

DISTRIBUTED COMPUTING COLUMN

Stefan Schmid
TU Berlin & T-Labs, Germany
`stefan.schmid@tu-berlin.de`

Fault-tolerant Distributed Systems in Hardware

Danny Dolev (Hebrew University of Jerusalem)

Matthias Függer (MPI for Informatics)

Christoph Lenzen (MPI for Informatics)

Ulrich Schmid (TU Vienna)

Andreas Steininger (TU Vienna)

Very large-scale integrated (VLSI) hardware designs can be seen as distributed systems at several levels of abstraction: from the cores in a multicore architecture down to the Boolean gates in its circuit implementation, hardware designs comprise of interacting computing nodes with non-negligible communication delays. The resulting similarities to classic large-scale distributed systems become even more accented in mission critical hardware designs that are required to operate correctly in the presence of component failures.

We advocate to act on this observation and treat fault-tolerant hardware design as the task of devising suitable distributed algorithms. By means of problems related to clock generation and distribution, we show that (i) design and analysis techniques from distributed computing can provide new and provably correct mission critical hardware solutions and (ii) studying such systems reveals many interesting and challenging open problems for distributed computing.

1 Introduction

Very large-scale integrated (VLSI) circuits bear several similarities with the systems studied by the distributed computing community:

- They are formed by an increasing number of interacting computing nodes.
- Communication delays are not negligible.
- The cost of communication, such as area and power consumption, is not negligible.

In fact, this view is correct at different levels of abstraction. We will elaborate on two such levels that significantly differ from each other with respect to the computational power of the system's nodes, their means of communication, and the problems they solve.

(I) Viewed from a low-level perspective, every digital circuit is a network of logic gates with delays, which continuously compute their current output state from their input history and interact via binary-valued, continuous-time signals. We stress the fact, however, that this is already a (convenient) abstraction, as real gates are electronic devices that process analog (continuous-valued) signals: A signal that is above a certain threshold voltage is considered high, otherwise low. Whereas analog signals (like a good clock signal) that swing fast from low voltages to high voltages are represented reasonably well by the resulting digital abstraction, this is obviously not the case for every signal: Just consider an analog signal that stays very close to the threshold voltage and just, e.g., due to very small noise, occasionally crosses it. It will turn out that this modeling inaccuracy causes serious problems both for synchronization and fault-tolerance, which must be considered explicitly.

Analogously to distributed computing, there are two fundamentally different ways to design digital circuits (i.e., algorithms in hardware), which correspond to synchronous and asynchronous algorithms in distributed computing.

The classic design paradigm relies on the synchronous computational model. It abstracts away the timing of gates and interconnects by considering gate outputs only at predetermined instants dictated by a central periodic clock signal. While this approach allows the designer to solely concentrate on the stable outputs of a network of gates, it relies critically on the guarantee that all signals have settled and all transients have vanished at the occurrence of the next clock transition. Inherently, such designs run at the speed of the clock period that is determined from worst-case bounds on gate and interconnect delays. Due to increasingly pronounced delay variations [52, 84] this results in highly conservative bounds and thus in considerable performance loss.

In contrast, designs that do not rely on the convenient discrete time abstraction provided by a clock signal are called *clockless* or asynchronous.¹ Such circuits must rely on different techniques to enforce some ordering between signal transitions. Suitable techniques range from aggressively timed circuits that explicitly use information on the delays along certain paths [80, 91, 97] to circuits that tolerate (certain) delay variations by means of some forms of handshaking. Prominent examples of the latter are *delay insensitive* (DI) circuits [72], speed-independent (SI) circuits and *quasi-delay insensitive* (QDI) circuits [74, 75]. While DI circuits are guaranteed to behave correctly in the presence of arbitrary gate and interconnect delay variations, SI resp. QDI circuits assume that all resp. certain signal forks in the interconnect are isochronic, i.e., have roughly equal propagation delays along all their fork teeth.

The robustness to delay variations in DI circuits comes at a high price, however: Martin [73] showed that the expressiveness of circuits that are DI at gate-level is severely limited. In fact, the only two-input gate allowed in such circuits is the C-Element, which is an AND gate for signal transitions; it produces a, say, rising transition at its output when it observed a rising transition at all its inputs. This clearly restricts the usability of DI circuits for real applications.

SI and QDI circuits, on the other hand, are Turing-complete [70]. Intuitively, the isochronic fork assumption guarantees that a gate whose output drives an isochronic fork implicitly performs a handshake with all its successor gates while just handshaking with one of its successors. A precise characterization of the conditions on the propagation delay that have to hold on certain paths in SI circuits was derived in [56].

(II) With the increasing number of computing nodes in *System-on-Chip* (SoC) and *Network-on-Chip* (NoC) architectures, problems that used to be relevant only in large-scale computer networks also become relevant within a single chip. Examples range from establishing a common time base over data communication and routing between nodes to load balancing.

In the hardware context, establishing a common time base among all nodes is of particular interest, because this sets the base for a synchronous computational model: Rather than being implemented entirely clockless, higher-level services like routing and load balancing could then also exploit synchrony properties. Unfortunately, however, the GHz clock speeds and sizes of modern VLSI circuits make it increasingly difficult to distribute a central clock signal throughout the whole circuit [43, 96]. Modern SoCs and NoCs hence typically rely on the *globally asynchronous locally synchronous* (GALS) approach [14], where different parts of a chip are clocked by different clock sources. Using inde-

¹We will use the term “clockless” in the following, as such circuits do not always allow for arbitrarily large and unknown delays.

pendent and hence unsynchronized clock domains would give away the advantages of global synchrony and also requires non-synchronous cross-domain communication mechanisms or synchronizers [59, 58]. A promising alternative is mesochronous clocking [79] (sometimes also called *multi-synchronous* clocking [93]) as it guarantees some upper bound on the skew between clock domains. In this article, we will thus focus on discussing methods for providing a common time in GALS architectures.

Fault-tolerance. Besides an increasing number of components and non-negligible communication costs both at gate and system level, there is a further trend that advocates the application of distributed computing methods for designing VLSI chips: the increasing susceptibility to failures. Indeed, fault-tolerance has been identified as a key challenge in the International Technology Roadmap for Semiconductors [52] for years. Besides the increasing susceptibility of nanometer VLSI technology to permanent failures caused by manufacturing process variations and excessive operating conditions (supply voltage, temperature) [62], steadily decreasing feature sizes and signal voltage swings also led to dramatically increased transient failure rates [16], caused by ionizing particles hitting the junction of transistors [6], electro-magnetic cross-talk between signal wires and supply-voltage variations caused by simultaneous switching activities [78, 85].

Unfortunately, even relatively simple faults unveil the very limited ability of the convenient digital signal abstraction to properly describe reality. For example, an out-of-spec output driver of a gate that drives a fork to two different gate inputs may be able to reach the threshold voltage at one input but not at the other, causing those to interpret the gate output inconsistently. Similarly, a *single-event transient* (SET) [6] caused by an ionizing particle that hits the output driver of such a gate may be visible at one input but not at the other, depending on the latter's input capacitances. It is hence apparent that classic benign failure models from distributed computing, where a message is either lost or transmitted correctly, do not cover such faults. In fact, faulty gates have to be assumed to potentially behave arbitrarily, i.e., Byzantine [86].

While there is a huge body of work on fault mitigation techniques at the technological level (like *silicon-on-insulator* (SOI) technology [76]) and gate level (like the SET-tolerant DICE latch [83]), keeping the overall error rate acceptable [89, 22] despite the tremendously increasing number of gates/cores on a single chip demands for additional architectural solutions. At the same time, solutions that require central knowledge of the current system state (i) become infeasible due to the high communication costs and (ii) would themselves form a single point of failure. Algorithmic solutions that use only local knowledge, studied by the distributed computing community for decades, are hence promising in this

context.

Classic architectural fault-tolerance approaches [87] like *Dual Modular Redundancy* (DMR) and *Triple Modular Redundancy* (TMR) fail in absence of a global time base, as it becomes unclear over which values to vote. Jang and Martin [53] adapted this method to QDI designs and applied it to build a micro-controller tolerating certain transient faults [54], in particular, *single-event upsets* (SEUs), where a state-holding device may flip its state due to a particle hit. Their solution duplicates each gate and adds two succeeding cross-coupled C-Elements whose inputs are connected to the outputs of the duplicated gates. In case of a spurious state transition of one of the duplicated gates, both C-Elements do not propagate the spurious output value until it is matched by the other gate's output also (which can be proved to eventually happen). While this method tolerates SEUs, it neither allows to tolerate SETs nor permanent faults.

Tolerating such faults necessarily requires to extend the circuit's control logic not to wait for all of its predecessors' outputs. In contrast to the AND-causality semantics of the C-Element, this requires OR-causality semantics. Interestingly, there has been research in this direction in a different context: In certain cases, a Boolean function's value can already be determined from a subset of its parameters. This fact can be used to speed up clockless circuits [17, 95]. Instead of waiting for all of a module's inputs to arrive, the module waits until its outcome is determined and then immediately produces a new output value. Care must be taken not to mix up current input data with lately arriving input data from a previous computation, however. The approach thus requires additional timing assumptions and ways to memorize which input values a module already took into account when computing its output.

A similar strategy can also be applied to design clockless circuits that tolerate a certain fraction of its input nodes to fail permanently. In particular, it has also been employed in the DARTS Byzantine fault-tolerant clock generation approach [49, 44] for mesochronous GALS architectures, which comprises a network of interconnected OR-causality gates whose outputs generate tightly synchronized clock pulses. The approach is discussed in more detail in Section 3.1.3.

The limit of the fault-tolerant hardware solutions discussed above is that they allow only a certain subset of the components to fail. Even if these components start to operate according to their specification again later on, their state may remain corrupted, preventing them from recommencing correct operation. For massive transient failures, which are likely to occur e.g. in space missions and may corrupt the entire system state, the above solutions are not adequate. To attack this problem, the concept of self-stabilizing distributed algorithms [33] has successfully been applied to digital circuits. A self-stabilizing circuit is guaranteed to eventually behave correctly again after its state was arbitrarily corrupted. For example, a self-stabilizing token passing algorithm implemented in clockless hard-

ware was presented in [11], and S. Dolev and Haviv [34] built and proved correct a self-stabilizing microprocessor.

From a robustness point of view, a combination of resilience to permanent faults and self-stabilization is most desirable: Appropriate solutions operate correctly in the presence of not too many permanent faults and even recover from massive transient disruptions. In [27, 28] the fault-tolerant and self-stabilizing pulse generation algorithm FATAL and its hardware implementation were presented and proven correct. This solution is discussed in more detail in Section 3.1.4.

Additional challenges. While we highlighted major similarities between VLSI designs and classic distributed systems, there are also important differences. In most cases, these advise against a naive implementation of distributed algorithms in hardware.

First and foremost, this is the absence of computationally powerful atomic actions at the gate level in clockless circuits: Explicit means such as handshaking must be employed to synchronize activities in such a circuit, which is not only costly but also imperfect in terms of synchronization accuracy. This, in turn, is also a source for a unique problem called *metastability*, which arises when a circuit must handle input signals that bear no well-defined timing relation to its own state transitions. Consider a simple R/W register in a shared memory system that may be read by a processor at the time it is written by some other processor. Distributed computing models based on regular or even atomic registers assume that either the previous or the newly written value is returned by the read. In reality, the situation is even worse than assumed for safe registers, which allow an arbitrary return value in this case: The register may reach a metastable state, which moves its output voltage to a non-digital level for an unpredictable time!

Another unique problem arises from the imperfect coverage of the digital abstraction for analog signals in the case of failures. In distributed computing, Byzantine behavior is considered the worst a component can do to a system. Unfortunately, in digital circuits, generating an arbitrary binary-valued continuous-time signal is not the worst behavior of a component. Rather, a component may produce an arbitrary analog signal on its output, e.g., an output voltage that remains very close to the threshold voltage arbitrarily long, which is actually one manifestation of metastability (creeping metastability) [7, 13]. We will discuss these issues in more detail in Section 2.

Structure of this article. We start in Section 2 with a discussion on the peculiarities of SoCs in comparison to classic distributed systems, and the challenges arising in the definition of an appropriate distributed system model. In Section 3,

we discuss the problem of obtaining a common time base for multi-synchronous GALS architectures, which is both fundamental to the solution of other problems and exemplarily captures the challenges of adapting distributed algorithms for use on a chip. The problem is divided into three parts: (i) Section 3.1 discusses the problem of generating synchronous pulses. (ii) Section 3.2 deals with establishing local counters. Here we provide more technical details, with the primary goal of illustrating how techniques from distributed computing find application in VLSI design. (iii) Section 3.3 finally is concerned with distributing the clock over a wider area. The work is concluded in Section 4.

