

Transition and Cancellation in Concurrency and Branching Time

Vaughan R. Pratt

Stanford University

We review the conceptual development of (true) concurrency and branching time starting from Petri nets and proceeding via Mazurkiewicz traces, pomsets, bisimulation, and event structures up to higher dimensional automata (HDAs), whose acyclic case may be identified with triadic event structures and triadic Chu spaces. Acyclic HDAs may be understood as extending the two truth values of Boolean logic with a third value $\top$ expressing *transition*. We prove the necessity of such a third value under mild assumptions about the nature of observable events, and show that the expansion of any complete Boolean basis L to $L_{\top}$ with a third literal $\hat{a}$ expressing $a = \top$ forms an expressively complete basis for the representation of acyclic HDAs. The main contribution is a new value $\times$ of *cancellation*, sibling to $\top$, serving to distinguish $a(b + c)$ from $ab + ac$ while simplifying the extensional definitions of termination $\checkmark\mathcal{A}$ and sequence $\mathcal{A}\mathcal{B}$. We show that every HDAX (acyclic HDA with $\times$) is representable in the expansion of $L_{\top}$ to $L_{\top\times}$ with a fourth literal $\hat{a}$ expressing $a = \times$.

1. Introduction

1.1. Sequential and Concurrent Behavior

What distinguishes sequential computation, or for that matter sequential behavior of any kind, from concurrent? The usual viewpoint draws the following temporal distinction.

Sequential behavior allows only one event to happen at a time. With concurrent behavior multiple events may occur simultaneously.

Implicit in this distinction is the traditional view of time as evolving steadily and independently, providing a background against which to observe the processing of information. This view makes the behavior of time itself sequential—no two nanoseconds can overlap—while allowing arbitrarily independent behavior for bits.

The point of view espoused in this paper views time and information more symmetrically as two complementary or dual spaces. We view the points of time not so much as temporal instants as instantaneous events serving as state deltas: an event changes information but not time. Dually the points of information space are seen as time deltas: during a given state time passes while information remains fixed. These are of course idealized events and states; physical events always take time even if only 10^{-20} seconds while physical states are never completely stationary.

While this more symmetric view of time and information favors neither sequential nor concurrent behavior, it does raise the question of whether they can be distinguished in a symmetric way. We offer the following symmetric distinction.

Sequential behavior synchronizes the evolution of time and information. With concurrent behavior time and information evolve more independently.

In particular the sequential events (transitions) and states of traditional automata theory alternate in lock step, allowing time to be defined abstractly as the number of state changes. Concurrency decouples this relationship, with the result that the passage of both time and information can then be measured in terms of each other only in a distributed way. With the representation adopted below of processes as matrices over a set K , time and information are obtained naturally by standardly lifting a suitable generalized metric on K to respectively rows and columns *in the same way* to yield metrics on the respective temporal and information spaces. We shall treat this further in a more categorically-oriented follow-up paper for CONCUR'02 focusing on enrichment.

1.2. Refinements of Atomic Time

The simplest nontrivial view of atomic time, the view of time from the perspective of an atomic event, is that of Figure 1(a). Time divides into just two regions or *values*, 0 and 1, according to whether the event has not or has happened respectively. The event may pass from 0 to 1 (the figure orients time as flowing upwards), but having done so it may not return. That is, the (truth) value of the atomic proposition P_a expressing “ a has happened” can only increase monotonically with time. We write K for the allowed values, so in Figure 1(a) $K = \{0, 1\}$.

	(a)	(b)	(c)	(d)
$a b \stackrel{?}{=} ab + ba$	$a \begin{bmatrix} 0101 \\ 0011 \end{bmatrix} = a \begin{bmatrix} 0101 \\ 0011 \end{bmatrix}$	$a \begin{bmatrix} 0\ulcorner 10\ulcorner 10\ulcorner 1 \\ 000\ulcorner\ulcorner 111 \end{bmatrix} \neq a \begin{bmatrix} 0\ulcorner 1010\ulcorner 1 \\ 000\ulcorner\ulcorner 111 \end{bmatrix}$	$a \begin{bmatrix} 0101 \\ 0011 \end{bmatrix} = a \begin{bmatrix} 0101 \\ 0011 \end{bmatrix}$	$a \begin{bmatrix} 0\ulcorner 10\ulcorner 10\ulcorner 1 \\ 000\ulcorner\ulcorner 111 \end{bmatrix} \neq a \begin{bmatrix} 0\ulcorner 1010\ulcorner 1 \\ 000\ulcorner\ulcorner 111 \end{bmatrix}$
$a(b+c) \stackrel{?}{=} ab + ac$	$a \begin{bmatrix} 0111 \\ 0010 \\ 0001 \end{bmatrix} = a \begin{bmatrix} 0111 \\ 0010 \\ 0001 \end{bmatrix}$	$a \begin{bmatrix} 0\ulcorner 111111 \\ 000\ulcorner 100 \\ 00000\ulcorner 1 \end{bmatrix} = a \begin{bmatrix} 0\ulcorner 111111 \\ 000\ulcorner 100 \\ 00000\ulcorner 1 \end{bmatrix}$	$a \begin{bmatrix} 0111 \\ 001\ulcorner \\ 00\ulcorner 1 \end{bmatrix} \neq a \begin{bmatrix} 0111 \\ 001\ulcorner \\ 0\ulcorner\ulcorner 1 \end{bmatrix}$	$a \begin{bmatrix} 0\ulcorner 111111 \\ 000\ulcorner 1\ulcorner\ulcorner \\ 000\ulcorner\ulcorner 1 \end{bmatrix} \neq a \begin{bmatrix} 0\ulcorner 111\ulcorner 1111 \\ 000\ulcorner 1\ulcorner\ulcorner\ulcorner\ulcorner \\ 0\ulcorner\ulcorner\ulcorner\ulcorner 00\ulcorner 1 \end{bmatrix}$

Figure 1. Refinements of atomic time

The process $a||b$, perform events a and b independently, permits the states of a and b to be any of the four combinations 00, 10, 01, or 11, which appear as the four columns in the first box of Figure 1. The process $ab + ba$, perform a and b in either order, permits the same states, so this choice of K satisfies $a||b = ab + ba$. The process $a(b+c)$, perform a and then one of b or c , permits for a, b, c the states 000, 100, 110, and 101, as does $ab + ac$, perform one of a then b or a then c , so this K also satisfies $a(b+c) = ab + ac$.

Figure 1(b) refines this view with the recognition of an intermediate state $\ulcorner$ of *transition*. K now has 3 states, whence $a||b$ has $3^2 = 9$ states. $ab + ba$ has the same states with the exception of the state $\ulcorner\ulcorner$ denoting a and b simultaneously in transition, and so $\ulcorner$

now distinguishes $a||b$ from $ab+ba$. But even though $\sqcap$ also furnishes $a(b+c)$ and $ab+ac$ with additional states, the new states are the same for both, which therefore remain indistinguishable. Other roles for $\sqcap$ include distinguishing asymmetric from symmetric conflict, representation of resource limits, and a satisfactory notion of running time of a parallel computation.

Figure 1(c) proposes a different way of adding a third state, namely the state $\times$ of *cancellation*. A choice may necessitate cancelling an event, e.g. choosing a in $a+b$ automatically cancels b immediately, even before a happens. No choice is made in $a||b$ whence there are no cancellations. But even though $ab+ba$ entails a choice, both a and b happen on both branches and so neither is cancelled in either case, leaving $a||b = ab+ba$.

Now in $a(b+c)$, after a happens (state 100) a choice is required, at which point one of b or c is cancelled (state 1×0 or $10\times$ respectively). In $ab+ac$ this choice is made at the beginning, at which time one of b or c is cancelled (state 0×0 or $00\times$ respectively, with state 100 now disallowed). So $\times$ distinguishes $a(b+c)$ from $ab+ac$. Other roles for $\times$ include a more satisfactory notion of final state leading to improved definitions of sequence and termination.

Figure 1(d) is the evident amalgamation of 1(b) and 1(c) (though since the relationship of 1(a) to 1(b) is embedding while to 1(c) it is quotient, it would have to be some sort of mixed-variance amalgamation or epi-mono square construction.). Tables 1-4, encountered along the way as our more detailed story unfolds, correspond respectively to Figures 1(a)-(d).

We associate with the states 0, $\sqcap$, 1, $\times$ of K the respective literals $\bar{a}$, $\hat{a}$, a , and $\not{a}$, asserting respectively $a = 0$, $a = \sqcap$, $a = 1$, and $a = \times$. We shall show that every subset of K^A is representable as a disjunction of conjunctions of these four literals.

The earliest explicit representation of nonperformance we are aware of is Genrich and Thiagarajan's notion of L-value for Petri nets (GT84). Much more recently, and closer to our notion of cancellation, Rodriguez and Anger (RA01) give an analogous account of branching time in terms of an expansion of Allen's interval algebra (All84) of the 13 possible relationships between two intervals sliding past each other. The traditional account of interval algebra derives the 13 relationships from three binary relations $<$, $=$, and $>$ each between two endpoints one from each interval. Rodriguez and Anger consider a fourth relationship $||$ in which the two endpoints are no longer comparable, as when the parallel tracks on which the intervals slide diverge at some point, deriving 19 interval relationships. Section 3.7 considers this notion of branching time in more detail, addressing the meaning of the symmetric relationship $||$ when reinterpreted as a value of K . We show that it would have the meaning of an event that blocks as opposed to being cancelled, with the result that using $||$ analogously to $\times$ to distinguish $a(b+c)$ from $ab+ac$ would cause d to block in $a(b+c)d$, a crucial distinction between the respective characterizations of branching time by $||$ and $\times$.

