jitter generator. To recover the output timing signal from the slave clock, an outgoing digital signal is selected from the unit controlled by the slave clock under test. This digital signal is connected to the receive portion of the Primary Rate Signal Simulator and Analyzer. In this unit, the receiver timing function is decoupled from the transmit timing used in the generator. The receiver extracts a frame timing signal from the input signal and provides this timing signal at an external port. This Frame timing signal is phase coherent with the outgoing timing from the slave clock under test.

E.3 Test categories

To adequately characterize the performance of a slave clock, a series of tests must be performed. In general, the tests fall into the three categories of operation described hereafter.

E.3.1 Ideal Testing

The purpose of this testing is to obtain a baseline performance measure for a slave clock. The model predicts that slave clocks under ideal conditions would likely produce a white noise PM phase instability. This white noise PM would be expected to be small because it represents the best case performance of a slave clock (clearly less than 1 μ s based on current relative TIE output requirements). It needs to be measured in the presence of realistic levels of network input jitter to assure acceptable jitter transfer.

In the test procedure described, the maximum bandwidth of the measurements is 1 Hz. In some slave clock designs, there is significant noise between this 1 Hz cutoff and the 10 Hz cutoff associated with jitter. It is important to evaluate the jitter in this band. One approach is to use a jitter test set externally referenced to the reference clock described in E.2.2. With a stable external reference, some jitter test sets can extend jitter measurement bandwidth down to 1 Hz.

E.3.2 Stress Testing

This area of testing is critically important to adequately evaluate slave clocks. The difficulty in this test is selection of the appropriate interruption events. For some slave clocks, any event that appears as a severely errored second will produce a phase buildout event. In some slave clocks, any outage or spurious noise spikes will perturb a counter in the phase detection, thereby producing a spurious phase hit which may or may not be phase build-out depending on its severity. On the other hand, slave clocks can be designed to observe the framing pulse position to extract phase. In such slave clocks, an interruption need not produce a phase build-out event unless there is an actual shift in the framing pulse position (for example a protection switch event). These general observations demonstrate the difficulty in selecting interruption test criteria. In addition, the nature of interruption events produced in networks is difficult to determine. As already mentioned, the allowable magnitude of severe error bursts in the network is quite high. It would be unwarranted to assume that actual link performance will be substantially better than 10 SES/day. Assuming 10 SES reflects a reasonable level to expect, the next question is what fraction of these SES events will produce degradation effects. In the absence of data, it is reasonable to assume a worse case scenario in which each interruption event produces a phase build-out event.

It is proposed that one minimum stress test which needs to be performed is to simulate an SES event with a short outage on the order of 100 ms at a rate of 10 SES per day in the presence of background input jitter. Typically an outage of this magnitude will force a slave clock to attempt to phase build – out without switching references. Precautions need to be taken to prevent reference switching under this testing scenario. Other stress inputs need to also be considered in evaluating a slave clock.

E.3.2.1 – Error Burst

An error burst in which the underlying timing waveform is not perturbed can be simulated. Under this condition, it would be advantageous for a slave clock not to phase build – out. Such a test would gain in importance if it is determined that the majority of error burst events are actually pure data errors with no perturbation in timing.

E.3.2.2 – Phase Hit

Phase hits are produced by protection activity as well as from other slave clocks. Phase hits are interruption events that would be expected to either force a phase build – out event or be inadvertently followed by the slave clock. In either case, they will degrade a slave clock's performance. This is an area for further study.

E.3.2.3 – Restart Events

Restart events are a phenomenon associated with certain slave clocks. A restart event is associated with a slave clock giving up its current state, and defaulting back to its initial conditions. The result is a transient event which can be significant. Restart events need not happen during normal slave clock operation and thus would not likely be included in a general slave clock testing plan. However, it is important that this behavior be better understood and controlled.

E.3.2.4 – Frequency Hit

It is important that slave clocks not follow references that exhibit large frequency hits. However the ability to detect frequency hits is closely tied to the selection of the tracking bandwidth of a given slave clock PLL. The solution to the problem will depend on the degree to which the stability of various slave clocks in a network can be standardized.