2 Modeling Issues

While we pointed out the similarities of VLSI circuits and fault-tolerant distributed systems in Section 1, a simple migration of classic solutions in distributed computing is not favorable and most of the time even infeasible. The most prominent obstacles are:

(i) Gates continuously compute their output state from their input states. They generate events, i.e., binary transitions, in a fully parallel way and are capable of very simple computations, such as the logical AND of two binary inputs, only. Any kind of event sequencing and atomic actions that group several binary transitions into more powerful computations requires explicit synchronization between the concurrently operating gates, e.g., by handshaking or timing assumptions. Note that this includes even “simple” computations such as the sum or the product.

(ii) Communication and computation is costly, especially if the proposed solution is meant to “only” provide low-level services to the application running on top. For example, clock generation algorithms must not use more than a few wires between nodes to be able to compete with classic clock distribution networks. Exchange of data, even a few bits, requires parallel or serial coding and decoding logic and thus typically cannot be afforded for low-level services. Rather, solutions must resort to signaling a few status bits only. Exchanging data of more than, say, 32 bits, is usually also difficult for high-level services.

(iii) Non-digital low-level effects must be taken into account. Every binary valued model necessarily abstracts from the analog signals in real gate implementations. While it is perfectly valid to resort to binary abstractions most of the time, these models come to their limits when synchronization and failures enter the picture: Marino [71] showed that any bistable element, e.g., a binary memory cell, may get stuck arbitrarily long in between its two stable states. This may result in spontaneous, unpredictably late transitions on its output, and even in an inconsistently perceived input at multiple successor gates. While classic designs prevent these scenarios by ensuring that certain timing constraints on the input signals

are not violated, this is not always possible and inherently cannot be assumed in fault-tolerant circuits.

In order to be able to predict the behavior of a circuit and reason formally about its correctness and performance at early design stages, i.e., before fabrication, a suitable circuit model is required. Clearly, any such model should be sufficiently simple to support fast predictions and formal analysis, while at the same time ensure that the results reflect reality sufficiently accurate. We will briefly sketch common approaches.

Discrete time state machines. Synchronously clocked circuits of any kind can be modeled by a discrete-time, discrete-value synchronous communicating state machine model, for which very efficient and fast timing prediction and formal analysis tools are available. Unfortunately, this abstraction does not cover all existing circuits. This is obvious for clockless circuits, but also for the timing analysis of clocked circuits, which is mandatory for validating the clock timing requirements for justifying the synchronous abstraction. The latter is particularly important with the advent of aggressively timed high-speed synchronous circuits, where clock speed must be traded against the increasing rate of manufacturing errors and other sources of timing failures. In that case one has to resort to continuous time models.

Continuous time models. Arguably, the most accurate models used in circuit design today are fully-fledged continuous-time analog valued ones as, e.g., instantiated by Spice [81]. However, excessive analog simulation times prohibit its use for analyzing more than a fairly small part of a VLSI circuit, over fairly short periods of simulated real-time. Discrete-value models, and specifically binary-valued ones, are hence an attractive alternative. Modern digital design approaches e.g. based on description languages such as VHDL [5] incorporate digital timing simulators that are typically based on zero-time Boolean gates interconnected by non-zero-delay channels. Popular instances are simple pure (i.e., constant-delay) and inertial delay channels (i.e., constant-delay channels that suppress short input pulses) [94], but also more elaborate ones like the Delay Degradation Model (DDM) [8] or the empirical Synopsis CCS Timing model [92]. Continuous time, discrete-value models can be either state-based or trace-based, as detailed in the following.

Clockless, state-based models. At the gate level, clockless circuits are typically modeled by a binary state vector representing the global circuit state and, potentially time-annotated, guard-action pairs [4, 72] that describe the gates. An execution, i.e., signal trace, of the circuit is a sequence of global states over time

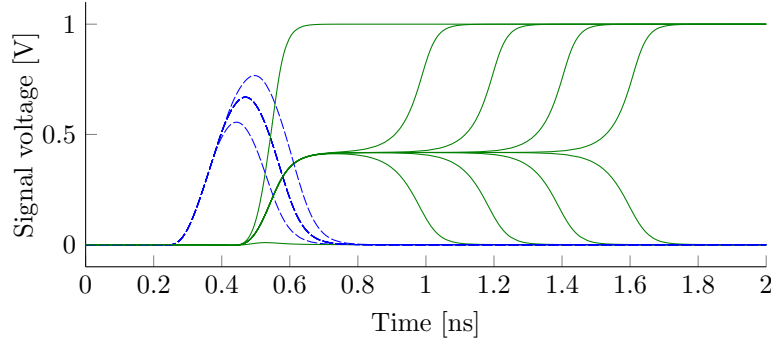


Figure 1: Analog simulation traces of a storage loop (e.g., a 2-input OR gate with the output fed back to one input) that can persistently memorize a high state of its (other) input. The blue dashed curves (the medium line actually corresponding to 8 almost identical pulses) show the real analog shape of short input pulse signals of different duration, the solid green ones give the corresponding output signals of the storage loop.

generated by a parallel execution of the guard-action pairs. Note that executions need not necessarily be unique, accounting for potential delay variations within the circuit. Like the models used in classic distributed computing, such as the Alur-Dill Timed Automata [2] and the Timed I/O Automata by Keynar et al. [55], these models all act on the explicitly given state-space of the underlying system.

While this view serves as a good level of abstraction in clockless circuits operated in closed environments, it comes to its limits when signals do not necessarily stabilize before they change again. For instance, consider a gate that produces a very short low-high-low pulse at its output: In reality, this means that the gate driver circuitry has only started to drive the output signal to high when it is turned off again. This results in a short, potentially non-full-swing waveform that may quite unpredictably affect the subsequent gate. An example is shown in the blue dashed pulse shape in Figure 1.

In this example, the subsequent gate is a *memory flag*, which persistently memorizes a high state at its input, until it is reset again to 0. A straightforward implementation is given by a *storage loop*, e.g. consisting of a 2-input OR gate with its output fed back to its other input. The solid green lines represent the output signals of the storage loop corresponding to the blue dashed inputs. The largest input pulse causes the storage loop to flip to the high state immediately, while the smallest one does not cause any effect on the initial low output. The medium input pulse, however, which actually represents 8 different ones that differ only marginally from each other, causes the loop to enter a metastable state: The input pulses are too short to allow the storage loop, which has some short but non-

zero delay, to settle to a stable value. Depending on minor variations of the input pulses, the metastable state eventually resolves after some unpredictable and possibly large resolution time. The memory flag does not operate as intended during this time, possibly affecting the downstream logic in a similar way.

Such situations cannot be prevented from happening in open environments in which a circuit cannot control all of its inputs. The same holds in fault-tolerant circuits, where the signals provided by faulty nodes may be arbitrary. Thus they must be reasonably covered by an appropriate digital circuit model. Unfortunately, however, this is not the case for any model used in modern timing circuit simulators today. Besides the complete lack of modeling and analysis support for fault-tolerance, it was shown in [48] that none of the existing analytic channel models, including the popular pure and inertial delay channels [94] as well as the DDM model [8], faithfully model the propagation of short pulses in physical circuits. Specifically, it has been shown that these models are inconsistent with possibility and impossibility results concerning the implementation of a one-shot inertial delay channel: a channel that, like an inertial delay channel, suppresses short pulses, but is required to work correctly only in presence of a single input pulse (one-shot).

Recently, however, a candidate delay model [47] based on *involution channels* has been proposed that does not have this shortcoming: It is not only consistent with the theoretical results on the one-shot inertial delay problem [48], but also achieves a promising level of accuracy [82]. As a consequence, there is a prospect of eventually identifying a tractable delay model that can form the basis for a comprehensive modeling framework for digital circuits.

Clockless, trace-based models. Existing frameworks for designing clockless digital circuits also have shortcomings at higher abstraction levels. In particular, we are not aware of any modeling framework (except the one we proposed in [27]) that supports fault-tolerance. Instead of blowing up the state space of existing state-based models – like Alur-Dill Timed Automata [2], Lamport’s TLA [63], Timed IO Automata by Keynar et al. [55], discrete abstractions for hybrid systems [3], or state-space-based control theory [64] – with error states and/or using non-determinism or probabilistic state transitions for modeling faults, it advocates the use of solely trace-based models, which focus on the externally visible behavior of a circuit only.

Examples of earlier trace-based approaches are Ebergen’s trace theory for clockless circuits [37] and Broy and Stølen’s FOCUS [12]. In both approaches, a module is specified exclusively in terms of the output signal traces that it may exhibit in response to a given input signal trace, without referring to internal state. The trace-based approach introduced in Dolev et al. [27] allows to express tim-

ing conditions via (dense) real-time constraints relating input/output signal transitions, and supports fault-tolerance by allowing (sub-)modules to behave erroneously, i.e., deviate from their specified behavior according to some fault model (e.g. Byzantine behavior [86]). It provides concepts for expressing the composition and implementation of circuits by other circuits, which also allow to rigorously specify self-stabilizing [33] circuits. The model has been used to precisely specify the modules of the Byzantine fault-tolerant and self-stabilizing FATAL clock generation algorithm, which will be described in Section 3.1.4, at a level of abstraction that allows for a direct implementation in hardware.

Compared to state-based approaches, it may be more involved to apply a behavioral approach, in particular, in settings like fully synchronous or fully asynchronous systems, where state-space-based descriptions are reasonably simple. After all, in general, it is even difficult to decide whether behavioral specifications match at interface boundaries [21]. On the other hand, it is much better suited for systems with a complex evolution of the system state over time and/or in which the internal state of system components is too complex or even unknown, which is typically the case for self-stabilizing algorithms. Another attractive side-effect is the inherent support of hierarchical design using (pre-existing or yet to be built) building blocks, without the need to define interface state machines. One can easily build a system and/or its *model* in a modular way, by composing sub-components and/or their models, whose implementation can be provided (typically by different designers) and verified at a later stage.

Nevertheless, it is understood that behavioral constraints translate into appropriate constraints on the modules' states implicitly. Although making this relation explicit is not part of our modeling framework, this definitely is part of the effort to implement a module and to prove that it indeed exhibits the specified behaviors. The latter may of course involve any appropriate technique, e.g. timed automata [2] and related verification techniques [1].

Open problems. Although the model of [27] explicitly avoids metastable upsets in fault-free executions, it cannot deal explicitly with metastable upsets and metastability propagation. The work by Ginosar [51], which provides several examples of synchronizer circuits where current prediction models drastically underestimate the probability of metastable upsets, shows the importance for such an extension. The challenge here is to bound metastability resolution times and propagation effects, potentially in a probabilistic manner, to be able to quantify upset probabilities and stabilization times.

Besides the need to extend the model [27] by standard tools like simulation relations and abstractions, the integration with a faithful digital circuit model like [82] remains a challenge. The ultimate goal is a comprehensive modeling frame-

work for modern digital circuits, which facilitates (semi-automated) formal verification of circuits, correctness proofs and accurate performance analysis as well as design parameter synthesis, ideally via supporting tools.

3 Generating and Distributing a System Clock

Simulating a synchronous system in the presence of both transient and permanent failures is a challenging task. The traditional approach to generating and distributing the clock signal, a *clock tree* [43], follows the master/slave principle: the signal of a single quartz oscillator is distributed to all logical gates by means of a tree topology. This approach is trivially self-stabilizing, but it must be abandoned due to the possibility of permanent faults; a failure of the tree's root (i.e., the oscillator) or a node close to the root breaks the entire system.

In the severe failure model considered in this article, this fundamental problem was first studied by S. Dolev and Welch [35, 36]. It was motivated by the observation that the assumption of only a fraction of the node being affected by transient faults is too optimistic for the typically long mission times (e.g., space missions) during which clock synchronization has to be provided.

The ultimate goal is to simulate synchronous rounds that are consistently labeled (at all correct nodes) by a round counter modulo $C \geq 2$, where C usually is fairly large. Dolev and Welch give a protocol that stabilizes in exponential time. While this does not seem very exciting at first glance, at this time the big surprise was that the problem was solvable at all!

For the sake of clarity of presentation, let us break down the task into three subproblems:

1. *Pulse Synchronization:* Simulating unlabeled synchronous rounds in a system with bounded communication delay and local clocks of bounded drift.
2. *Counting:* Computing consistent round counters in a synchronous system with unnumbered rounds.
3. *Clock Distribution:* Distributing pulses and/or round numbers efficiently, i.e., using a low-degree topology.

We remark that it is not imperative to follow this structure when solving the problem. However, the discussion will reveal why this is a fairly natural decomposition of the task.