1.3. Rationale

For state-of-the-art software and hardware verification, the traditional state-oriented view of computation is all that is needed. In this view a computer system, implemented in

either hardware, software, or some mix thereof, is understood abstractly as an edge-labeled graph $G = (V, E)$. Each vertex $u \in V$ denotes a possible *state* of the system. Each edge $e \in E$ from u to v is labeled with a symbol or *action* $\lambda(e) = \mathbf{a}$ indicating that if the system is in state u and action $\mathbf{a} = \lambda(e)$ occurs, the next state of the system will be v (or *can be* v in the case of nondeterministic computation).

Correctness of the system so modeled is defined in terms of predicates at program points L that are required to hold of the system state x whenever the program counter (whether real or virtual) in x is at L . Only states as the effects of transitions are observable, not the transitions themselves.

The canonical application of this viewpoint is verification, whose function originally was conceived as an imprimatur certifying the whole program correct once and for all, but which in recent years has retargeted logical laws to serve as in-laws perpetually finding fault with the system, rooting out bugs in parallel with evolution of the system design while exposing tacit assumptions of the system designers. This view serves its intended purposes well and has become the basis for a thriving (albeit still cottage) verification industry.

There are however two assumptions tying the system designer's hands here, namely fixed granularity and sequential observation.

Fixed granularity. This is the assumption that actions of the system are built up from atomic actions whose start and end are both observable, but not the period in between—the transition is quicker than the eye. This assumption justifies the equation $a||b = ab+ba$ identifying independent occurrence of atoms a, b with their occurrence in either order. A natural name for the latter is atomic mutual exclusion.

The assumption of fixed granularity is called into question by the widespread practice of maintaining confidentiality of source code. Confidentiality has two benefits, a competitive one of trade secrecy and a technological one of being able to improve the implementation without compromising the specification. However it also has the effect of presenting some system components as atomic to the user even though the implementor sees them as built up from smaller atoms. While it certainly *suffices* for system correctness to organize such user-atomic components so as to satisfy atomic mutual exclusion, this requirement can in many cases impose a sufficient burden on both implementors and users as to justify asking whether the user *must* assume atomic mutual exclusion in order to verify larger systems built from such components. The study of concurrent computation bifurcates according to the answer, with this paper firmly in the camp dedicated to exploring what is possible without the assumption of fixed granularity, and hence without assuming atomic mutual exclusion.

Sequential observation. This is the assumption that every dispute can be resolved by one referee. From a legal standpoint it might seem as though the nationals of any country with a Supreme Court would be on safe ground with this assumption, at least within their own borders, since that court will ultimately resolve any dispute of sufficient significance. But the resolution may come too late to be useful, or the one referee may be overwhelmed with such requests, or the centralized arbiter may recognize a prior period of time where one distributed system observed another with results that the arbiter can see in hindsight

could not have been obtained at the time by a single observer even in principle regardless of available resources (PP96).

The appropriate operator combining two distributed systems one observing the other is orthocurrence (Pra85; Pra86; CCMP91; GP93). Orthocurrence can be understood either symmetrically as the *interaction* $\mathcal{A} \otimes \mathcal{B}$ of two processes $\mathcal{A}$ and $\mathcal{B}$, or ostensibly asymmetrically as the *observation* $\mathcal{A} \multimap \mathcal{B}^\perp$ of states of process $\mathcal{B}$ from vantage points of process $\mathcal{A}$, or the dual observation $\mathcal{B} \multimap \mathcal{A}^\perp$ (giving a sense in which observation is really symmetric) (Pra01). Besides being the only viable candidate proposed to date for this role, it has the additional benefits of an intuitive definition in terms of crossword puzzles, and attractive algebraic and logical properties, specifically both k -2 logic (corresponding to k -adic Chu spaces) and Girard's linear logic (Gir87).

Now orthocurrence makes no sense for the standard state-based or transition-system view of computation because its events consist of pairs of events each taken from one of two interacting or "orthocurrent" systems. Product does appear in automata theory, but only for concurrence, which forms the product of state sets; in contrast orthocurrence multiplies event sets. The incompatibility of orthocurrence with the standard state-based view of behavior makes it of predictably limited interest for systems designed from that perspective.

We maintain nonetheless that orthocurrence as the multiplicative form of concurrency is every bit as important as concurrence, its additive counterpart. (This distinction is that of multiplicative-additive linear logic or MALL, where *concurrence* and *orthocurrence* are called respectively *plus* and *tensor* and a third operation *perp* is introduced whose action on processes is to interchange the perspectives of observer and observed, with the side effect of interchanging events and states.)

Our standard example of orthocurrence has long been that of trains passing through stations, which this paper spices up with examples involving pigeons choosing holes subject to various constraints. Further evidence for the utility of orthocurrence appears in (Pra00; RA01) where Allen interval algebras for various models of time, previously obtained tediously by ad hoc manual methods (AR91; RA93b; RA93a), are all obtained uniformly and automatically using orthocurrence.

To these two assumptions of the standard model one might add a third, discrete time. The standard model assumes that time passes in discrete steps, justified by the behavior of digital systems which are understood abstractly as so behaving. Yet flip-flops, Schmitt triggers, etc. are implemented from analog transistors, signal processing and image processing involve the digitization of analog signals, and buses need to decide quickly which of two independently arriving pulses arrived first on two request lines in order to grant a shared resource to one requester without inadvertently granting it also to the other. The transition-based view of computation has been coerced by many to this analog viewpoint in recent years, but not gracefully in the sense that the usual approach has been to interleave small intervals of the continuum, which only exacerbates the problems created by the assumptions of fixed granularity and sequential observation. A suitable framework for this purpose that is more robust than interleaved intervals is generalized metric spaces (CCMP91).

2. Extant Models of Behavior

2.1. Petri Nets

The earliest extant model of concurrency is the **Petri net** (Pet62), a structure consisting of statements[†] and transitions. Truth values of statements are natural numbers, and a statement holds when its truth is nonzero. A marking or *state* of a Petri net is a truth assignment to its statements. An *event* is the firing of a transition. Firing is regulated by associating to each transition two sets of statements, its preconditions and its postconditions. When all the preconditions of a transition hold it can fire, and when it does so it decrements its preconditions and increments its postconditions.

A sequential run of a Petri net is a sequence of firings. A concurrent run is a so-called *occurrence* or *causal net*, an acyclic Petri net each of whose statements is postcondition to just one transition and precondition to just one other, whose transitions now constitute events in the sense that they can each fire at most once. Variations in width of a run express the varying degrees of concurrency during the run, i.e. the number of transitions that can be firing concurrently.

2.2. Mazurkiewicz Traces

Just as many different state automata can all exhibit the same sequential behavior, so can many different Petri nets all manifest the same concurrent behavior. One therefore seeks a more abstract notion of behavior that minimizes irrelevant implementation detail.

Just as formal languages abstract state automata by modeling the latter's runs as strings, so do **Mazurkiewicz traces** (Maz77) abstract Petri nets by modeling their runs as equivalence classes of strings. Consider for example a Petri net consisting of two isolated transitions labeled respectively **a** and **b**, a transition being considered isolated when it has a true precondition and a false postcondition each shared with no other transition. Whereas the sequential behavior of this Petri net would consist of the two strings **ab** and **ba**, its Mazurkiewicz behavior would consist of the single Mazurkiewicz trace resulting from identifying these two strings by way of indicating that *a* and *b* had fired independently.

The basis for this identification is *action independence*, defined as an irreflexive symmetric binary relation I on an alphabet Σ of actions. Action independence induces a congruence $\cong_I$ (with respect to concatenation) on the monoid Σ^* of all finite strings on Σ , namely the least congruence $\cong$ for which **ab** $\cong$ **ba** for all **(a, b)** $\in I$. (We distinguish between actions **a** as event *labels* and events *a* as action *instances* using boldface vs. italics respectively.) That is, two strings are I -congruent when one can be obtained from the other by a series of exchanges of adjacent independent actions. The quotient $\Sigma^* / \cong_I$ consisting of the congruence classes of $\cong_I$ constitutes the monoid of all Mazurkiewicz traces over Σ induced by I . A concurrent behavior is then a set of such traces, i.e. a subset of $\Sigma^* / \cong_I$.

A fundamental limitation of Mazurkiewicz traces is that independence is global: if

[†] *Stellen*, traditionally translated more literally as places.

actions **a** and **b** are independent anywhere they are independent everywhere. Consider the four actions one would use to model the transmission and receipt of bits through an asynchronous (buffered) channel, namely **T₀**, **T₁** (transmission of 0 or 1) and the corresponding receipts **R₀**, **R₁**. Receipt of the first bit necessarily depends on its transmission, but for proper asynchrony that same receipt should be independent of the transmission of the second bit. So if both transmitted bits are the same, say zero, both bits are transmitted via the same action, namely **T₀**. But then the action of receipt of the first bit, namely **R₀**, must be both dependent on and independent of **T₀** as the common action of both the first and second transmitted bits, an impossibility for Mazurkiewicz traces.