E.3.3 Holdover Testing

In holdover testing, the objective is to estimate the initial frequency offset and the drift of the slave clock model. The initial frequency offset is dependent on the accuracy of the frequency estimate obtained in the control loop, and the frequency setability of the local oscillator. It is important to test holdover from a reasonable stress condition prior to holdover to capture the control loop's capability of obtaining an accurate frequency estimate.

In determining the drift estimate, one critical factor for a quartz oscillator is that it typically takes observation intervals lasting over days to obtain a statistically significant drift estimator. This is a hard reality that cannot be avoided. In addition, attention must be placed on the temperature and power supply conditions maintained during the test.

Annex F (informative)

Signalling for management of synchronization

F.1 Presentation

Strategy 3.1 consists of choosing the clock references, assigning configuration parameters, describing reactions of a node in case of failures and the definition of signalling information to be exchanged

F.1.1 Configuration parameters

Two configuration parameters have been identified :

- (1) Connection or no connection to a public node.
- (2) Potential for a node to become the master in a dual relationship at both ends of a transparent private link.

The direction of the clock enslavement depends on these parameters, they cannot be changed during the state machine progress; a reset of the node state machine is mandatory.

Interworking with PINXs without signalling is possible : these PINXs are said to send the normal value. A PINX without signalling but using strategy 1 or strategy 2 shall be considered as a public node. This shall be taken into account during the configuration of the adjacent nodes.

F.1.2 Reactions of the node

At the end of the configuration phase, each node of the network has clock references and the information necessary to know which signalling information it has to send.

According to all the combinations of the configuration parameters, seven different states are identified : state 1 to state 7.

F.1.3 Reference clock switching and restoral

Clock switching and restoral shall comply with subclause 1.6.. Automatic switch back is permitted but not described in the SDL diagrams.

F.2 Description of the states

F.2.1 Initial states

At the end of the configuration phase, a node enters one of seven possible states : State 1 to State 7. Their defined states depend on the type of link (public or private), and on the potential of the node to become the master.

In these states, nodes are enslaved through their main input p.

The following table presents the configuration parameters of a node in each initial state of the state machines.

Table F.1 – Node configuration parameters

	link type behind p	Potential to become master of the node behind p	link type behind s	Potential to become master of the node behind s
State 1	public		public	
State 2	public		private	yes
State 3	public		private	no
State 4	private	yes	private	yes
State 5	private	yes	private	no
State 6	private	yes	private	no
State 7	private	yes	private	yes

F.2.2 Slave states

After having lost p, when a node is enslaved through the secondary input s, it runs into a slave state.

In the slave states, events are taken into account according to the logical configuration parameters of the node. Three slave states have been identified, they are called Slave 1, Slave 2, and Slave 3.

F.2.3 Autonomous state

After having lost p and s, a node is autonomous in a free running or hold over mode.

F.2.4 Wait states

In the wait states, a node is waiting for an answer or an acknowledge from the other node (signalling exchange). The enslavement input (p or s) is not yet modified.

F.3 Description of the events

F.3.1 Failure of links

These events occur when the signal behind p or s is no more a valid clock source. They are called "p fails" or "s fails" in the state machine.

F.3.2 Signalling information

Info	Comments
Hold on request	The node sending that information shall become master if no better clock source is found in the private network.
Yes	Request ACK
Idle	no change in the state machine is required
No	Request NACK
Enslavement request	Information send by a node when it has lost its master and when it has not the potential to become the master.
Free running	A node working with its own clock informs all the adjacent nodes with this signal.
Default value	The default signal which allows interworking with a PINX without signalling

F.3.3 Time out

A time out shall be defined to avoid dead lock situations when a node does not receive any answer to its signalling message. This time out shall be set when the signalling message is sent, and shall be reset when the answer to the message is received. At the end of the time out, the node shall enter a free running mode, and shall work with its own clock.

F.4 SDL representation of the state machine

The relevant exchanges are represented in the following state machine. Only the possible transitions are shown.

If the node, behind a link is a public node or a private node without signalling the default value is sent. This does not appear in the state machine.

Furthermore, unless stated in the state machine, default signal shall be sent ("idle" during state 1 to 7, slave and wait, and "free running" in the autonomous state)