3.1 Pulse Synchronization

In the pulse synchronization problem, we are given a fully connected system of n nodes, $f < n/3$ of which may be Byzantine faulty. Nodes communicate by messages that are delayed between 0 and 1 time units,² which also accounts for any delay incurred by local computations. Each node $i \in \{1, \dots, n\}$ is equipped with a local clock $C_i : \mathbb{R}_0^+ \rightarrow \mathbb{R}$ of bounded drift, i.e.,

$$\forall t > t' : t - t' \leq C_i(t) - C_i(t') \leq \vartheta(t - t')$$

for a constant $\vartheta > 1$. As we require self-stabilization, the initial states of the nodes, including the values of their local clocks, are arbitrary.

Pulse synchronization now requires nodes to regularly trigger *pulses* in a synchronized way. For a time $T \geq 0$, denote by $t_i^{(k)}$ and $i \in \{1, \dots, n\}$, $k \in \mathbb{N}$, the time when node i generates its k^{th} pulse at or after time T (we omit T from the notation). A pulse synchronization algorithm of precision Δ , accuracy bounds $A_{\min}$, $A_{\max}$, and stabilization time S satisfies in any execution that there is some time $T \leq S$ so that

$$\text{precision: } \forall i, j, k : |t_i^{(k)} - t_j^{(k)}| \leq \Delta \text{ and}$$

$$\text{accuracy: } \forall i, k : A_{\min} \leq |t_i^{(k+1)} - t_i^{(k)}| \leq A_{\max}.$$

Here it is implicit that indices i, j refer to correct nodes only, as we cannot make any guarantees on Byzantine nodes' behavior. Note that typically $A_{\min}$ will be a given parameter and the goal is to minimize $A_{\max} - A_{\min}$ as a function of $A_{\min}$ (or vice versa). Due to the drifting local clocks and delayed messages, indistinguishability arguments show that always $\Delta \geq 1$ and $A_{\max} \geq \max\{\vartheta A_{\min}, 1\}$.

3.1.1 Approaches by the Distributed Community

The results from [36] prompted the question whether pulse synchronization could also be solved efficiently, i.e., with a small stabilization time. In a series of papers, the stabilization time was first reduced to polynomial [20] and then linear [18, 29] in n .³ These works also revealed that randomization is not essential to solve the problem: the latter algorithm is based on running multiple instances of deterministic consensus concurrently.

Together, these results indicated that the problem could admit solutions suitable for hardware implementation. However, none of the above algorithms was a

²This is a normalization. In all existing algorithms, the maximum delay affects stabilization times, etc. linearly.

³Linear stabilization time was claimed earlier [19], but the algorithm contained non-trivial errors that were fixed in [18].

good candidate, due to unacceptable stabilization time [36], the need for highly accurate local clocks [20], or message size $\Theta(n \log n)$ and too involved local computations [18, 29]. Malekpour provides an alternative linear-time solution [68, 69] with small messages and simple computations, but uses a simplified model (in particular, it is assumed that $\vartheta = 1$, i.e., there is no clock drift).

3.1.2 Approaches by the Hardware Community

Frequently, designers of fault-tolerant architectures consider the clocking mechanism sufficiently reliable and hence do not add any measures for fault-tolerance. The typical rationale is that a clock distribution network has very strong drivers and is therefore not susceptible to transient disturbances. Permanent defects, on the other hand, will make the system stop operation completely, which may be considered safe in some cases. In reality, however, the clock distribution infrastructure is already so complicated that a “partial” defect can easily occur (imagine a driver responsible for a sub-net failing). Moreover, considering that the clock tree is virtually always the most widespread network, showing the highest activity (in terms of transition frequency), it is not so clear why it should not be affected from transient faults as well. These arguments become even more dominant when talking about requirements of, e.g., a failure probability smaller than 10^{-9} per hour. For such a degree of reliability, it is unrealistic to assume that the system can be just “tried out” before being used, and the cost associated with a design error can be extremely high.

As a single clock source like a crystal oscillator constitutes a single point of failure, practitioners aiming for fault-tolerant clocking often turn to the alternative of using multiple clock sources. While this approach is indeed capable of solving the fault-tolerance issue, it at the same time introduces a new problem, namely that of synchronization. In the single-clock case we have a single timing domain to which all activities are synchronized.⁴ Within this synchronous domain it is easy to perform communication based on known time bounds, to establish a clear ordering/precedence of events, and to avoid metastability caused by setup/hold time violations (i.e., too short input pulses) at storage elements. When using multiple clock sources, we immediately leave this safe area. It does not matter whether we use multiple clocks with the same nominal frequency or not – the only important distinction is whether the clocks are correlated (i.e., originate at the same source) or uncorrelated. In the latter case, one can never reason about their relative phase (which is essential for avoiding setup/hold time violations), which makes it mandatory to use explicit synchronizers that can, beyond causing performance

⁴The difficulty of providing this time information with the required phase accuracy all over a large system is, besides the fault-tolerance aspect, a key reason why this globally synchronous design paradigm is being challenged.

and area overheads, never attain complete protection from metastable upsets in the general case [60, 71].

With respect to the precision of existing approaches to synchronization, distinguishing “microticks” and “macroticks” has become common. Ultimately, this boils down to dealing with a large value of Δ by dividing the clock frequency (which is between $1/A_{\max}$ and $1/A_{\min}$ in our model) with respect to communication, so that $\Delta \ll 1/A_{\min}$. In other words, slowing down communication sufficiently, one can make the system work despite large Δ . However, this is obviously detrimental to performance, and one hence must strive for minimizing Δ . The software-based clock synchronization mechanisms found in practical applications like the Time-Triggered Protocol TTP [61] or FlexRay [42, 45] rely on adjusting local microtick counters appropriately to attain synchrony on macrotick level. However, both protocols are, essentially, variants of synchronous approximate agreement [32]. Hence, they require pre-established synchrony for correct operation, implying that they are not self-stabilizing.

Modern application-specific integrated circuits (ASICs) are typically composed of many internally synchronous function blocks that “solve” the issue of synchronization in an even simpler way. Instead of relying on any kind of synchronization between different clock domains, these blocks communicate without making any assumptions on timing (one needs still to avoid buffer overflows, however). This paradigm is called *globally asynchronous locally synchronous* (GALS) in the literature [14]. The intention here is mainly to mitigate the clock distribution problem, but this also provides a foundation for establishing fault-tolerance. Due to the consequent existence of multiple clock domains, such architectures feature copious amounts of synchronizers (to avoid metastability) and arbiters (to establish precedence). This works in practice, but comes with the associated performance penalties. Moreover, due to the current lack of tractable models accounting for metastability, there is no satisfying answer to the question “how many synchronizers suffice?”

What is needed to get rid of the performance and area overheads incurred by the GALS approach is correlated clocking all over the system, even if the phase is not perfectly matched in all places. Such clocks are called *mesosynchronous*. Probably the most natural implementation of such a distributed correlated clock source is a ring oscillator. The underlying principle is to use the delay along a cyclic path (gates plus interconnect) as a time reference. More specifically, such a path is implemented in an inverting fashion (odd number of inverting elements), such that the level is periodically inverted with the loop delay defining the half period of the oscillation. Examples of such approaches are the circuits presented by Maza et al, [77] and Fairbanks et al. [38]. They are ring oscillators in that they both exploit the fact that a circle formed by an odd number of inverters will oscillate, and the frequency of the produced clock is determined by circuit delays. In

contrast to the simple basic ring oscillator scheme, multiple “rings” are connected to form meshes of inverters that distribute “clock waves”, thereby also generating new ones. In forming the mesh, output signals of inverters are joined by simply hardwiring them and forked by wire forks.

While these approaches can indeed provide a fast clock that is perceived as “correlated” all over the system, the clock is still not intended and claimed to be fault-tolerant by the authors.

3.1.3 DARTS

One may view the above hardware realizations of distributed clock generators as very simple distributed algorithms, in which the algorithmic steps are determined by the laws of superposition at the merging points. From a theoretical point of view, this results in extremely limited options for the designer of the algorithm. Thus, it is quite tempting to try out more sophisticated algorithms and prove strong(er) fault-tolerance properties. To this end, a suitable algorithm must be chosen and the hardwiring be replaced by an implementation based on logic gates.

This idea was at the heart of the DARTS project [50]. The overall goal of the project was to implement the fault-tolerant synchronization algorithm by Srikanth and Toueg [90] in hardware. The pseudo-code of the algorithm, given in Algorithm 1, is executed at each node of a fully connected network of $n > 3f$ nodes, where f is the number of Byzantine faulty nodes it can tolerate. The code is extremely simple, yet one should not be fooled: its implementation in hardware required to overcome non-trivial obstacles [44].

Algorithm 1 Pseudo-code of a node to synchronously generate round(k) messages.

Upon bootup

- 1: $k \leftarrow 0$;
- 2: broadcast round(0);

Upon reception of a message

- 3: **if** received round(ℓ) from at least $f + 1$ distinct nodes with $\ell > k$ **then**
 - 4: broadcast round($k + 1$), $\dots$, round(ℓ);
 - 5: $k \leftarrow \ell$;
 - 6: **end if**
 - 7: **if** received round(k) from at least $2f + 1$ distinct nodes **then**
 - 8: broadcast round($k + 1$);
 - 9: $k \leftarrow k + 1$;
 - 10: **end if**
-

According to the algorithm’s assumptions, a fully connected communication structure was established, which also provides the highest possible degree of fault-tolerance. The implementation of the merging points, i.e., the actual algorithm, was done in clockless logic. This avoids the issue of having to synchronize the local clock source to the correlated “global” time provided by the algorithm (otherwise one would have to rely on synchronizers again), but in turn requires a careful algorithmic design and timing analysis of the entire system [49]. Interestingly, this means that the only timing sources in DARTS are lower and upper bounds on wire delays, without any formal local clocks. Thus, it is quite close in spirit to the solutions inspired by ring oscillators discussed previously. The hardware implementation of a DARTS node is shown in Figure 2.

The implementation of these hardware nodes, which were called “tick generation algorithm (TG-Alg) nodes,” was a very insightful example for how difficult it is to map algorithms that were, at best, developed with a software implementation in mind, to hardware. Assumptions that seem simple at the level of an algorithm may turn out extremely painful when having to be realized in hardware. Examples here are the existence of unbounded counters (such as “k” in Algorithm 1), the request for mutual exclusive execution of tasks, the generous use of operators (multiplication is expensive to implement in hardware), etc.

The identification of relevant timing constraints was a challenging issue in the design of the DARTS prototype ASIC as well. Recall that metastability can, in principle, not be avoided (in the general case) for uncorrelated clock sources. However, one can show that in fault-free executions, metastability does not occur. This is not straight-forward due to the following cyclic dependencies: Under the assumption of proper function of the algorithm one can rely on its guaranteed properties (e.g. precision) when establishing the timing closure to avoid setup/hold time violations. In turn, the freedom from setup/hold time violations is a prerequisite for correct functionality.⁵ Note that such timing guarantees are essential, as metastability, a possible result of setup/hold violations, results in unspecified behavior not covered by the analysis of the algorithm.

Several key techniques were applied for overcoming the above challenges:

- the use of difference counters for the cycle number, thus mitigating the problem of unlimited counting;
- the implementation of these counters through Muller pipelines, thus avoiding metastability issues that would arise from concurrent increase and decrease operations of the same counter;

⁵For DARTS, “only” an inductive argument was required. When turning to self-stabilization, establishing that the most basic timing assumptions eventually hold tends to be the most difficult aspect of the reasoning.

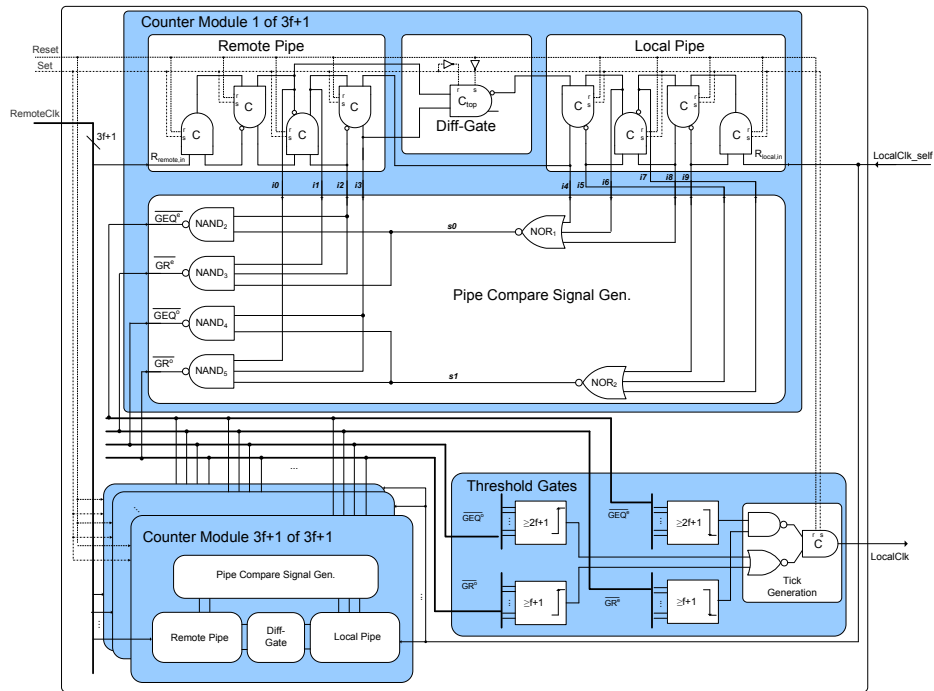


Figure 2: Hardware implementation of a DARTS node in clockless logic.