2.3. Pomsets

Pomsets (Gra81; Pra82; Pra86; Gis88) overcome this limitation of Mazurkiewicz traces by being based on independence not of actions but of action instances or *events*. A pomset or partially ordered multiset is a Σ -labeled poset $(A, \leq, \lambda)$ where $\leq$ partially orders a set A of events, with $a \leq b$ indicating that a necessarily happens before b , and $\lambda : A \rightarrow \Sigma$ labels each event $a \in A$ with an action $\lambda(a) \in \Sigma$. A string over Σ is the special case of a pomset which is finite and linearly ordered. A Mazurkiewicz trace in $\Sigma^* / \cong_I$ is the special case of a pomset which is finite and for any pair a, b of distinct events, if a and b are order-incomparable then $\lambda(a)$ and $\lambda(b)$ are independent, and if $\lambda(a)$ and $\lambda(b)$ are independent and $a \leq b$ then there exists a third event c satisfying $a < c < b$. When I is empty all events are order-comparable, i.e. the order is linear, in which case Mazurkiewicz traces are ordinary strings as expected. Strings, Mazurkiewicz traces, and pomsets thus form increasingly more general kinds of behaviors.

From the perspective of true concurrency, pomsets (and hence strings and Mazurkiewicz traces) represent behaviors that are not only deterministic but nonbranching: there is no element of choice whatsoever. Whereas “width” in an automaton results from disjunctive branching (just one branch is taken), the visually similar “branches” in a pomset are all taken concurrently, thereby constituting conjunctive branching. We will shortly settle on processes containing both kinds of branching. In that setting it is both customary and natural to hide conjunctive “branching” by lumping together strings, Mazurkiewicz traces, and pomsets under the general heading of *traces*, understood as deterministic nonbranching behaviors, and to reserve the term “branching” exclusively for its disjunctive form.

2.4. Branching Time and Bisimulation

Prior to the last two decades or so, nondeterminism and choice were customarily modeled abstractly in terms of sets of traces denoting the possible alternative runs. The abstract behavior of an automaton was taken to be the set of strings it accepted, and that of a Petri net the set of its permitted firing sequences. This catered for nondeterminism and branching by effectively making all choices blindly at the outset. To handle the case of a decision predicated on the outcome of a test performed during the behavior, if that outcome turns out to conflict with the choice made in the beginning, the trace *blocks*,

which means not that it terminates successfully but rather that it is abandoned as having been a bad choice, like a prospector giving up on an unproductive mining claim. The distinction between successful termination and blocking is observable in the construct $\mathcal{A}\mathcal{B}$, perform $\mathcal{A}$ then $\mathcal{B}$, in that when process $\mathcal{A}$ terminates successfully process $\mathcal{B}$ may proceed, whereas if $\mathcal{A}$ blocks then so does $\mathcal{A}\mathcal{B}$ and $\mathcal{B}$ does not get to start.

Models of behavior that make their decisions in the beginning in this way are said to obey **trace semantics**. The litmus test for trace semantics is the equation $a(b + c) = ab + ac$, meaning that choosing one of b or c after performing a is the same as choosing one of ab (perform a then b) or ac . Students of automata theory will recognize this equation as holding for formal languages (sets of strings) where ab is concatenation and $a + b$ union, but it also holds for processes as sets of traces of any kind including pomsets. *Every model of computation considered thus far obeys trace semantics.*

Now it is common sense that a prematurely made decision may adversely limit one's options by committing one to a path before sufficient information is available to justify that decision. In that respect $a(b + c)$ is more prudent than $ab + ac$, making $a(b + c) = ab + ac$ inappropriate.

A finer equivalence than trace equivalence that recognizes the impact of decision timing is **bisimilarity** (Mil80; Par81), which satisfies only $(a + b)c = ac + bc$ and not $a(b + c) = ab + ac$.

Bisimilarity is defined not on events but on states of an automaton or transition graph, via the notion of simulation. For the purpose of this definition we take an automaton to be a vertex-and-edge-labeled graph $\mathcal{X} = (X, E)$ whose vertices are states $x \in X$ labeled with predicates $\lambda_X(x)$ and whose edges in E are transitions from x to x' labeled with actions $\mathbf{a}$, denoted $x \xrightarrow{\mathbf{a}} x'$. A binary relation $R \subseteq X \times Y$ relating states of $\mathcal{X}$ to states of $\mathcal{Y}$ is a **simulation** between $\mathcal{X}$ and $\mathcal{Y}$ when for all (x, y) in R , (i) $\lambda_X(x) = \lambda_Y(y)$, and (ii) for every transition $x \xrightarrow{\mathbf{a}} x'$ in E_X there exists a like-labeled transition $y \xrightarrow{\mathbf{a}} y'$ in E_Y , the *simulating* transition, for which $(x', y') \in R$. State y in $\mathcal{Y}$ **simulates** state x in $\mathcal{X}$ when there exists a simulation between $\mathcal{X}$ and $\mathcal{Y}$ containing (x, y) . For automata with a specified initial state, when the initial state of $\mathcal{Y}$ simulates the initial state of $\mathcal{X}$ we say that $\mathcal{Y}$ simulates $\mathcal{X}$.

For an example of simulation consider the automaton $\mathcal{X} = \begin{array}{c} \bullet \\ \swarrow \xrightarrow{a} \bullet \searrow \xrightarrow{b} \bullet \\ \nwarrow \xrightarrow{a} \bullet \searrow \xrightarrow{c} \bullet \end{array}$ realizing the behavior $ab + ac$ and $\mathcal{Y} = \begin{array}{c} \bullet \\ \swarrow \xrightarrow{a} \bullet \searrow \xrightarrow{b} \bullet \\ \nwarrow \xrightarrow{a} \bullet \searrow \xrightarrow{c} \bullet \end{array}$ realizing $a(b + c)$, with states labeled vacuously (the same predicate on every state). Each of their states is reached from its parent initial state by exactly one of the strings ϵ, a, ab, ac . Take $R \subseteq X \times Y$ to consist of all pairs (x, y) such that x and y are reached by the same string (one such pair for each state of $\mathcal{X}$, in fact R is a surjection from X onto Y that is injective except for identifying the two states of $\mathcal{X}$ reached by a). R is easily verified to be a simulation, and furthermore it relates initial states and so is a simulation of $\mathcal{X}$ by $\mathcal{Y}$. However the converse of R is not a simulation of $\mathcal{Y}$ by $\mathcal{X}$ because the (x, y) in R for which x is the midpoint of ab is such that y has a c transition from it while x does not. Moreover there exists no simulation at all of $\mathcal{Y}$ by $\mathcal{X}$, the price of $\mathcal{X}$'s imprudently early decision.

A **bisimulation** is a simulation whose converse is also a simulation. Processes with initial states are bisimilar when there is a bisimulation between them that makes their initial states bisimilar. Semantics making finer distinctions than trace equivalence have

been provided for both CSP (BHR84) and CCS (Mil80). Such semantics are distinguished from trace or linear time semantics under the rubric of branching-time semantics, with bisimilarity being the finest such, at least for ordinary yes-no nondeterminism as opposed e.g. to probabilistically governed choices.

Bisimilarity breaks the symmetry of time inherent not only in traditional automata theory based on sets of strings but also in accounts of concurrency based on sets of Mazurkiewicz traces and sets of pomsets—in general on sets of traces. It characterizes branching time intensionally by making us write $(a + b)c \cong ac + bc$ expressing the bisimilarity of $(a + b)c$ and $ac + bc$ and $a(b + c) \not\cong ab + ac$ for the non-bisimilarity of $a(b + c)$ and $ab + ac$ when an extensional characterization would permit the conceptually simpler $(a + b)c = ac + bc$ and $a(b + c) \neq ab + ac$.

2.5. Event Structures

Our program at this point is to characterize both concurrency and branching time extensionally. We take as our basic framework for this program that of event structures (NPW81; Win80; Win86; Win88a).

An **event structure**[‡] $(A, \leq, \#)$ consists of a partial order $(A, \leq)$ [§] equipped with a symmetric irreflexive binary relation $\#$ of *conflict*, satisfying the following condition.

Axiom ES. For all events $a, b, c \in A$, $a \# b$ and $b \leq c$ implies $a \# c$.

A **configuration** of an event structure $(A, \leq, \#)$ is a conflict-free downset $x \subset A$.[¶] That is, given $a \in x$, if $b \leq a$ then $b \in x$ (the notion of downset), and if $a \# b$ then $b \notin x$ (the meaning of conflict).

The set X of all configurations of an event structure, called a *family* of configurations, can be understood as the states of an acyclic automaton realizing the event structure. The initial state of this automaton is the empty set of states, which is vacuously conflict-free and a downset and hence a configuration. Its transitions are the inclusions between configurations, oriented to pass from the smaller state to the larger.

Event structures as such are unlabeled. They may be labeled with labels drawn from an action alphabet Σ , exactly as with pomsets. A **labeled event structure** $(A, \leq, \#, \lambda)$ is an event structure together with a labeling function $\lambda : A \rightarrow \Sigma$. As for pomsets, labeling specifies for each event the action it performs. Many events may perform the same action, e.g. while a communication channel may transmit any number of bits each transmission is an instance of one of two actions $\mathbf{T}_0$ and $\mathbf{T}_1$ denoting transmission of respectively 0 or 1.

Labels carry over to the family of configurations in a straightforward way. A transition from x to y is labeled with the multiset of labels of the events in $y - x$ (those events that

[‡] Subsequently described as *prime coherent* event structures, to distinguish this basic class from larger classes of event structures introduced later.

[§] Note the absence of cardinality restrictions: A may be finite (even empty), countably infinite, or uncountable.

[¶] Configurations represent states. For uniformity we shall write states as lower-case $x, y, z, \dots$ independently of whether they take the form of vertices of a graph or configurations of an event structure. We reserve upper case X for sets of states.

happened during the passage from x to y). This multiset can be defined as the restriction of λ to $y - x$.

Labels make it easy to distinguish $\mathbf{a}||\mathbf{b}$ from $\mathbf{ab} + \mathbf{ba}$ and $\mathbf{a}(\mathbf{b} + \mathbf{c})$ from $\mathbf{ab} + \mathbf{ac}$. For the former we take $\mathbf{a}||\mathbf{b}$ to consist of two events labeled respectively $\mathbf{a}$ and $\mathbf{b}$, and $\mathbf{ab} + \mathbf{ba}$ to consist of four events labeled respectively $\mathbf{a}$, $\mathbf{b}$, $\mathbf{b}$, and $\mathbf{a}$ with the event order and with every event of one branch of the choice in conflict with every event of the other. For the latter we take $\mathbf{a}(\mathbf{b} + \mathbf{c})$ to consist of three events labeled $\mathbf{a}$, $\mathbf{b}$, and $\mathbf{c}$, with $\mathbf{a}$ preceding both $\mathbf{b}$ and $\mathbf{c}$, and with $\mathbf{b}$ and $\mathbf{c}$ in conflict, and we take $\mathbf{ab} + \mathbf{ac}$ to consist of the evident four-event event structure.

Now these distinctions are made by the questionable practice of breaking what is intuitively a single event into multiple events, the copies of which are then kept track of via the labels. The question then arises as to whether a more conservative accounting of events is possible without having to duplicate an event which then needs to be kept track of with labels. We answer this in the affirmative by accounting for both concurrency and branching time entirely with unlabeled event structures. As will become clear in due course, Figure 1 shows that this is not possible with ordinary or dyadic event structures, but becomes possible when further qualia besides 0 and 1 are added to time, corresponding to the constructs *before* (may happen), *during* (is happening), *after* (has happened), and *instead of* (won't happen). The states themselves are those of *initialization*, *transition*, *termination*, and *cancellation*.

2.6. Nature of Prime Coherent Event Structures

Before embarking on this project it will help fix ideas to bide a while with the dyadic case, and furthermore with the original prime coherent subcase.

The nature of prime coherent event structures is brought out by their following properties, leading up to the observation of Corollary 5 that prime coherent event structures are extensional (distinct event structures have distinct interpretations).

Lemma 1. Every principal downset of an event structure is a configuration of that event structure.

(A downset is *principal* when it contains its sup, equivalently when it contains an element a such that the downset can be expressed as $\{b | b \leq a\}$, written $\downarrow a$.)

Proof. For if not, that downset must contain a conflicting pair both below a . Two applications of Axiom ES yield $a \# a$, contradicting irreflexivity of $\#$. $\square$

Corollary 2. $a \leq b$ and $a \# b$ are mutually exclusive. (For otherwise the principal downset $\downarrow b$ would not be a configuration.)

Lemma 3. If a and b do not appear together in any configuration of an event structure then $a \# b$ in that event structure.

Proof. If a and b do not appear together in any configuration then $\downarrow a \cup \downarrow b$ must not be a configuration. But it is a downset and hence can only fail to be a configuration by

containing two events in conflict. By Lemma 1 one must be in $\downarrow a$ and the other in $\downarrow b$. Two applications of Axiom ES then yield $a\#b$. $\square$

Lemma 4. If a does not appear in any configuration of an event structure without b then $b \leq a$ in that event structure.

Proof. By Lemma 1, $\downarrow a$ is a configuration, which by hypothesis contains b , so $b \leq a$. $\square$

Corollary 5. An event structure is uniquely determined by its family of configurations (straightforward given Lemmas 3 and 4).

Lemma 6. The number of event structures on a given set having one or two events is respectively one or four.

Proof. A singleton admits only one partial order and one irreflexive binary relation. A doubleton admits three partial orders and two irreflexive symmetric binary relations for a total of six, but the two of these in which the events are in conflict and the order is linear are ruled out by Corollary 2. The four survivors are then seen by inspection to be event structures. $\square$

The four two-event structures are succinctly characterized as those that omit no configurations or one non-empty configuration. These are clearly nonisomorphic with the exception of the two linearly ordered event structures, so up to isomorphism there are only three event structures on two events.

2.7. Event Structures via Logic

Event structures have evolved with time and experience to successively larger classes, most notably through the work of Glynn Winskel (Win80; Win82; Win86; Win88b). A uniform framework from which to view this evolution casts an event structure as a Boolean operation^{||} whose satisfying assignments meet certain plausible restrictions. One such restriction might insist on the all-zero assignment (the initial state) being among the satisfying assignments. Another might require that every satisfying assignment be reachable from the all-zero assignment via a *monotone Hamming* path through the satisfying assignments, one that proceeds by changing one atom at a time from false to true. (Prime coherent event structures meet both of these restrictions.) The evolution of more general event structures can then be described simply as the gradual removal of such restrictions.

In more detail, an event structure $\mathcal{E}$ on a set A of events is viewed as a Boolean operation $f_{\mathcal{E}}$ on A , namely a function of type $2^A \rightarrow 2$, equivalently a set of sets of type 2^{2^A} . The elements of A are reinterpreted as atoms or propositional variables. Event a is understood as an atomic proposition P_a expressing “ a has happened.” We abbreviate P_a to a and rely on context to disambiguate. The configurations of $\mathcal{E}$ then become the

^{||} Boolean operations are used in preference to Boolean formulas because the former as restricted are in one-one correspondence with event structures while with the latter the correspondence is many-one.

satisfying assignments of $f_{\mathcal{E}}$, those assignments $\alpha : A \rightarrow 2$ of truth values to atoms such that $f_{\mathcal{E}}(\alpha) = 1$.

Prime coherent event structures, the class treated above, are translated as follows. Each temporal precedence $a \leq b$ is rendered as the implication $b \rightarrow a$ between atoms (or $\bar{a} \rightarrow \bar{b}$ between literals), namely if b has happened then a has happened. (The implication arrow thus points backwards in time.) Each conflict $a \# b$ is expressed as the proposition $\bar{a} \wedge \bar{b}$, or $a \rightarrow \bar{b}$, or $b \rightarrow \bar{a}$. And Axiom ES simply disappears, being absorbed into the general Boolean framework as the application of transitivity to $a \rightarrow \bar{b}$ and $\bar{b} \rightarrow \bar{c}$.

We then define $T_{\mathcal{E}}$ to be the least literal-implication theory on A containing the above implications of $\mathcal{E}$. To this end we define a **literal-implication theory on A** to be a reflexively transitively closed set of implications $P \rightarrow Q$ between literals P, Q each of the form a or $\bar{a}$ for some $a \in A$. We further require that such a theory also be closed under *modus tollens*, meaning that if $P \rightarrow Q$ is in the theory then so is $\bar{Q} \rightarrow \bar{P}$.

The associated Boolean operation $f_{\mathcal{E}}$ holds (evaluates to 1) just where every implication in $T_{\mathcal{E}}$ holds, i.e. $f_{\mathcal{E}}$ is the denotation of $\bigwedge T_{\mathcal{E}}$.

We say that a Boolean formula on A holds in an event structure $\mathcal{E}$ when it is a consequence of the Boolean operation $f_{\mathcal{E}}$ associated to $\mathcal{E}$, that is, when it holds in every assignment (of truth values to atoms of A) satisfying $f_{\mathcal{E}}$.

Any literal-implication theory of a prime coherent event structure must satisfy the following three conditions. It may not contain any implication $\bar{a} \rightarrow b$ from a negative literal to a positive (equivalently disjunctions of positive literals) because these are falsified by the empty (initial) configuration. It may not contain any implication $a \rightarrow \bar{a}$ as this expresses $a \# a$ which by Lemma 3 violates irreflexivity of $\#$. And it may not contain both implications $a \rightarrow b$ and $b \rightarrow a$ (equivalently $a \equiv b$ cannot hold in $\mathcal{E}$), by Lemma 4 and antisymmetry of $\leq$.

Theorem 7. Every literal implication theory satisfying the above three conditions arises as the translation of some prime coherent event structure.