- a careful mix of event-based and state-based logic;
- the separated treatment of rising and falling signal edges in order to substantially relax the timing constraints.

The project succeeded in developing a working prototype chip, demonstrating that it was indeed feasible to use a distributed algorithm for generating a system-wide clock *and* prove that its *implementation* provides respective guarantees:

- bounded precision and accuracy,
- tolerance of Byzantine faults, and
- use of standard hardware libraries, with one exception: a C-Element must be added.

While the third property might seem like an oddball here, one should be aware that novel circuit components need to be designed on transistor level, layouted, characterized and validated (by simulation) as well. The existing standard libraries had to be augmented by a C-Element for DARTS.

While DARTS constituted a breakthrough in terms of bringing theory and practice closer to each other, the resulting solution exhibits a number of deficiencies calling for further research:

- full connectivity between nodes;
- lack of recovery from transient faults: even if only a minority of nodes undergoes transient faults at any time, there is no mechanism to recover to a correct state;
- too small frequency, limited by the propagation delay of a long loop;
- non-trivial initialization due to fairly strict demands on initial synchrony.

3.1.4 FATAL

In light of the theoretical results from Section 3.1.1 and the proof-of-concept that Byzantine fault-tolerance is viable in low-level hardware provided by DARTS, the obvious next question is whether self-stabilization can be added on the circuit level, too. This was answered affirmatively in [26].

The FATAL algorithm builds on the idea of adding a recovery mechanism to a pulse generation mechanism based on threshold voting. On an abstract level, the FATAL algorithm can be viewed as an implementation of the Srikanth-Toueg algorithm (c.f., Algorithm 1) that avoids having to keep book on tick/pulse numbers by making sure that the time between pulses is sufficiently large: instead of broadcasting round(k) messages, we simply broadcast round messages in Algorithm 1. Another interpretation is that of having a DARTS system that runs slow enough to operate with “pipes of length one”, i.e., simple memory cells.

The basic principle is illustrated in Figure 3, depicting a state machine each node runs a copy of. Circles represent states, arrows state transitions that happen when the condition next to the arrow is satisfied, and the box with “propose” on the state transition from *pulse* to *ready* indicates that the nodes’ memory flags are reset during this state transition. Each node continuously broadcasts whether it is in state *propose* or not, and when a node perceives another in this state according to this signal (including itself), its respective memory flag is set (i.e., each node has one memory cell for each node, including itself). The condition “ 3ϑ local time passed” means that node i will transition from *pulse* to *ready* at the time t when its local clock reads $C_i(t) = C_i(t') + 3\vartheta$, where t' is the time when it switched to *pulse*. Nodes generate a pulse when switching to *pulse*.

It is not hard to verify that, if all nodes start in *ready* with their memory flags cleared, this algorithm will solve pulse synchronization with $\Delta = 2$, $A_{\min} = 3 + 3\vartheta$, and $A_{\max} = 3 + 3\vartheta + 3\vartheta^2$. By making nodes wait longer in one of the transitions by

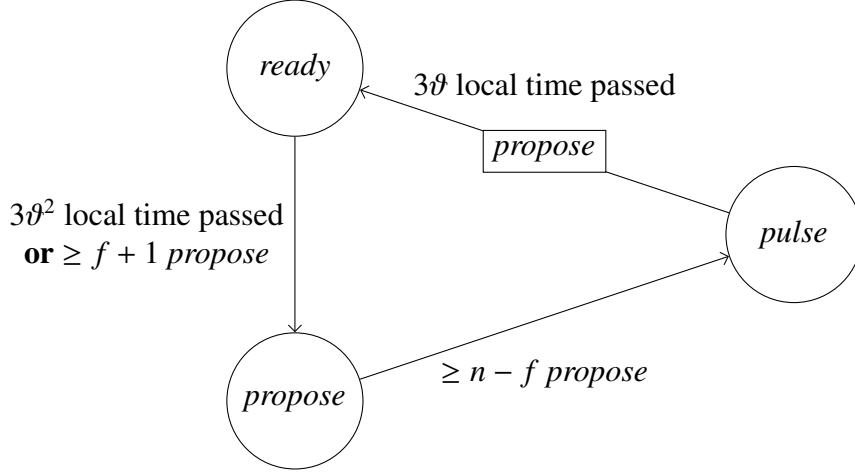


Figure 3: Simple pulse synchronization requiring consistent initialization.

ϑx local time, we can actually have $A_{\min} = 3 + 3\vartheta + x$ and $A_{\max} = 3 + 3\vartheta + 3\vartheta^2 + \vartheta x$, for any $x \geq 0$, i.e., $A_{\max} \rightarrow \vartheta A_{\min}$ for $x \rightarrow \infty$.

We believe that the recovery mechanism developed for FATAL is a potential key building block of further improvements in the future. In [26], the above basic algorithm is modified so that the task of “stabilizing” the routine, i.e., getting it into a (global) state as if it had been initialized correctly, is reduced to generating a single pulse *by an independent algorithm*. More precisely, all correct nodes need to trigger a “start stabilization” event within a reasonably small time window and then not trigger another such event for $\Theta(1)$ time in order to guarantee stabilization.

The easier task of generating a single “helper pulse” for the purpose of recovery from transient faults is then solved by relying on randomization. The solution used in [25] generates such a pulse with probability $1 - 2^{-\Omega(n)}$ within $O(n)$ time, resulting in an overall stabilization time of $\Theta(n)$. Hence, the algorithm matches the best known stabilization time of $O(n)$. The improvement lies in the communication complexity and the amount of local computations: Apart from a few memory flags for each other node, each node runs a state machine with a constant number of states; each node continuously broadcasts only a few bits about its own state. Moreover, the algorithm can operate with arbitrary values of ϑ , permitting to use very simple oscillators as local clock sources.

In [25], the approach is implemented and evaluated in hardware. The experiments confirm the theoretical results from [26]. However, the algorithm cannot be used for clocking as-is, for several reasons:

- The algorithm generates pulses every $\Theta(1)$ time, but the involved constants are impractically large. Naive application of the approach would result in

slowing down systems by several orders of magnitude.

- The pulses are anonymous, i.e., the counting problem discussed in the next section needs to be solved.
- The system is fully connected, which is infeasible in large circuits.

These issues will be discussed next.

3.2 Counting

Once pulse synchronization is solved, it can be used to simulate synchronous rounds: One adjusts the accuracy lower bound $A_{\min}$ such that it allows for the maximal sum of the communication delay between neighbors, the local computations for a round, and a safety margin proportional to Δ (recall that a pulse is not issued at all nodes precisely at the same instant of time). However, due to the strong fault model, it is non-trivial to achieve agreement on a round counter. Round counters are highly useful for, e.g., applying pre-determined time division multiple access (TDMA) schemes to shared resources (memory, communication network, etc.) or scheduling synchronized operations (measurements, snapshots, etc.) that are to be executed regularly.

We will now discuss how a self-stabilizing round counter can be constructed in a synchronous system with $f < n/3$ Byzantine faults. The problem of *C-counting*, where C is an integer greater than 2, is formalized as follows. In each round $r \in \mathbb{N}$, each node i outputs a counter $c_i(r) \in \{0, \dots, C - 1\}$. The algorithm stabilizes in $S \in \mathbb{N}$ rounds, if for all $r \geq S$ we have

agreement: $\forall i, j : c_i(r) = c_j(r)$ and

counting: $\forall i : c_i(r + 1) = c_i(r) + 1 \bmod C$.

In this subsection, the discussion will be more technical, with the goal of illustrating how the fault-tolerance techniques that have been developed by the distributed computing community are canonical tools for designing fault-tolerant algorithms in hardware. We remark that the inclined reader should feel free to skip the technical proofs, as they are not essential to the remaining discussion in this article.

3.2.1 Equivalence to Consensus

The task of (binary) *consensus* requires that, given an input $b_i \in \{0, 1\}$ at each node at the beginning of round 1, each correct node computes an output o_i satisfying

agreement: $\forall i : o_i = o$ for some $o \in \{0, 1\}$ (we refer to o as the output),

validity: if $\forall i : b_i = b$ then $o = b$, and

termination: all (correct) nodes eventually set their output (exactly once) and terminate.

In practice, one usually requires explicit bounds on when nodes terminate. By an R -round consensus algorithm, we will refer to an algorithm in which all correct nodes terminate by the end of round R .

The counting problem is equally hard as consensus with respect to asymptotic time complexity. We show this for deterministic algorithms and binary consensus algorithms here, but extensions to non-binary consensus and randomized algorithms are straightforward.

Lemma 3.1 (Counting solves consensus). *Given an algorithm for C -counting stabilizing in R rounds, binary consensus (with $f < n/3$ Byzantine nodes) can be solved in R rounds.*

Proof. Denote by $\mathbf{x}(0)$ and $\mathbf{x}(1)$ state vectors such that the counting algorithm would properly count starting from 0 and 1, respectively (regardless of the subset of faulty nodes). Such states must exist, because after stabilization the algorithm will count modulo C and Byzantine nodes may pretend correct operation to avoid detection until after stabilization. Given an input vector $\mathbf{b} \in \{0, 1\}^n$, initialize each (correct) node i with state $x_i(b_i)$ and run the algorithm for R rounds. Then each node outputs $c_i(R) - R \bmod 2$.

Clearly, this algorithm terminates in R rounds and, by the agreement property of the counting protocol, all nodes output the same value. Hence it remains to show that this output value is valid, i.e., equals b if $b_i = b$ for all correct nodes. This follows from the choice of $\mathbf{x}(0)$ and $\mathbf{x}(1)$ and the counting property, which implies that, for all correct nodes i , $c_i(R) = R \bmod C$ if $b = 0$ and $c_i(R) = R + 1 \bmod C$ if $b = 1$. $\square$

The other direction was shown in [30]. We present a simpler variant in the following lemma. It makes use of consensus for non-binary values, i.e., $b_i, o_i \in \{0, \dots, C-1\}$ (this case can be reduced to binary consensus in a straightforward manner).

Lemma 3.2 (Consensus solves counting). *Given a synchronous consensus algorithm for inputs $0, \dots, C-1$ terminating in R rounds that tolerates $f < n/3$ Byzantine nodes, C -counting can be solved with stabilization time $O(R)$.*

Proof. Given the consensus algorithm, we solve C -counting as follows. In each synchronous round, we start a new consensus instance that will generate an output value $c_i(r + R)$ at each node i exactly R rounds later (which will double as node

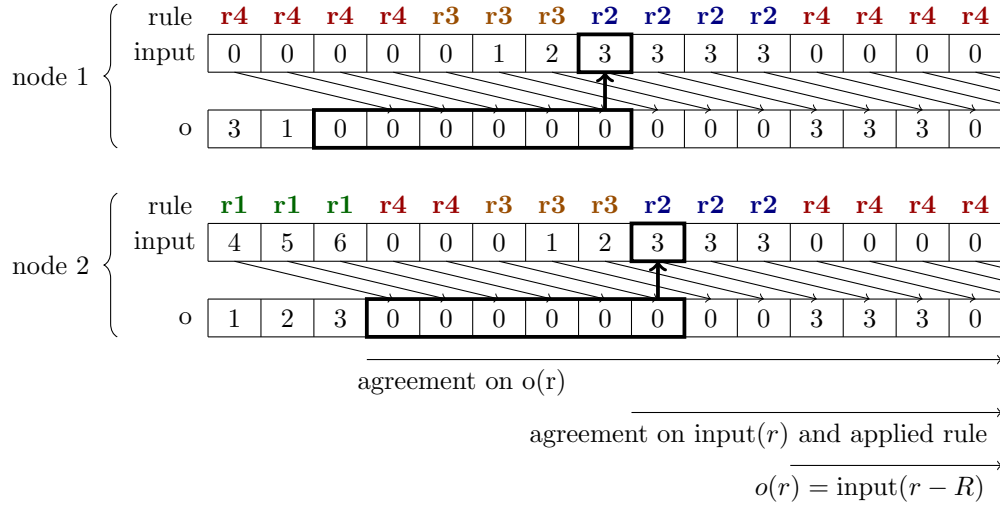


Figure 4: Part of an execution of two nodes running the C -counting algorithm given in the proof of Lemma 3.2, for $C = 8$ and $R = 3$. The execution progresses from left to right, each box representing a round. On top of the input field the applied rule (1 to 4) to compute the input is displayed. Displayed are the initial phases of stabilization: (i) after R rounds agreement on the output is guaranteed by consensus, (ii) then agreement on the input and the applied rule is reached, and (iii) another R rounds later the agreed upon outputs are the agreed upon inputs shifted by 3 rounds.

i 's counter value). Note that, while we have no guarantees about the outputs in the first R rounds (initial states are arbitrary), in all rounds $r \geq R$ all correct nodes will output the same value $o(r) = o_i(r)$ (by the agreement property of consensus). Hence, if we define the input value $F_i(r)$ of node i as a function of the most recent $O(R)$ output values at node i , after $O(R)$ rounds all nodes will start using identical inputs $F(r) = F_i(r)$ and, by validity of the consensus algorithm, reproduce these inputs as output R rounds later (cf. Figure 4). In light of these considerations, it is sufficient to determine an input function F from the previous $O(R)$ outputs to values $0, \dots, C - 1$ so that counting starts within $O(R)$ rounds, assuming that the output of the consensus algorithm in round $r + R$ equals the input determined at the end of round r .