Proof. Given such a theory T on A , construct the event structure $\mathcal{E}$ on A for which $a \leq b$ holds when $b \rightarrow a$ is in T and for which $a \# b$ holds when $a \rightarrow \bar{b}$ is in T . Claim: $T = T_{\mathcal{E}}$. This follows because for every $a \in A$, $a \rightarrow a$ and $\bar{a} \rightarrow \bar{a}$ must be present by reflexivity, while $a \rightarrow \bar{a}$ and $\bar{a} \rightarrow a$ must be absent by two of the conditions. For all $a \neq b$ we only need consider all implications from signed a to signed b and no implications from signed b to signed a since the latter set is completely determined from the former by *modus tollens*. $\bar{a} \rightarrow b$ is forbidden, and the conditions allow at most one of $a \rightarrow b$, $\bar{a} \rightarrow \bar{b}$ (equivalent to $b \rightarrow a$), and $a \rightarrow \bar{b}$. That one will clearly be in $T_{\mathcal{E}}$ if and only if it was in T . $\square$

A useful formula not expressible with prime coherent event structures is $b \rightarrow (a \vee c)$. This asserts that b cannot happen until one of a or c has happened, e.g. the vending machine will not produce a soft drink until either a dollar bill or four quarters are inserted. Yet another is $(a \wedge c) \rightarrow b$, which asserts that a and c are in conflict (cannot both happen) until b happens, e.g. you can't buy both the candy and the pen you want with the dollar you have until you get a second dollar. This could be called either temporary

conflict or deferred concurrency. Neither of these are expressible as a conjunction of literal implications. In view of the evident utility of such constructs it is natural to look for ways to broaden the definition of event structures to admit them.

In this view of the evolution of the class of event structures by removal of the less natural or undesired restrictions, there is no clear-cut boundary for what counts as “natural.” This makes the inevitable limit that of no restrictions, i.e. any infinitary Boolean operation may be allowed. This is the point of view taken in (GP93) and (vGP95), which study event structures at this level of generality from the viewpoint of extensional Chu spaces over a two-letter alphabet $K = \{0, 1\}$, a connection that is not the central point of this paper but whose notation and concepts we shall borrow from.

Taking this viewpoint here, we may summarize the development thus far by characterizing event structures on a set A of events as precisely the Boolean operations on A , that is, entities of type 2^{2^A} . The automaton associated with such an event structure has for its states the set X of satisfying assignments of that operation, and its multi-event transitions or **steps** are all pairs (x, y) of states for which $x \leq y$ coordinatewise: each event may stay in its present state, or pass from 0 to 1, but not from 1 to 0. Note that the step relation is transitive for now, but this will change later on when we introduce $\sqsubset$.

We make a point of distinguishing $f : 2^A \rightarrow 2$ from $g : 2^B \rightarrow 2$ for $A \neq B$, even when f and g are both constantly true or both constantly false. Now $2^A \rightarrow 2$ and 2^{2^A} are isomorphic via the obvious pairing of $f : 2^A \rightarrow 2$ with $(A, \{x \in 2^A \mid f(x) = 1\})$. We shall abuse type and speak as though f and its matching (A, X) were identical. From now on we treat event structures and (abstract) Boolean propositions as one and the same.

One benefit of this identification is simplification of concepts. One drawback that we shall go into in more detail shortly is that the naming conventions of process algebra and logic are not fully compatible: the constant 0 and the atom a denote different entities in each. We deal with this by assuming the process algebra interpretation by default, indicating any exceptions explicitly.

One way of presenting (A, X) is as an $A \times X$ matrix whose entry at (a, x) is the value of a in state x . Thus the independent behavior $a \parallel b$ of two states a, b is presented as the matrix $\begin{smallmatrix} a & \begin{bmatrix} 1 & 0 & 1 \\ 0 & 0 & 1 \end{bmatrix} \\ b & \begin{bmatrix} 0 & 1 & 1 \\ 0 & 0 & 1 \end{bmatrix} \end{smallmatrix}$ while the sequence ab removes the 01 state to give $\begin{smallmatrix} a & \begin{bmatrix} 1 & 1 \\ 0 & 1 \end{bmatrix} \\ b & \begin{bmatrix} 0 & 1 \\ 0 & 1 \end{bmatrix} \end{smallmatrix}$. (This matrix viewpoint brings out the duality of events and states mentioned in the introduction better than either the event structure or Boolean formula viewpoint.) Since the columns arise as subsets of A , any two columns of such a matrix must be distinct (whence there can be at most exponentially many more columns than rows); this however is the only restriction on such matrices for the case of general Boolean operations.

We call two formulas $\mathcal{A}, \mathcal{B}$ **isomorphic** when there exists a bijection $g : A \rightarrow B$, understood as a 1-1 renaming of atoms, such that $X = \{y \circ g \mid y \in Y\}$, that is, when they have the same satisfying assignments modulo the renaming.

We conclude this section with a systematic enumeration of the types this paper is primarily about, namely processes, events, states, and values.

At the top level are *processes* $\mathcal{A}, \mathcal{B}, \mathcal{C}$. In a more categorically oriented version of this paper these would form a category furnished with structure via morphisms between

processes. In this version processes are organized into a process algebra via the operations treated in the next section.

The principal constituents of a process are its *events* a, b, c , forming a set A , and its *states* x, y, z forming a set X . Thus far we have viewed states extensionally as sets of events, making A the only primitive set with $X \subseteq 2^A$. However they can also be viewed intensionally as independent entities sibling to events, so that we now have two primitive sets A and X . In this case we no longer have set membership defining the relationship between events and states; instead we provide its intensional counterpart explicitly, as an arbitrary binary relation or matrix $R \subseteq A \times X$ between A and X . Natural notations for $R(a, x)$ are $a.x$ as inner product in a Hilbert space as mathematicians write it, or $\langle a|x \rangle$ as physicists write it, with the value however being not a complex number but a truth value.

The extensional view is the special case of the intensional view with $X \subseteq 2^A$ and $R = \in$. In this paper we forego the appealing symmetry of the intensional view^{††} and stick to the extensional view, whose benefit is that a process is just a pair (A, X) with $X \subseteq 2^A$ instead of a triple (A, R, X) with $R \subseteq A \times X$.

Implicit in the above are the values 0 and 1 taken on by the formulas $a \in x$ and $R(a, x)$. As these values will later be joined by $\sqcap$ and $\times$, we make them more explicit here, collecting them as the set $K = \{0, 1\}$ of *values*. We omit the customary prefix "truth" as these values are schizophrenic, having temporal significance when varying down a column of R and propositional or truth significance when varying across a row.

Summarizing in reverse order, the types are *values*, *events*, *states*, *processes*, and *process operations*.

Besides the above, we have also encountered *labels* $\mathbf{a}, \mathbf{b}, \mathbf{c}$ forming an alphabet Σ . However these only appear with labeled event structures, which we touch on only tangentially in this paper. We also have *logical operations*, which play an important role in expressiveness, permitting every process to be named even when $|K| = 4$ as we shall show later, but which are less natural for programming purposes than the process operations, the topic of the next section.

2.8. Process Algebra

Process algebra (BK84; BW90; BPe00) provides a way of assembling complex processes from simpler ones. This is in contrast to Amir Pnueli's temporal logic (Pnu77), which describes a single complex process from the perspective of a single neutral observer of that process' universe. This distinction has been described by Pnueli as respectively exogenous vs. endogenous specification, constituting the crucial element of temporal logic that distinguishes it from dynamic logic (Pra76; HKT00), a similar modal logic of programs extending temporal logic with the ability to state properties of compositional or algebraic programs.

Temporal logic's neutral-observer viewpoint intrinsically entails an interleaving view

^{††} The extensional view can be symmetrized by forbidding repeated rows ("little chu" in the jargon of Chu spaces (Bar91)), thus allowing events to be viewed as sets of states.

of the progress of the universe. Process algebra restores compositionality of programs in a way that works equally well for interleaving semantics and true concurrency, even when both the observer and observed are distributed as catered for by orthocurrence. Process algebra is a popular subject with a rich literature, many workshops, a comprehensive handbook (BPe00), and many skeptics whose purposes are adequately served by the interleaving viewpoint of concurrency as discussed in Section 1.3.

Before considering the operations for growing bigger processes from smaller, let us consider the smallest processes, those without a plurality of events, of which there are six. The two without any events are (in process algebra notation) 0 and $\bullet$, having respectively one and no states, corresponding respectively to the truth constants (zeroary operations) 1 and 0 . (So without any events at all we already have a naming discrepancy between the notations of process algebra and logic.)

The four processes with one event are, in logical notation, a , $\bar{a}$, 1 , and 0 (the latter two as constant unary operations). Formula a holds only when event a is 1 , and therefore denotes the “stuck-at-one” process which starts in the state in which a has already happened and stays there. Formula $\bar{a}$ is the counterpart for the process that stays in one state in which a has not happened. Formula 1 is the unconstrained process, having both possible states 0 and 1 . Formula 0 is the inconsistent process, having no states.

Of the above four, process algebra notation only offers a name for the third, the unconstrained process. When processes are named only up to isomorphism, the name 1 has the same denotation in process algebra and logic.^{††} In this paper however we will generally not consider processes as being defined only up to isomorphism. Instead we name one-event processes by their one event, since when there are several events a, b, c floating around we will want to keep track of which processes have which event. We therefore call each such process by the name of its one event.

We then have a second naming discrepancy: logically a denotes the already-done process, whereas process algebraically a denotes the unconstrained process.

We now define four basic process algebra connectives: *concurrency* $\mathcal{A} \parallel \mathcal{B}$, *sequence* $\mathcal{A} \mathcal{B}$, *choice* $\mathcal{A} + \mathcal{B}$, and *orthocurrence* $\mathcal{A} \otimes \mathcal{B}$. For all four definitions we assume $\mathcal{A} = (A, X)$ and $\mathcal{B} = (B, Y)$ where $X \subseteq 2^A$ and $Y \subseteq 2^B$. Given our identification of event structures with Boolean operations, this makes process algebra connectives also logical connectives, though not necessarily via the most obvious connections— $\mathcal{A} + \mathcal{B}$ is not $\mathcal{A} \vee \mathcal{B}$ for example, concurrent process algebra does not enjoy the logical simplicity of sequential dynamic logic (Pra76).

We define $\mathcal{A} \wedge \mathcal{B}$ as $(A \cup B, \{z \in 2^{A \cup B} \mid z \restriction A \in X \wedge z \restriction B \in Y\})$, where $z \restriction A$ denotes the restriction of z (whose domain is $A \cup B$) to the subdomain A . One extreme of this is when A and B are disjoint, in which case $\mathcal{A} \wedge \mathcal{B}$ is isomorphic to $(A + B, X \times Y)$, understanding $(x, y)(a) = x(a)$ and $(x, y)(b) = y(b)$, which is coproduct in the category of extensional Chu spaces over 2 . The other extreme is when $A = B$, for which $\mathcal{A} \wedge \mathcal{B}$ simplifies to $(A, X \cap Y)$, ordinary conjunction as intersection of state sets. Hence conjunction is

^{††} Well, almost: the latter requires the additional information that this is the unary constant 1 as opposed to some other arity. This process incidentally is the tensor unit when these processes are organized into a $*$ -autonomous category (Bar79; GP93) to form a model of linear logic (Gir87).

idempotent: $\mathcal{A} \wedge \mathcal{A} = \mathcal{A}$. Similarly we define $\mathcal{A} \vee \mathcal{B}$ as $(A \cup B, \{z \in 2^{A \cup B} \mid z \restriction A \in X \vee z \restriction B \in Y\})$, similarly idempotent. Both $\wedge$ and $\vee$ are clearly also commutative, and almost as clearly associative.

For occasional use in this paper we also define “external” implication $\mathcal{A} \vdash \mathcal{B}$, $\mathcal{A}$ entails $\mathcal{B}$, as simply $X \subseteq Y$. This only makes sense when $A = B$ but we shall only need it for this case in what follows. (There is more than one internal implication, one being $\mathcal{A} \multimap \mathcal{B}$, definable in terms of orthocurrence and perp as $(\mathcal{A} \otimes \mathcal{B}^\perp)^\perp$. It is not needed for this paper however.)

Concurrence $\mathcal{A} \parallel \mathcal{B}$ is defined as $\mathcal{A} \wedge \mathcal{B}$. This makes $\mathcal{A} \parallel \mathcal{B}$ the joint behavior of $\mathcal{A}$ and $\mathcal{B}$: the events of both are performed subject to the constraints of each. Concurrence inherits the commutativity, associativity, and idempotence of $\wedge$; in particular $a \parallel a = a$, a corollary of this being an algebra of unlabeled events. We shall discuss this disconcerting equation after treating the other operations.

In the case $\mathcal{A} = a$, $\mathcal{B} = b$ (viewing a and b as processes, not formulas, the practice we follow for process algebra) there are no constraints on a or b separately and hence $a \parallel b$ is equally unconstrained, i.e. is the constantly true binary operation.

Sequence $\mathcal{A}\mathcal{B}$ is defined as $\mathcal{A} \wedge \mathcal{B} \wedge (B = 0 \vee \check{\mathcal{A}})$. Here $B = 0$ denotes the process $(B, \{0^B\})$ having just the one all-zero state, expressible as $\bigwedge_i \bar{b}_i$ taken over all atoms $b_i \in B$. However $\check{\mathcal{A}}$ is a more delicate notion, which we define for this dyadic case as $(A, \{x \in X \mid \forall z \in X[x \leq z \rightarrow x = z]\})$. So $\check{\mathcal{A}}$ holds at just those states at which $\mathcal{A}$ holds and for which there is no greater (later) state at which $\mathcal{A}$ holds; these are considered the final states of $\mathcal{A}$. $B = 0 \vee \check{\mathcal{A}}$ holds at those states where either $\mathcal{B}$ has not yet started or $\mathcal{A}$ has finished. The meaning of $\mathcal{A}\mathcal{B}$ is therefore $\mathcal{A} \parallel \mathcal{B}$ subject to the additional constraint that $\mathcal{B}$ remain in its initial state until $\mathcal{A}$ has entered a final state.

Whereas $B = 0$ is a local or first-order notion of initiality, $\check{\mathcal{A}}$ is a global or second-order notion of finality necessitated by the limited scope of dyadic logic, $K = \{0, 1\}$. The *cancel* state $\times$ introduced later will permit a simpler and more robust first-order notion of termination.

When $\mathcal{A} = a$ and $\mathcal{B} = b$ the corresponding formula simplifies to $1 \wedge 1 \wedge (\bar{b} \vee a)$ (1 denoting *true* when used in formulas), or $b \rightarrow a$. This suggests the logical interpretation of sequence $\mathcal{A}\mathcal{B}$ as a new internal implication $\mathcal{B} \rightarrow \mathcal{A}$.

Just as $a \parallel a = a$ we also have $aa = a$, unlike the situation with the string **aa** which in the present context is understood as two consecutive actions constituting a labeled event structure with two distinct events. Viewed in this light, the notational conventions of automata theory make the letters of a regular expression such as $(\mathbf{a}(\mathbf{a} + \mathbf{b}))^*$ actions rather than events; a finite regular expression understood as a labeled event structure will always have finitely many actions, but in general infinitely many action instances or events.

Unlike concurrence, the above definition of sequence makes it not idempotent but only transitive ($\mathcal{A}\mathcal{A} \vdash \mathcal{A}$), witness $\mathcal{A} = ab$ which allows $a = 1$, $b = 0$ unlike $abab$ which forces $a = b$. It can be made idempotent by replacing $B = 0$ by $B - A = 0$ in the definition, but this destroys associativity: $a(ba)$ is ab while $a(ba)$ is $a \equiv b$. We will defer the further pursuit of this issue until we add $\times$ to K , permitting a better definition of $\check{\mathcal{A}}$. Sequence is of course not commutative.

Normally concurrence and sequence are defined in terms of the disjoint union or marked sum of their event sets. This is the appropriate combination in the case of labeled event structures where the labels are the primary identifying features and the identities of the individual events fade into the background. The present approach is therefore something of a novelty, especially the resulting idempotence or near-idempotence of these operators, which are customarily thought of as analogous more to addition than to union. The difference is in our focus on events a, b instead of actions $\mathbf{a}, \mathbf{b}$. Whereas multiple references to an event a all denote the same event, multiple references to an action $\mathbf{a}$ are understood to denote distinct instances of that action, i.e. distinct events. Thus whereas aa denotes the single event a preceding itself (vacuous by reflexivity of precedence), $\mathbf{a}\mathbf{a}$ denotes two instances of action $\mathbf{a}$ one preceding the other. This paper restricts attention to events.

Choice $\mathcal{A} + \mathcal{B}$ is defined as

$$A \cup B = 0 \vee (\mathcal{A} \neq 0 \wedge B - A = 0) \vee (\mathcal{B} \neq 0 \wedge A - B = 0).$$

Here $A = 0$ denotes $(A, \{0^A\})$, the process having just the one state 0^A in which all events are zero, while $\mathcal{A} \neq 0$ denotes $(A, X - \{0^A\})$, $\mathcal{A}$ stripped of its zero state if any, “ $\mathcal{A}$ under way.” The states of $\mathcal{A} + \mathcal{B}$ can therefore be partitioned into three disjoint blocks: the initial state in which all events of $A \cup B$ are zero; those states whose restriction to A is a nonzero state of $\mathcal{A}$ and for which those events of B not in A are all zero; and likewise with $\mathcal{A}$ and $\mathcal{B}$ interchanged. The role of the initial state is to permit $\mathcal{A} + \mathcal{B}$ to remain undecided prior to taking its turn. Once it is $\mathcal{A} + \mathcal{B}$ ’s turn, one of $\mathcal{A}$ or $\mathcal{B}$ must be selected and all events absent from the selected process cannot occur.

Substituting a for $\mathcal{A}$ and b for $\mathcal{B}$ in this definition of choice yields $(\bar{a} \wedge \bar{b}) \vee (a \wedge \bar{b}) \vee (b \wedge \bar{a})$, that is, $\neg(a \wedge b)$, expressing conflict $a \# b$.

For A, B disjoint and at least one of $\mathcal{A}$ or $\mathcal{B}$ having the all-zero state, $\mathcal{A} + \mathcal{B}$ coincides with the choice operation $\mathcal{A} \sqcup \mathcal{B}$ defined in (GP93; Gup94), namely as a 2×2 block matrix whose two diagonal blocks are $\mathcal{A}$ and $\mathcal{B}$ and whose off-diagonal blocks are all zero.