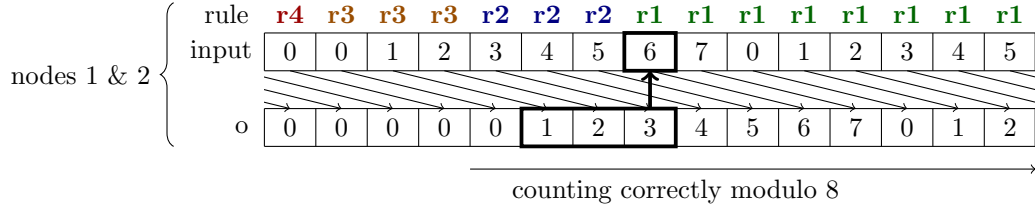


Figure 5: Extension of the execution shown in Figure 4. Nodes have already agreed upon inputs and outputs so that the latter just reproduce the inputs from R rounds ago. The rules now make sure that the nodes start counting modulo 8 in synchrony, always executing rule 1.

We define the following input function, where all values are taken modulo C :

$$\text{input}(r) := \begin{cases} c + R & \text{if } (o(r - R + 1), \dots, o(r)) = (c - R + 1, \dots, c) \\ x + R & \text{if } (o(r - 2R + 1 - x), \dots, o(r)) = (0, \dots, 0, 1, \dots, x) \\ & \text{for some } x \in \{0, \dots, R - 1\} \\ x & \text{if } (o(r - R + 1 - x), \dots, o(r)) = (0, \dots, 0) \\ & \text{for maximal } x \in \{0, \dots, R - 1\} \\ 0 & \text{else .} \end{cases}$$

In the setting discussed above, it is straightforward to verify the following properties of input:

- Always exactly one of the rules applies, i.e., input is well-defined.
- If the outputs counted modulo C for $2R$ consecutive rounds, they will do so forever (by induction, using the first rule); c.f. Figure 5.
- If this does not happen within $O(R)$ rounds, there will be R consecutive rounds where input 0 will be used (by the third and the last rule), c.f. Figure 5.
- Once R consecutive rounds with input 0 occurred, inputs $1, \dots, 2R$ will be used in the following $2R$ rounds (by the second and third rule).
- Finally, the algorithm will commence counting correctly (by the first rule).

Overall, if each node i computes its input $F_i(r)$ from its local view of the previous outputs using input, the algorithm will start counting correctly within $S \in O(R)$ rounds. $\square$

These two lemmas imply that there is no asymptotic difference in the round complexity of consensus and the stabilization time of counting. However, note

that the conversion of a consensus algorithm into a counting algorithm given by Lemma 3.2 is very “lossy” in terms of communication complexity and computational efficiency, as R instances of consensus run concurrently! Hence, the main question is whether there are fast solutions to counting that are efficient in terms of communication and computation as well.

3.2.2 Counting Using Shared Coins

The pulse synchronization algorithm by S. Dolev and Welch [36] is conceptually based on a counting algorithm given in the same article, yielding an exponential time randomized solution.

This was improved by Ben-Or et al. [9]. We explain a simpler variant of the algorithm here. The solution is based on *shared coins*. A (weak) shared coin guarantees that there are probabilities $p_0, p_1 > 0$ so that with at least probability p_0 , all correct nodes output 0, and with at least probability p_1 , all correct nodes output 1. We call $p := \min\{p_0, p_1\}$ the *defiance* of the shared coin. Moreover, we require that the value of the coin is not revealed before the round in which the outputs are generated, so that faulty nodes cannot exploit this knowledge to prevent stabilization.

Observe that we neither require $p_0 + p_1 = 1$ nor that all correct nodes always output the same value. In particular, a trivial shared coin with defiance 2^{-n} is given by each node flipping a coin independently. The algorithm from [36] essentially makes use of this trivial shared coin, which results in its expected exponential stabilization time.

The first step of the algorithm from [9] is to solve 2-counting.

Lemma 3.3 (2-counting from shared coin). *Given a stream of weak shared coins with constant defiance, 2-counting can be solved with constant expected stabilization time.*

Proof. In each round r , each node i

1. broadcasts $c_i(r)$;
2. if it received at least $n - f$ times value $c_i(r) - 1 \bmod 2$ in round $r - 1$, it sets $c_i(r + 1) := c_i(r) + 1 \bmod 2$; and
3. otherwise, $c_i(r + 1)$ is set to the output of the shared coin at i in round r .

Before we prove the claim, note that Step 2 depends on messages that were broadcasted in round $r - 1$ instead of messages broadcasted in step 1 during the same round r . The reason for this is to avoid that the faulty nodes exploit so-called *rushing*: As the value of the coin for round r is revealed in round r , faulty

nodes may exploit this information to affect the outcome of the broadcast (in terms of what correct nodes observes) exactly so that in Step 3 the “wrong” action is taken by correct nodes relying on the coin. By referring to the broadcast of the previous round instead, the faulty nodes are forced to commit to an outcome of the broadcast before the coin is revealed, making sure that with probability at least p the “right” action is taken by correct nodes in Step 3.

To see that the algorithm indeed stabilizes, observe first that it cannot happen that, in the same broadcast, a correct node sees value 0 at least $n - f$ times and another correct node sees value 1 at least $n - f$ times: this implies that at least $n - 2f$ correct nodes each have $c_i(r) = 0$ and $c_i(r) = 1$, respectively, but we have only $n - f < n - f + (n - 3f) = 2(n - 2f)$ correct nodes (here we used that $f < n/3$). Assume that $c \in \{0, 1\}$ is the unique value such that some node receives c at least $n - f$ times in round $r - 1$. If there is no such value, choose c arbitrarily. With probability at least p_c , all correct nodes set $c_i(r + 1) := c + 2 \bmod 2 = c$ in round r . Similarly, in round $r + 1$ all nodes set $c_i(r + 2)$ to $c + 1 \bmod 2$ with probability at least p_{1-c} . Once this happened, the clocks of correct nodes will start counting deterministically, as always Step 2 will be executed. Hence, the algorithm stabilizes with (independent) probability $p_0 p_1 \geq p^2$ every other round. $\square$

Once 2-counting is available, the generalization to C -counting is achieved by a similar approach. The key difference is that we now use a two-round protocol controlled by the output of the 2-counting algorithm to solve C -counting. We remark that the algorithm essentially performs a gradecast ([40]) followed by achieving a consistent choice with constant probability using the shared coin.

Lemma 3.4 (*C -counting from shared coin and 2-counting*). *Given a stream of weak shared coins with constant defiance and a 2-counter, C -counting can be solved with constant expected stabilization time.*

Proof. In each round r , each node i performs the following steps.

1. If the 2-counter reads 0:
 - (a) broadcast $c_i(r)$;
 - (b) if received value $c \neq \perp$ at least $n - f$ times, set helper variable $h_i(r) := c$, otherwise $h_i(r) := \perp$;
 - (c) if $b_i(r - 1) = 1$ or the shared coin shows 1 at i in round r , set $c_i(r + 1) := c_i(r) + 1 \bmod C$, and otherwise $c_i(r + 1) := 0$.
2. If the 2-counter reads 1:
 - (a) broadcast $h_i(r - 1)$;

- (b) if received value $c \neq \perp$ at least $n - f$ times, set $c_i(r + 1) = c + 1 \bmod C$ and $b_i := 1$;
- (c) else if received value $c \neq \perp$ at least $n - 2f$ times, set $c_i(r + 1) = c + 1 \bmod C$ and $b_i(r) := 0$;
- (d) else set $c_i(r + 1) := 1$ and $b_i(r) := 0$.

To see why this stabilizes with constant probability within $O(1)$ rounds, observe that the following holds once the 2-counter stabilized:

- If in an even round r all correct nodes agree on the clock value and have $b_i(r - 1) = 1$, the algorithm will count correctly forever.
- The same holds if they agree on the clock and the shared coin shows 1 at all nodes in round r .
- As $f < n/3$, there can be at most one value $c \neq \perp$ with correct nodes setting $h_i(r) := c$ in an even round r .
- If any correct node i receives this unique value c at least $n - f$ times in the subsequent odd round $r + 1$, all correct nodes receive c at least $n - 2f$ times. Hence, *either* it holds that (b) or (c) applies at all correct nodes *or* (c) or (d) apply at all nodes.
- In the first case, all correct nodes have the same clock value. Hence, the shared coin showing 1 in round $r + 2$ guarantees stabilization.
- In the second case, all correct nodes set $b_i(r + 1) := 0$. Hence, if the coin shows 0 at all nodes in round $r + 2$, they all set $c_i(r + 3) := 0$ and, subsequently $c_i(r + 4) := 1$. If the coin shows 1 at all nodes in round $r + 4$, this implies stabilization.

Hence, the algorithm stabilizes with (independent) probability $\min\{p_1, p_0 p_1\} \geq p^2$ within every 4 rounds once the 2-counter counts correctly. $\square$

Composing the two algorithms yields a C -counter with expected constant stabilization time. We stress the similarity of the routine to solutions to consensus based on shared coins [88]; the ideas and concepts developed for consensus translate directly to the counting problem, even if it might be harder in terms of the required communication.

Unfortunately, this approach to solving counting suffers from the same problem as consensus algorithms based on shared coins: theoretically sound protocols that generate shared coins based on the private randomness of the nodes are highly expensive in terms of communication and computation.

3.2.3 Constructing Large Counters from Small Counters

There are several techniques for constructing large counters from small counters, indicating that the key challenge is to obtain a 2-counter. One is given by Lemma 3.4, which however necessitates the presence of a shared coin. Another one is very basic, but inefficient time-wise, as C enters the stabilization time as a factor.

Lemma 3.5 (*C-counting from 2-counting [31]*). *Given a 2-counting algorithm with stabilization time S , for any $k \in \mathbb{N}$ we can solve 2^k -counting with stabilization time $2^k S$ and at most factor 2 more communication.*

Proof. The proof goes by induction over k , the base case being covered by assumption. For the step, we simply execute the 2-counting algorithm slower, by performing one round when the 2^k -counter switches to 0. This way, concatenating the clock bit of the slow 2-counter and the 2^k -counter, we obtain a 2^{k+1} -counter. The stabilization time is $2^k S$ for the slowed-down 2-counter plus the stabilization time of the 2^k -counter, yielding by induction a total stabilization time of $\sum_{l=1}^k 2^l S < 2^{k+1} S$. The communication bounds of the 2-counting algorithm together with the slow-down yield the claim concerning the amount of communication.⁶ $\square$

A simple variant on the theme achieves faster stabilization at the cost of increased message size.

Lemma 3.6 (*C-counting from 2-counting, faster stabilization*). *Assuming we are given a 2-counting algorithm with stabilization time S , for any $k \in \mathbb{N}$ we can solve 2^k -counting with stabilization time $2^k + kS$ and at most factor k more communication.*

Proof. The proof goes by induction over k , the base case being covered by assumption. For the step, we execute another copy of the 2-counting algorithm with a minor change: If the already constructed 2^k -counter reads 0, we skip a round of the 2-counting algorithm. Thus, the 2-counter will proceed by $2^k - 1 \bmod 2 = 1$ every 2^k rounds. The 2^{k+1} -counter is now given by the 2^k -counter and an additional leading bit, which is the value the 2-counter had when the 2^k -counter most recently was 0. By the above considerations, the 2^{k+1} -counter will count correctly once (i) both counters stabilized and (ii) the 2^k -counter had value 0 once after this happened.

⁶Note that one can also ensure that the maximum message size does not increase by more than factor 2, by shifting the communication pattern so that no more than 2 instances of the 2-counting algorithm communicate in the same round.

The stabilization time bound now follows: once the 2^k -counter is correctly operating, the 2-counter stabilizes within $S + 1$ rounds, and the 2^k -counter will become 0 again within another 2^k rounds; summation yields $\sum_{l=0}^k 2^l + S < 2^{k+1} + (k + 1)S$ rounds for stabilization of the constructed 2^{k+1} -counter. The communication bound is trivially satisfied. $\square$

Even if 2-counting can be solved efficiently, these techniques are slow if C is large. Motivated by this issue, in [46] large clocks are constructed from small ones by encoding clock values over multiple rounds. This enables to increase the clock range exponentially. Specifically, the paper provides two main results. The first is essentially a reduction to consensus (with only one instance running at any given time), and it is similar to the approach taken in Lemma 3.4. The key changes are the following:

1. To enable 1-bit messages, broadcasts of clock values are replaced by $\lceil \log C \rceil$ rounds each in which the clock bits are transmitted sequentially.
2. Instead of relying on shared coins, nodes run an instance of consensus with the variables b_i determined in odd “rounds” as input. The output of the consensus algorithm is used by all nodes to decide whether c (shifted by the number of rounds passed) is the next clock value or 0.
3. In all other rounds, clock values are locally increased by 1 modulo C .

Due to the use of consensus, the correctness argument becomes simpler. If the consensus algorithm outputs 1, there must be a node that used input 1 and therefore received $n - f$ times c in the second broadcast. This implies that all nodes received $n - 2f \geq f + 1$ times c and therefore can determine the new clock value. Otherwise, the algorithm is certain to default to resetting clocks to 0.

This approach replaces the need for a shared coin with the need for an efficient consensus algorithm and a sufficiently large counter. We instantiate the result for the phase king protocol [10] in the following corollary.

Corollary 3.7 (Large counters from small counters and consensus [46]). *Given a C -counter for $C \in \mathcal{O}(n)$ sufficiently large, a $2^{\Omega(C)}$ -counter with stabilization time $\mathcal{O}(n)$ can be constructed deterministically using 1-bit broadcast messages.*

Note that one can combine this corollary with Lemma 3.5 or Lemma 3.6 to construct large counters from 2-counters. In [46], a randomized alternative to these lemmas is given that constructs larger counters from an $\mathcal{O}(1)$ -counter at smaller overhead. Using either of the two lemmas to obtain the required $\mathcal{O}(1)$ -counter, the following corollary can be derived.

Corollary 3.8 (Large counters from 2-counters using randomization [46]). *Given a 2-counter, C -counting can be solved with expected stabilization time $O(n + \log C)$ and $O(\log^* C)$ broadcasted bits per node and round.*

3.2.4 Counting from Pulse Synchronization

Ironically, the obstacle of solving 2-counting disappears if it is feasible to remove one level of abstraction and exert some control over how (perfect) synchrony is simulated. More concretely, in all existing pulse synchronization algorithms one can increase $A_{\min}$ (the minimum time between consecutive pulses) at will, so that $A_{\max}$ grows proportionally. In particular, this can be used to allow for more than a single synchronous round to be simulated in between pulses. Initializing the simple (non-self-stabilizing) pulse synchronization algorithm given in Figure 3 consistently, we thus can allow for sufficient time to generate a tunable number C of “sub-pulses” before the next pulse occurs. Counting locally modulo C by re-initializing the counter to 0 at each pulse and increasing it by 1 on each sub-pulse, we can use the sub-pulses as round delimiters for simulating synchronous rounds with a self-stabilizing C -counter that comes “for free”.

To be precise, this counter does not come entirely for free; apart from the additional logic, increasing $A_{\min}$ may also result in increasing the stabilization time of the pulse synchronization algorithm. However, one can obtain a 2-counter or, in fact, any $O(1)$ -counter, without asymptotically affecting the stabilization time of the underlying pulse synchronization algorithm. The techniques for constructing larger counters based on small counters given in [46] then can take it from there.

From an abstract perspective, this can be seen as an implementation of the Srikanth-Toueg algorithm [90] that counts only up to $O(1)$ and then is restarted. This approach is followed by FATAL⁺, an extended version of FATAL analyzed in [26] and implemented and evaluated in [25]. Owing to the simplicity of the algorithm given in Figure 3, the sub-pulses can actually be produced at a higher frequency and with better precision than offered by the basic FATAL algorithm.

We remark that the method of coupling the two algorithms in FATAL⁺ may be of independent interest. We expect that it can also be applied to couple FATAL or FATAL⁺ to non-stabilizing pulse synchronization protocols based on approximate agreement, like the earlier discussed TTP and FlexRay protocols. This bears the promise of obtaining a pulse synchronization protocol that (i) can run at even higher frequencies (i.e., $A_{\min}$ is smaller) and (ii) achieves a precision in the order of the *uncertainty* of the communication delay, i.e., if messages are underway between $1 - \varepsilon$ and 1 time unit, then $\Delta \in O(\varepsilon)$. This is to be contrasted to algorithms like DARTS or FATAL, which use explicit voting for resynchronization at each pulse and therefore have $\Delta \geq 1$ even if there is a lower bound on the communication delay. Note that the uncertainty of the communication delay is known to

be a lower bound on the worst-case precision of any clock synchronization protocol [67], implying that $\Delta \in \Omega(\varepsilon)$ is unavoidable.

3.2.5 Constructing Counters from Scratch

Modifying the system on such a deep level as how the clock signal is provided may not always be possible, e.g., when one designs a submodule in a larger circuit. In this case, one may still have to solve 2-counting directly.

Recent research has started to tackle this issue. In [31], efficient solutions for the special case of $f = 1$ are designed and proved optimal in terms of the trade-off between stabilization time and number of local states using computer-aided design. However, as the space of algorithms to consider is huge, this method does not scale; even the case $f = 2$ is currently beyond reach.

In [66], a recursive approach is taken to avoid that each node participates in $\Theta(R)$ concurrent instances of an R -round consensus algorithm used for establishing agreement on clock values. The target is to, in each step of the recursion, boost the *resilience* of the protocol to faults, while increasing neither the number of required nodes nor the time for stabilization too much. The result is an algorithm of slightly suboptimal resilience $f < n^{1-o(1)}$, but linear stabilization time $O(f)$ and only $O(\log^2 n / \log \log n + \log C)$ state bits per node. These state bits are broadcasted to all other nodes in each round. For deterministic algorithms, this implies an exponential improvement in communication complexity as compared to the solution from Lemma 3.2, since deterministic consensus algorithms satisfy that $R > f$ (see [41]).

To sketch the idea of the approach, consider a system of n nodes. Each node i runs a 0-resilient C_i -counter (for some C_i we will determine shortly). This is nothing but a local counter modulo C_i : it is increased in each round, and it works correctly so long as i does not fail. We use these counters to let nodes determine temporary leaders that will assist with stabilization if required; once the system is stabilized, the corresponding communication will be ignored. The current leader's local counter is used to provide a temporarily working counter to all nodes. This counter is used to run an $O(f)$ -round consensus algorithm, the phase king protocol [10], to agree on the counter values. It is straightforward to show that agreement cannot be destroyed by this mechanism once it is achieved, even if the temporary counter produces garbage later on.

In short, this mechanism reduces the task to ensuring that eventually a correct leader will emerge and persist for $R \in O(f)$ rounds. We achieve this as follows: Node 1 will cycle through all possible leaders, where it keeps “pointing” to the same node for $\Theta(R)$ consecutive rounds. Node 2 does the same, albeit slower by a factor of $2n$. This guarantees that, for any other node j , nodes 1 and 2 eventually consider it the leader for R consecutive rounds. We proceed inductively, slowing

down the “leader pointer” of node j by a factor of $(2n)^{j-1}$ compared to the one of node 1. Clearly, eventually all correct nodes will point to a correct node for R consecutive rounds.

The downside of this approach is that the stabilization time is exponential, since the slowest pointer takes $R \cdot (2n)^n$ rounds to complete a single cycle. Here the recursion comes into play. Instead of using single nodes, on each level of the recursion one groups the nodes into a small number $k \in O(1)$ of clusters. Each cluster runs an f -resilient counter that is used to determine to which leader the node currently points. The “leaders” now are also clusters, meaning that the slowest clock takes $R \cdot (2k)^k \in O(R)$ rounds for a cycle. Now the same principle can be applied, assuming that we can also convince correct nodes in blocks with more than f faults to point to the “correct” leader the blocks with at most f faults will eventually choose. Requiring that fewer than half of the k blocks have more than f faults, this is ensured by an additional majority vote. The resilience of the compound algorithm therefore becomes $\lceil k/2 \rceil (f + 1) - 1$ (one fault less than required to make half of the blocks contain $f + 1$ faults). Crunching numbers and tuning parameters, one obtains the claimed result.

Maybe the most interesting aspect of this construction is its resemblance to recursive approaches for improving the communication complexity of consensus protocols [10, 15, 57]. The additional challenge here is the lack of a common clock, which is overcome by relying on the guarantees from the lower levels together with the simple leader election mechanism explained above. From this point of view, the construction can be interpreted as a natural generalization of the recursive phase king algorithm given in [10]. Accordingly, one may hope that also for counting, it is possible to achieve optimal resilience in a similar way.

3.3 Clock Distribution

All the algorithms so far assume full connectivity, which is unrealistic if the number of nodes is not quite small. In principle, one could seek for solutions to the pulse synchronization and counting problems in low-degree topologies directly. However, it is much easier to solve these tasks in a small, fully connected “core” and then distribute the results redundantly using a sparse interconnection topology. The advantage is that for the distribution problem, it now is sufficient to have pre-defined master/slave relations, i.e., a pre-defined direction of propagation of the clock signal throughout the system. This greatly simplifies the job, as it circumvents the need for reaching any sort of agreement: the clock information is plainly dictated by the nodes further upstream.

When using a sparse network, we must also decrease our expectations in terms of fault-tolerance. Solving clock synchronization (or consensus, for that matter) in the presence of f faults requires minimum node degrees of $2f + 1$, as otherwise

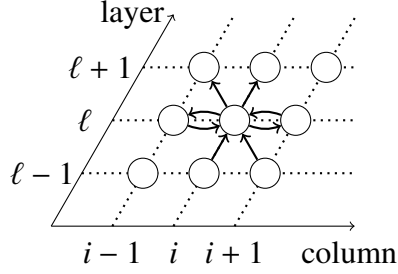


Figure 6: Node i in layer ℓ of a HEX grid and its incident links. The node propagates a pulse when having received it from both nodes on the previous layer. If one fails, the second signal is provided by one of its neighbors in the same layer.

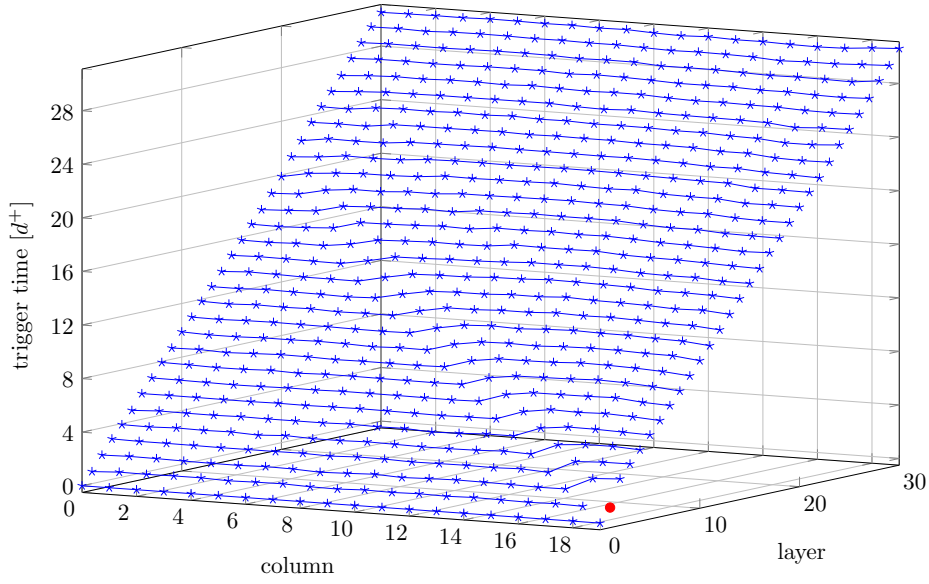


Figure 7: Pulse propagation in HEX with a single Byzantine node. The figure shows pulse trigger times of nodes in a grid: initially nodes (0 to 19) in layer 0 generate pulses in synchrony, feeding these pulses into the grid. The Byzantine faulty node 19 in layer 1 generates a “ripple” in trigger times that is clearly visible in the next layer, but smooths out over the following layers.

it may happen that a correct node does not have a majority of correct neighbors, rendering it impossible to falsify a claim jointly made by its faulty neighbors [23]. Respecting this, we require only that, for a given parameter f , the system tolerates up to f faulty nodes in the neighborhood of correct nodes.