At the other extreme, $A = B$, we have $\mathcal{A} + \mathcal{A} = \mathcal{A} + 0$ (where 0 is the process with no events and one state), and if $\mathcal{A}$ has the zero state we have $\mathcal{A} + \mathcal{A} = \mathcal{A}$. So choice is idempotent up to presence of the zero state. By the symmetry of the definition it is commutative. To see that it is associative we may describe $\mathcal{A} + \mathcal{B} + \mathcal{C}$ as having for its nonzero states the nonzero states of exactly one of the three arguments, with all events not in that argument held at zero.

Orthocurrence $\mathcal{A} \otimes \mathcal{B}$ is $(A \times B, \{z \in K^{A \times B} \mid z(\lambda, \forall) \in X \wedge z(\forall, \lambda) \in Y\})$, where $z(\lambda, \forall)$ denotes $\lambda a. z(a, b)$ with b universally quantified in the containing proposition, and dually for $z(\forall, \lambda)$. This makes $z(\lambda, \forall)$ and $z(\forall, \lambda)$ respectively the columns and rows of the $A \times B$ matrix z , making z in effect a “bilinear” form on $A \times B$, analogous to Halmos’ approach to defining tensor product of vector spaces (Hal74, §25). The states of $\mathcal{A} \otimes \mathcal{B}$ may be thought of as all $A \times B$ crosswords that can be formed by taking the states of $\mathcal{A}$ as the “down” dictionary and the states of $\mathcal{B}$ as the “across” dictionary.

For $\mathcal{A} = a, \mathcal{B} = b$, orthocurrence gives an unconstrained process consisting of the single event (a, b) . Any unconstrained one-event process acts as the unit for orthocurrence, and as such corresponds to the linear logic constant 1 .

Orthocurrence is a versatile operation which extends smoothly to ternary ($\mathcal{A} \otimes \mathcal{B} \otimes \mathcal{C}$) and higher arities as well to continuous (real-valued) time (CCMP91). It is associative because the states of $\mathcal{A} \otimes \mathcal{B} \otimes \mathcal{C}$, viewed as “bricks,” are constrained symmetrically by each of X , Y , and Z along the respective axis. It is symmetric in the sense of being commutative up to isomorphism. It is not at all idempotent however.

Orthocurrence can be understood as either interaction $\mathcal{A} \otimes \mathcal{B}$ or observation $\mathcal{A} \multimap \mathcal{B} = (\mathcal{A} \otimes \mathcal{B}^\perp)^\perp$ where $\mathcal{B}^\perp$ is the transpose or dual of $\mathcal{B}$ and $\multimap$ is the previously mentioned internal implication (Pra01). As interaction it describes the collision of systems of particles or events, while as observation it describes the dependence of view on viewpoint: the state $z(a, \lambda)$ of $\mathcal{B}$ in $\mathcal{A} \otimes \mathcal{B}$ (or point or event $z(a, \lambda)$ of $\mathcal{B}$ in $\mathcal{A} \multimap \mathcal{B}$) as seen from $\mathcal{A}$ depends on which point a of A it is seen from. Whereas interaction is an explicitly symmetric notion, observation is an ostensibly asymmetric notion with an implicit symmetry corollary to the linear logic isomorphism $\mathcal{A} \multimap \mathcal{B} \cong \mathcal{B}^\perp \multimap \mathcal{A}^\perp$. Although orthocurrence is not customarily regarded as one of the basic process algebra operations, the author has maintained for many years (Pra85; Pra86; CCMP91) that it should be, long before realizing (GP93) that this created a compelling link between process algebra and linear logic.

The operations of concurrence and orthocurrence make no assumptions about K other than that it is a set. However orthocurrence differs from concurrence (as defined in this paper) in that it makes no structural difference whether the two arguments are disjoint. Sequence depends on state 0 being viewed as initial and state 1 as final, and also involves a second-order notion of finality for process states. Choice depends only on state 0 being initial.

The definitions are for the unlabeled case. We have paid little attention to the labeled case because labeling is a simple process with no impact at all on the portion of process algebra treated here. Labels for the events of processes built with concurrence, sequence, and choice are straightforward because events retain their identity between the inputs and output of those operations: the labels on the events of the result are simply those on the events of the inputs. For orthocurrence, labeling is defined as $\lambda(a, b) = (\lambda(a), \lambda(b))$, namely the label at a pair of events is the pair consisting of their respective labels as provided in the processes input to orthocurrence.

We have not suggested any particular choice of iterative or recursive construct because the topic even of least fixed points, let alone greatest, optimal, etc., of terms built up from processes understood as arbitrary Boolean operations using the above process algebra connectives is sufficiently substantial as to warrant its own paper. Here the reader will need to settle for the account of $\mathbf{a}^*$, $(\mathbf{ab})^*$, and $\mathbf{a} + \mathbf{b}^*$ in Section 3.4 to get at least some idea of what an infinitely generated process should look like, and how it is affected by the additional states $\perp$ and $\times$.

The original Chu semantics took the point of view when combining processes that they had disjoint sets of events, or more precisely that events from distinct processes were unique entities bearing no relationship to each other. The recognition here that processes may share events permits a straightforward account of terms such as $ba + ca$ and $ab + ac$. Such sharing is certainly expressible categorically with pushouts in place of

coproducts, but ordinary union and intersection achieve the same end with less fuss and greater familiarity to most.

These definitions give rise to the semantics shown in Table 1 below for particular formulas in atoms a, b, c , expressed as a matrix whose rows are the atoms and whose columns are the permitted states (the same notation as used for Chu spaces).

Note that this process algebra is an algebra of *unlabeled* event structures, in the sense that its variables range not over actions $\mathbf{a}$ but events a . Thus in expressions such as $ab + ac$ and $ab + ba$ containing more than one occurrence of a variable, all instances of that variable refer not just to the same action but to the same event.

The atomic process a has two states, 0 for not done and 1 for done. The states of $a||b$ are all subsets of $\{a, b\}$. ab forbids the state 01 in which b happened without a , while $a + b$ forbids the state 11 in which both a and b happened. The set of states of $a|(b + c)$ is (isomorphic to) the cartesian product of the state sets of a and $b + c$, six states. $(b + c)a$ forbids the state 100 of $a|(b + c)$ in which a has happened and neither of b or c have happened; this is equivalent to $ba + ca$ and not at all controversial. $a(b + c)$ similarly forbids 010 and 001 and similarly is equivalent to $ab + ac$, which is consistent with trace semantics but not with bisimilarity. And $ab + ba$ is equivalent to $a||b$, consistent with interleaving concurrency but not with true concurrency.

a	$a \begin{bmatrix} 01 \\ 1 \end{bmatrix}$	$ba + ca$	$a \begin{bmatrix} 00101 \\ 01100 \\ c00011 \end{bmatrix}$
$a b$	$a \begin{bmatrix} 0101 \\ 0011 \end{bmatrix}$	$a(b + c)$	$a \begin{bmatrix} 0111 \\ b0010 \\ c0001 \end{bmatrix}$
ab	$a \begin{bmatrix} 011 \\ b001 \end{bmatrix}$	$ab + ac$	$a \begin{bmatrix} 0111 \\ b0010 \\ c0001 \end{bmatrix}$
$a + b$	$a \begin{bmatrix} 010 \\ b001 \end{bmatrix}$	$ab + ba$	$a \begin{bmatrix} 0101 \\ b0011 \end{bmatrix}$
$a (b + c)$	$a \begin{bmatrix} 010101 \\ b001100 \\ c000011 \end{bmatrix}$	$ab \otimes cd$	$ac \begin{bmatrix} 011111 \\ ad000111 \\ bc001011 \\ bd000001 \end{bmatrix}$
$(b + c)a$	$a \begin{bmatrix} 00101 \\ b01100 \\ c00011 \end{bmatrix}$	$(a + b) \otimes (c + d)$	$ac \begin{bmatrix} 0000101 \\ ad0001010 \\ bc0010010 \\ bd0100001 \end{bmatrix}$

Table 1. 2-valued process algebra

The orthocurrence $ab \otimes cd$ describes the flow of ab through cd . Each of a and b passes each of c and d , and these four visits constitute the events of $ab \otimes cd$, respectively (a, c) , (a, d) , (b, c) , and (b, d) . If a and b are trains and c and d stations then this describes two trains moving along a track between two stations. As can be seen from the states, train a arrives at station c first, then a arrives at d concurrently with b visiting c , and finally b arrives at d .

The orthocurrence $(a + b) \otimes (c + d)$ realizes the analogous process with choice in place of sequence. If a and b are pigeons and c and d are pigeonholes, then $a + b$ represents choosing a pigeon (chosen at each hole) while $c + d$ represents the choice of hole (chosen by each pigeon). With these constraints we can either put pigeon a in hole c and pigeon b in hole d , or a in d and b in c , as reflected in the two final states.

A variant of the one-pigeon-one-hole discipline allows multiple pigeons per hole while not permitting the same pigeon in two holes simultaneously. With two pigeons a, b and two holes c, d the only way this can happen is as $(a||b) \otimes (c + d)$, which works out to be

$$\begin{array}{l} ac \\ ad \\ bc \\ bd \end{array} \begin{array}{l} 000010011 \\ 000101100 \\ 001000101 \\ 010001010 \end{array}$$
 The final states are the last four. Pigeon a has to be in either hole c or hole d , and likewise for pigeon b , a total of four possibilities.