Following these two paradigms – local fault-tolerance and directed clock propagation – and requiring a “nice” interconnect topology (planarity, connections to physically close nodes) led to HEX [24]. In a HEX grid, each node has 6 neighbors arranged in a hexagon, and the clock information spreads along the layers of the grid, cf. Figure 6. Nodes simply wait for the second signal indicating the current clock pulse, where nodes in the same layer send redundant signals in case one of the nodes in the preceding layer fails. Accordingly, a HEX grid tolerates $f = 1$ local fault, while having small node degrees and a planar topology.⁷

In order to prove precision bounds for HEX, we assumed that communication delays vary by at most $\varepsilon \ll 1$. Clearly, this implies that the precision cannot deteriorate by more than ε per layer. Surprisingly, a much stronger bound of $1 + \Delta_0 + O(\varepsilon^2 W)$ can be shown on the precision of adjacent nodes, where Δ_0 reflects the precision of the core and W is the width of the HEX grid.

This is an example of a non-trivial *gradient property*, a concept introduced by Fan and Lynch [39]. Finding clock distribution mechanisms that are suitable for hardware realization, fault-tolerant, and provide non-trivial gradient properties is of great interest, as a strong gradient property enables adjacent chip components to communicate at smaller delays in a synchronous, clock-driven fashion. In particular, it directly affects the time required to simulate a synchronous round and hence the operational frequency of the entire system.

Unfortunately, the worst-case precision of HEX deteriorates by $\Theta(f)$ in the presence of f faults, cf. Figure 7. While the simulations show that this is most likely overly pessimistic [65], the adverse effects of even a single fault are problematic in comparison to the surprisingly good performance of the system in absence of faults. We hope that topologies that feature at least $2f + 1$ links to the preceding layer will offer much better precision in face of faults; the idea is illustrated in Figure 8.

Open problems. In Section 3.2.2 we discussed efficient approaches to construct 2-counters from shared coins. While generating shared coins assuming Byzantine failures is prohibitively costly in terms of communication, it is interesting whether there are efficient shared coin protocols under weaker failure assumptions that are realistic for the considered hardware setting.

Clearly, the search for clock distribution topologies that can be implemented

⁷Clearly, the principle can be generalized to larger values of f adding edges, but degrees increase and planarity is lost.

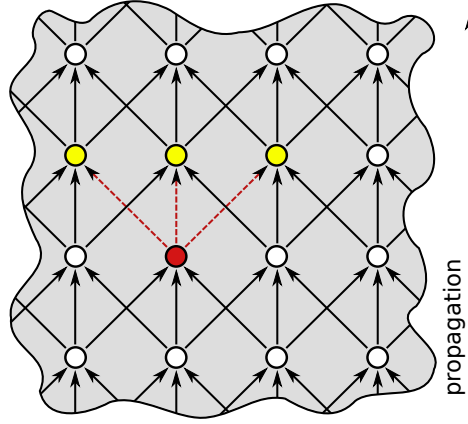


Figure 8: Local structure of a clock propagation approach similar too HEX. Using three connections from the previous layer helps to further mitigate the effect of faults on precision, as the redundant third clock signal does not propagate along a longer path.

with a small number of layers, balanced link delays and sufficiently high degree to tolerate more than 1 local node failure is of central interest. It is also not clear of how to adapt the pulse triggering rules in these HEX-variants to obtain optimal guaranteed precision bounds.

Improving the precision of fault-tolerant self-stabilizing approaches to clocking is an important challenge to achieve utility in practice. As mentioned earlier, it is promising to couple existing solutions with weak precision bounds to algorithms based on approximate agreement to combine high precision and self-stabilization.

Last but not least, an important question is how to verify the correctness of designs prior to production. Striving for algorithms that are sufficiently simple to be implemented in practice bears the promise of enabling formal verification of (parts of) the resulting systems. Given that suitable models can be devised, a grand challenge is the full verification of a fault-tolerant clocking mechanism bottom to top, from gates and wires up to the synchronous abstraction.

4 Conclusion

Due to the continuously increasing scale and complexity of today’s VLSI circuits, it becomes insufficient to ensure their reliability by fault mitigation techniques at technological and gate level only, as manufactures will not be able to support the combination of exponentially growing numbers of transistors and decreasing fea-

ture size indefinitely. Error correction, on the other hand, is restricted to storing information, neglecting the issue of dependable computation. Hence, one must strive for algorithmic fault-tolerance above the gate level, but below the abstraction of a synchronous, computationally powerful machine.

The distributed computing community has developed many concepts and algorithmic ideas that can be applied to VLSI circuits once we see them as what they truly are: distributed systems in their own right. The main challenges are

- to adapt and extend the existing theory beyond the abstraction of computationally powerful nodes;
- to devise models of computation that account for faults and metastability in a tractable manner;
- to come up with simple, yet efficient algorithms suitable for hardware implementation; and
- to reason about their correctness in ways ensuring that produced chips *will* work.

We believe that the examples given in this article demonstrate that the existing techniques are essential tools in tackling these challenges. The task that lies ahead is to fill the gap between fault-tolerance in theory and the design of practical, dependable hardware.

References

- [1] R. Alur, C. Courcoubetis, N. Halbwachs, T. A. Henzinger, P.-H. Ho, X. Nicollin, A. Olivero, J. Sifakis, and S. Yovine. The Algorithmic Analysis of Hybrid Systems. *Theoretical computer science*, 138(1):3–34, Feb. 1995.
- [2] R. Alur and D. L. Dill. A theory of timed automata. *Theoretical Computer Science*, 126(2):183–235, 1994.
- [3] R. Alur, T. Henzinger, G. Lafferriere, and G. Pappas. Discrete abstractions of hybrid systems. *Proceedings of the IEEE*, 88(7):971–984, July 2000.
- [4] J. H. Anderson and M. G. Gouda. A new explanation of the glitch phenomenon. *Acta Informatica*, 28(4):297–309, 1991.
- [5] P. J. Ashenden. *The designer's guide to VHDL*, volume 3. Morgan Kaufmann, 2010.
- [6] R. Baumann. Radiation-Induced Soft Errors in Advanced Semiconductor Technologies. *IEEE Transactions on Device and Materials Reliability*, 5(3):305–316, Sept. 2005.

- [7] S. Beer, R. Ginosar, J. Cox, T. Chaney, and D. M. Zar. Metastability challenges for 65nm and beyond; simulation and measurements. In *Design, Automation & Test in Europe Conference & Exhibition (DATE)*, pages 1297–1302. IEEE, 2013.
- [8] M. Bellido, J. Chico, and M. Valencia. *Logic-timing Simulation And the Degradation Delay Model*. Imperial College Press, 2006.
- [9] M. Ben-Or, D. Dolev, and E. N. Hoch. Fast Self-Stabilizing Byzantine Tolerant Digital Clock Synchronization. In *27th Symposium on Principles of Distributed Computing (PODC)*, pages 385–394, 2008.
- [10] P. Berman, J. A. Garay, and K. J. Perry. *Bit Optimal Distributed Consensus*, pages 313–321. Plenum Press, New York, NY, USA, 1992.
- [11] G. Brown, M. Gouda, and C. lin Wu. Token systems that self-stabilize. *IEEE Transactions on Computers*, 38(6):845–852, Jun 1989.
- [12] M. Broy and K. Stølen. *Specification and Development of Interactive Systems: Focus on Streams, Interfaces, and Refinement*. Springer-Verlag New York, Inc., 2001.
- [13] T. J. Chaney and C. E. Molnar. Anomalous behavior of synchronizer and arbiter circuits. *IEEE Transactions on Computers*, 22(4):421–422, 1973.
- [14] D. M. Chapiro. *Globally-Asynchronous Locally-Synchronous Systems*. PhD thesis, Stanford University, 1984.
- [15] B. A. Coan and J. L. Welch. Modular Construction of a Byzantine Agreement Protocol with Optimal Message Bit Complexity. *Information and Computation*, 97(1):61–85, 1992.
- [16] C. Constantinescu. Trends and Challenges in VLSI Circuit Reliability. *IEEE Micro*, 23(4):14–19, 2003.
- [17] J. Cortadella and M. Kishinevsky. Synchronous Elastic Circuits with Early Evaluation and Token Counterflow. In *44th Annual Design Automation Conference (DAC)*, pages 416–419, New York, NY, USA, 2007. ACM.
- [18] A. Daliot and D. Dolev. Self-Stabilizing Byzantine Pulse Synchronization. *Computing Research Repository*, abs/cs/0608092, 2006.
- [19] A. Daliot, D. Dolev, and H. Parnas. Linear Time Byzantine Self-Stabilizing Clock Synchronization. In *7th International Conference on Principles of Distributed Systems (OPODIS)*, volume 3144 of *LNCS*, pages 7–19. Springer Verlag, Dec 2003. A revised version appears in Cornell ArXiv: <http://arxiv.org/abs/cs.DC/0608096>.
- [20] A. Daliot, D. Dolev, and H. Parnas. Self-Stabilizing Pulse Synchronization Inspired by Biological Pacemaker Networks. In *6th Symposium on Self-Stabilizing Systems (SSS)*, pages 32–48, 2003.
- [21] L. de Alfaro, T. A. Henzinger, and M. Stoelinga. Timed Interfaces. In *Embedded Software (EMSOFT)*, pages 108–122, 2002.
- [22] A. Dixit and A. Wood. The Impact of New Technology on Soft Error Rates. In *IEEE Reliability Physics Symposium (IRPS)*, pages 5B.4.1–5B.4.7, Apr 2011.

- [23] D. Dolev. The Byzantine Generals Strike Again. *Journal of Algorithms*, 3:14–30, 1982.
- [24] D. Dolev, M. Függer, C. Lenzen, M. Perner, and U. Schmid. HEX: Scaling Honeycombs is Easier than Scaling Clock Trees. In *25th ACM Symposium on Parallelism in Algorithms and Architectures (SPAA)*, 2013.
- [25] D. Dolev, M. Függer, C. Lenzen, M. Posch, U. Schmid, and A. Steininger. Rigorously Modeling Self-Stabilizing Fault-Tolerant Circuits: An Ultra-Robust Clocking Scheme for Systems-on-Chip. *Journal of Computer and System Sciences*, 80(4):860–900, 2014.
- [26] D. Dolev, M. Függer, C. Lenzen, and U. Schmid. Fault-tolerant Algorithms for Tick-generation in Asynchronous Logic: Robust Pulse Generation. *Journal of the ACM*, 61(5):30:1–30:74, 2014.
- [27] D. Dolev, M. Függer, M. Posch, U. Schmid, A. Steininger, and C. Lenzen. Rigorously modeling self-stabilizing fault-tolerant circuits: An ultra-robust clocking scheme for systems-on-chip. *Journal of Computer and System Sciences*, 80(4):860–900, 2014.
- [28] D. Dolev, M. Függer, U. Schmid, and C. Lenzen. Fault-tolerant Algorithms for Tick-generation in Asynchronous Logic: Robust Pulse Generation. *Journal of the ACM*, 61(5):30:1–30:74, Sept. 2014.
- [29] D. Dolev and E. Hoch. Byzantine Self-Stabilizing Pulse in a Bounded-Delay Model. In *9th Symposium on Stabilization, Safety, and Security of Distributed Systems (SSS)*, volume 4280, pages 350–362, 2007.
- [30] D. Dolev and E. Hoch. On Self-stabilizing Synchronous Actions Despite Byzantine Attacks. In *21st Symposium on Distributed Computing (DISC)*, pages 193–207, 2007.
- [31] D. Dolev, J. H. Korhonen, C. Lenzen, J. Rybicki, and J. Suomela. Synchronous Counting and Computational Algorithm Design. In *15th Symposium on Stabilization, Safety, and Security of Distributed Systems (SSS)*, pages 237–250, 2013.
- [32] D. Dolev, N. A. Lynch, S. S. Pinter, E. W. Stark, and W. E. Weihl. Reaching Approximate Agreement in the Presence of Faults. *Journal of the ACM*, 33:499–516, 1986.
- [33] S. Dolev. *Self-Stabilization*. MIT Press, 2000.
- [34] S. Dolev and Y. Haviv. Self-stabilizing microprocessor: analyzing and overcoming soft errors. *IEEE Transactions on Computers*, 55(4):385–399, April 2006.
- [35] S. Dolev and J. L. Welch. Self-Stabilizing Clock Synchronization in the Presence of Byzantine Faults (Abstract). In *14th Symposium on Principles of Distributed Computing (PODC)*, page 256, 1995.
- [36] S. Dolev and J. L. Welch. Self-Stabilizing Clock Synchronization in the Presence of Byzantine Faults. *Journal of the ACM*, 51(5):780–799, 2004.

- [37] J. C. Ebergen. A formal approach to designing delay-insensitive circuits. *Distributed Computing*, 5(3):107–119, 1991.
- [38] S. Fairbanks and S. Moore. Self-Timed Circuitry for Global Clocking. In *11th International Symposium on Asynchronous Circuits and Systems (ASYNC)*, pages 86–96, 2005.
- [39] R. Fan and N. Lynch. Gradient Clock Synchronization. In *23rd ACM Symposium on Principles of Distributed Computing (PODC)*, pages 320–327, 2004.
- [40] P. Feldman and S. Micali. Optimal algorithms for Byzantine agreement. In *ACM Symposium on Theory of Computing*, pages 148–161, 1988.
- [41] M. J. Fischer and N. A. Lynch. A Lower Bound for the Time to Assure Interactive Consistency. *Information Processing Letters*, 14:183–186, 1982.
- [42] FlexRay Consortium et al. FlexRay communications system-protocol specification. *Version 2.1*, 2005.
- [43] E. G. Friedman. Clock Distribution Networks in Synchronous Digital Integrated Circuits. *Proceedings of the IEEE*, 89(5):665–692, 2001.
- [44] G. Fuchs and A. Steininger. VLSI Implementation of a Distributed Algorithm for Fault-Tolerant Clock Generation. *Journal of Electrical and Computer Engineering*, 2011(936712), 2011.
- [45] M. Függer, E. Armengaud, and A. Steininger. Safely Stimulating the Clock Synchronization Algorithm in Time-Triggered Systems – A Combined Formal and Experimental Approach. *IEEE Transactions on Industrial Informatics*, 5(2):132–146, 2009.
- [46] M. Függer, M. Hofstätter, C. Lenzen, and U. Schmid. Efficient Construction of Global Time in SoCs despite Arbitrary Faults. In *16th Conference on Digital System Design (DSD)*, pages 142–151, 2013.
- [47] M. Függer, R. Najvirt, T. Nowak, and U. Schmid. Towards Binary Circuit Models That Faithfully Capture Physical Solvability. In *Design, Automation, and Test in Europe (DATE)*, 2015.
- [48] M. Függer, T. Nowak, and U. Schmid. Unfaithful Glitch Propagation in Existing Binary Circuit Models. In *19th International Symposium on Asynchronous Circuits and Systems (ASYNC)*, pages 191–199, 2013.
- [49] M. Függer and U. Schmid. Reconciling Fault-Tolerant Distributed Computing and Systems-on-Chip. *Distributed Computing*, 24(6):323–355, 2012.
- [50] M. Függer, U. Schmid, G. Fuchs, and G. Kempf. Fault-Tolerant Distributed Clock Generation in VLSI Systems-on-Chip. In *6th European Dependable Computing Conference (EDCC)*, pages 87–96, 2006.
- [51] R. Ginosar. Fourteen Ways to Fool Your Synchronizer. In *9th International Symposium on Asynchronous Circuits and Systems (ASYNC)*, pages 89–96, 2003.
- [52] International Technology Roadmap for Semiconductors, 2012. <http://www.itrs.net>.

- [53] W. Jang and A. J. Martin. SEU-Tolerant QDI Circuits. In *11th International Symposium on Asynchronous Circuits and Systems (ASYNC)*, pages 156–165, 2005.
- [54] W. Jang and A. J. Martin. A soft-error-tolerant asynchronous microcontroller. In *13th NASA Symposium on VLSI Design*, 2007.
- [55] D. K. Kaynar, N. Lynch, R. Segala, and F. Vaandrager. *The Theory of Timed I/O Automata*. Morgan & Claypool Publishers, 2006.
- [56] S. Keller, M. Katelman, and A. J. Martin. A Necessary and Sufficient Timing Assumption for Speed-Independent Circuits. In *15th Symposium on Asynchronous Circuits and Systems (ASYNC)*, pages 65–76, 2009.
- [57] V. King and J. Saia. Breaking the $O(N^2)$ Bit Barrier: Scalable Byzantine Agreement with an Adaptive Adversary. *Journal of the ACM*, 58(4):18:1–18:24, 2011.
- [58] D. J. Kinniment. *Synchronization and Arbitration in Digital Systems*. Wiley, Chichester, 2008.
- [59] D. J. Kinniment, A. Bystrov, and A. V. Yakovlev. Synchronization Circuit Performance. *IEEE Journal of Solid-State Circuits*, SC-37(2):202–209, 2002.
- [60] L. Kleeman and A. Cantoni. On the Unavoidability of Metastable Behavior in Digital Systems. *IEEE Transactions on Computers*, C-36(1):109–112, 1987.
- [61] H. Kopetz and G. Bauer. The time-triggered architecture. *Proceedings of the IEEE*, 91(1):112–126, 2003.
- [62] I. Koren and Z. Koren. Defect tolerance in VLSI circuits: Techniques and yield analysis. *Proceedings of the IEEE*, 86(9):1819–1838, Sep 1998.
- [63] L. Lamport. The Temporal Logic of Actions. *ACM Transactions on Programming Languages and Systems*, 16(3):872–923, 1994.
- [64] E. A. Lee and P. Varaiya. *Structure and Interpretation of Signals and Systems*. LeeVaraiya.org, 2nd edition, 2011.
- [65] C. Lenzen, M. Perner, M. Sigl, and U. Schmid. Byzantine Self-Stabilizing Clock Distribution with HEX: Implementation, Simulation, Clock Multiplication. In *6th Conference on Dependability (DEPEND)*, 2013.
- [66] C. Lenzen, J. Rybicki, and J. Suomela. Towards Optimal Synchronous Counting. In *34th Symposium on Principles of Distributed Computing (PODC)*, 2015.
- [67] J. Lundelius and N. Lynch. An Upper and Lower Bound for Clock Synchronization. *Information and Control*, 62(2-3):190–204, 1984.
- [68] M. Malekpour. A Byzantine-Fault Tolerant Self-stabilizing Protocol for Distributed Clock Synchronization Systems. In *9th Conference on Stabilization, Safety, and Security of Distributed Systems (SSS)*, pages 411–427, 2006.
- [69] M. Malekpour. A Self-Stabilizing Byzantine-Fault-Tolerant Clock Synchronization Protocol. Technical report, NASA, 2009. TM-2009-215758.

- [70] R. Manohar and A. J. Martin. Quasi-delay-insensitive circuits are turing-complete. Technical report, Pasadena, CA, USA, 1995.
- [71] L. Marino. General Theory of Metastable Operation. *IEEE Transactions on Computers*, C-30(2):107–115, 1981.
- [72] A. J. Martin. Compiling communicating processes into delay-insensitive VLSI circuits. *Distributed Computing*, 1(4):226–234, 1986.
- [73] A. J. Martin. The Limitations to Delay-insensitivity in Asynchronous Circuits. In *Sixth MIT Conference on Advanced Research in VLSI*, AUSCRYPT '90, pages 263–278, Cambridge, MA, USA, 1990. MIT Press.
- [74] A. J. Martin. Synthesis of asynchronous VLSI circuits. Technical report, DTIC Document, 2000.
- [75] A. J. Martin and M. Nystrom. Asynchronous Techniques for System-on-Chip Design. *Proceedings of the IEEE*, 94(6):1089–1120, Jun 2006.
- [76] D. Mavis and P. Eaton. SEU and SET Modeling and Mitigation in Deep Submicron Technologies. In *45th Annual IEEE International Reliability physics symposium*, pages 293–305, April 2007.
- [77] M. Maza and M. Aranda. Interconnected Rings and Oscillators as Gigahertz Clock Distribution Nets. In *14th Great Lakes Symposium on VLSI (GLSVLSI)*, pages 41–44, 2003.
- [78] M. S. Maza and M. L. Aranda. Analysis of Clock Distribution Networks in the Presence of Crosstalk and Groundbounce. In *International IEEE Conference on Electronics, Circuits, and Systems (ICECS)*, pages 773–776, 2001.
- [79] D. G. Messerschmitt. Synchronization in Digital System Design. *IEEE Journal on Selected Areas in Communications*, 8(8):1404–1419, 1990.
- [80] C. Myers and T. H. Y. Meng. Synthesis of timed asynchronous circuits. *IEEE Transactions on Very Large Scale Integration (VLSI) Systems*, 1(2):106–119, June 1993.
- [81] L. W. Nagel and D. Pederson. SPICE (Simulation Program with Integrated Circuit Emphasis). Technical Report UCB/ERL M382, EECS Department, University of California, Berkeley, 1973.
- [82] R. Najvirt, M. Függer, T. Nowak, U. Schmid, M. Hofbauer, and K. Schweiger. Experimental Validation of a Faithful Binary Circuit Model. 2015. (appears in Proc. GLSVLSI'15).
- [83] R. Naseer and J. Draper. DF-DICE: A scalable solution for soft error tolerant circuit design. *IEEE International Symposium on Circuits and Systems (ISCAS)*, May 2006.
- [84] S. Nassif, K. Bernstein, D. Frank, A. Gattiker, W. Haensch, B. Ji, E. Nowak, D. Pearson, and N. Rohrer. High Performance CMOS Variability in the 65nm Regime and Beyond. In *Electron Devices Meeting, 2007. IEDM 2007. IEEE International*, pages 569–571, Dec 2007.

- [85] A. K. Palit, V. Meyer, W. Anheier, and J. Schloeffel. Modeling and Analysis of Crosstalk Coupling Effect on the Victim Interconnect Using the ABCD Network Model. In *19th IEEE International Symposium on Defect and Fault Tolerance in VLSI Systems (DFT)*, pages 174–182, Oct 2004.
- [86] M. Pease, R. Shostak, and L. Lamport. Reaching Agreement in the Presence of Faults. *Journal of the ACM*, 27:228–234, 1980.
- [87] M. Peercy and P. Banerjee. Fault Tolerant VLSI Systems. *Proceedings of the IEEE*, 81(5):745–758, May 1993.
- [88] M. O. Rabin. Randomized byzantine generals. In *24th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 403–409, 1983.
- [89] P. Shivakumar, M. Kistler, S. Keckler, D. Burger, and L. Alvisi. Modeling the Effect of Technology Trends on the Soft Error Rate of Combinational Logic. *International Conference on Dependable Systems and Networks (DSN)*, pages 389–398, 2002.
- [90] T. K. Srikant and S. Toueg. Optimal Clock Synchronization. *Journal of the ACM*, 34(3):626–645, 1987.
- [91] K. Stevens, R. Ginosar, and S. Rotem. Relative timing [asynchronous design]. *IEEE Transactions on Very Large Scale Integration (VLSI) Systems*, 11(1):129–140, Feb 2003.
- [92] Synopsis. CCS Timing. Technical white paper v2.0, 2006.
- [93] P. Teehan, M. Greenstreet, and G. Lemieux. A Survey and Taxonomy of GALS Design Styles. *IEEE Design and Test of Computers*, 24(5):418–428, 2007.
- [94] S. H. Unger. Asynchronous Sequential Switching Circuits with Unrestricted Input Changes. *IEEE Transactions on Computers*, 20(12):1437–1444, 1971.
- [95] A. Yakovlev, M. Kishinevsky, A. Kondratyev, L. Lavagno, and M. Pietkiewicz-Koutny. On the models for asynchronous circuit behaviour with OR causality. *Formal Methods in System Design*, 9(3):189–233, 1996.
- [96] C. Yeh, G. Wilke, H. Chen, S. Reddy, H. Nguyen, T. Miyoshi, W. Walker, and R. Murgai. Clock Distribution Architectures: a Comparative Study. In *7th Symposium on Quality Electronic Design (ISQED)*, pages 85–91, 2006.
- [97] T. Yoneda and C. Myers. Synthesis of Timed Circuits Based on Decomposition. *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems*, 26(7):1177–1195, July 2007.