We have tacitly capitalized on two-valued logic here. While the final states of $(a||b) \otimes (c + d)$ are as claimed, they are constructed from states of $a||b$ that are not always final (but $c + d$ is kosher). The problem arises in any hole that does not receive two pigeons. The one final state of $a||b$ has both pigeons, which is fine for a hole receiving two pigeons but not one receiving one or no pigeons. We could try to fix this by replacing $a||b$ by $0 + a + b + (a||b)$ so as to allow each hole to receive anywhere from 0 to two pigeons, but when we do so we find that nothing changes. $K = \{0, 1\}$ is too small to appreciate these fine distinctions, allowing us to be sloppy without penalty. The introduction of $\times$ later on will oblige us to count pigeons more carefully.

Three pigeons a, b, c any two of which can occupy the same hole can be described as $(a||b) + (b||c) + (c||a) = \begin{array}{l} a \\ b \\ c \end{array} \begin{array}{l} 0101010 \\ 0011001 \\ 0000111 \end{array}$. But we still can't put a pigeon in two holes d, e at the same time, so we have $((a||b) + (b||c) + (c||a)) \otimes (d + e) =$

$$\begin{array}{l} ad \\ ae \\ bd \\ be \\ cd \\ ce \end{array} \begin{array}{l} 000000100000001111000111 \\ 0000010000011110000111000 \\ 0000100001100010001011001 \\ 0001000110001000100100110 \\ 0010000010100100010101010 \\ 0100000101010001000010101 \end{array}.$$

This process has six final states (the six at the right each with three 1's), corresponding to the three choices of lone pigeon and the two holes it can go in.

Again we have the sloppiness permitted by $K = \{0, 1\}$ that $(a||b) + (b||c) + (c||a)$ is indistinguishable from $a + b + c + (a||b) + (b||c) + (c||a)$. Again the appearance of $\times$ later on will oblige greater care in counting pigeons.

2.9. Higher Dimensional Automata

An ordinary state automaton or transition system can be viewed as an edge-labeled graph. Geometrically the vertices and edges are both cells, of dimension respectively 0 and 1, with the vertices supplying the edges with boundaries, two per edge except in the case of self-loops.

A **higher dimensional automaton** (Pra91) or HDA is an automaton that admits cells of higher dimension, having as their boundaries cells of lower dimension representing the immediately preceding or following more quiescent states. Related prior art includes Papadimitriou's geometric treatment of concurrency control (Pap86), van Glabbeek and Vaandrager's ST-bisimulation (GV87), and the deterministic asynchronous automata of Shields (Shi85). Since its introduction higher dimensional automata have been treated by a number of authors (vG91; GJ92; GC93; Gou93; Gun94; Gou95b; Gou95a; Gou96a; Gou96b; SC96; BJ96; Tak96; SC96; FGR98; Pra00) and have led to two conferences and a special issue of this journal (Gou00).

The basic example of an HDA is the representation of independence $a||b$ as a solid square. The boundary of this square represents the sequential or interleaving interpretation of $a||b$, equivalent to $ab + ba$, while the interior represents the temporally overlapping or independent or concurrent occurrences of a and b .

More generally, a cubic cell of dimension n expresses the simultaneous occurrence of

n events, and its boundaries, as cubes of dimension less than n , represent simultaneous occurrence of some of those events with the remaining events each either not yet started or completed.

In full generality a higher dimensional automaton is a *cubical complex*, namely a set of cubes permitted to share boundaries, creating the possibilities of sequencing, branching, rejoining, and cycles.[†] The case of all cells of dimension at most one is just that of ordinary graphs (with multiple edges permitted between any two vertices). Each cube may be finite- or infinite-dimensional, and there may be infinitely many of them.

The usual notion of transition as an edge of a graph now gives way to that of a transition as any cell of nonzero dimension, with the dimension indicating the number of events participating jointly in that transition. A transition t may be said to run between its boundaries, the precise definition of which we give only for acyclic HDAs, which we now define.

2.10. Acyclic HDAs and Triadic Event Structures

Higher dimensional automata seem at first sight orthogonal to event structures. Indeed at the talk where the author first presented HDAs (Pra91), Boris Trakhtenbrot asked from the front row at question time what relationship these higher dimensional automata had to event structures[‡], for which we had no good answer. More recently we have described (Pra00) what we feel to be the correct connection, via acyclic HDAs and triadic event structures, which the present paper develops further. We sketch the essentials of this connection before proceeding to the main results of this paper.

In outline the reconciliation of HDAs and event structures on A is accomplished by restricting the former while generalizing the latter. HDAs are restricted to their acyclic case by defining them not as a set of cubes allowing sharing of cells but rather as a subset of a single cube 3^A . And event structures are generalized from two-valued to three-valued states, dyadic to triadic, by expanding the set $2 = \{0, 1\}$ of possible values for membership of events in states to $3 = \{0, \sqcup, 1\}$. The new state $\sqcup$ is understood as *transition*, making our logical formulation three-valued, with $\sqcup$ the “edge” in what electrical engineers call edge-triggered logic connecting the “levels” 0 and 1.

An **acyclic HDA** is a pair $\mathcal{A} = (A, X)$ where A is a set of events and X is a subset of 3^A whose elements are understood as the states of $\mathcal{A}$. In particular the cube 3^A itself is the acyclic HDA $(A, 3^A)$ on A realizing the independent or unrestricted occurrence of the events of A . Each A -tuple (configuration, state) x of 3^A is a cell whose dimension is the number of $\sqcup$ s in x and whose boundaries are all tuples in 3^A obtainable from x by setting one or more $\sqcup$ s in x to either 0 or 1. What makes 3^A acyclic is the linear ordering $0 < \sqcup < 1$ on 3, which lifts to the natural partial ordering on 3^A entirely analogously to that on 2^A .

[†] A very elegant characterization beyond the scope of this paper of a cubical complex is as a functor $F : \mathbf{Bip} \rightarrow \mathbf{Set}$ where $\mathbf{Bip}$ is the category of bipointed sets or 0-0 algebras, algebras with two constants and no nonzeroary operations. The category of cubical complexes and their homomorphisms can then be taken to be the functor category $\mathbf{Set}^{\mathbf{Bip}}$, a topos by virtue of being a presheaf category.

[‡] More precisely, to event spaces, a variant of event structures we had recently introduced.

A *triadic event structure* is defined exactly as for an acyclic HDA. Henceforth we use the terms interchangeably, preferring the latter when viewing the process geometrically.

The event a as the proposition “ a has occurred” now has three truth values. If a has not yet started the truth value is 0. If a is ongoing or in transition the truth value is $\top$. And if a has finished its value becomes 1.

The removal of states from 3^A may be understood as furnishing 3^A with a certain structure, namely that structure which is disrespected by precisely the removed states. This is the view of a state as a morphism from A to 3 respecting whatever structure is imputed to A by the omission of certain states. While in many cases such structure may be independently characterized in more conventional ways, e.g. with $\leq$ and $\#$ as with prime coherent event structures, this subsetting or “sculpture” approach to specifying structure has three important benefits: simplicity, generality, and categorical algebra. This last arises by formulating acyclic HDAs as triadic Chu spaces, inheriting their notion of morphism $f : (A, X) \rightarrow (B, Y)$ as a function $f : A \rightarrow B$ that is continuous in the sense that for all $y \in Y$, $y \circ f \in X$ ($- \circ f$ being the inverse-image function f^{-1}). For a more detailed account of this perspective visit <http://chu.stanford.edu/>.

Triadic event structures make it very simple to distinguish $a||b$ from atomic mutual exclusion $ab + ba$. The difference is that only the former admits the state $\top\top$ ($a = \top$, $b = \top$). In all other respects they are the same.

Although this distinction is a central benefit of the intermediate state, there also exist other phenomena not expressible with two states that become expressible with three, for example asymmetric conflict and numeric semaphores.

Asymmetric conflict. Asymmetric conflict is the condition that if a has happened then b cannot happen. Two-valued logic admits only the interpretation that the state $\top\top$ ($a = \top$, $b = \top$) is disallowed, which automatically makes this condition symmetric. However three-valued logic admits the meaning that the state $1\top$ is disallowed. That is, the actual occurrence of b ($b = \top$ as opposed to its completion $b = 1$) is ruled out by the completion of a . If b happens first taking us to 01 , a can still happen afterwards via $\top 1$ then 11 , a route that avoids the forbidden $1\top$.

Strong asymmetric conflict. This is asymmetric conflict with the additional condition of mutual exclusion, i.e. state $\top\top$ is proscribed. This begins to look like sequential composition ba , but differs from it in that a without b is permitted. In fact strong asymmetric conflict is expressible in basic process algebra as $a + ba$ and so is not needed as a primitive, unlike its weaker cousin. Whereas dyadic event structures identify all three of $a + ba$, $ab + ba$, and $a||b$, triadic event structures distinguish all three.

Numeric Semaphores. When n processes are competing for m resources, for example three children taking turns riding two ponies, it is natural to regulate the processes with a semaphore initialized to m . Each process obtains one of the resources from the pool by decrementing the semaphore, and returns it to the pool by incrementing it. The requirement that the semaphore remain nonnegative at all times limits allocation to the available resources.

This system is representable with triadic event structures by the requirement that the number of $\top$ ’s (events in transition) in a state be at most m , or from the HDA perspective that dimension not exceed m . The 3-children-2-ponies example has the three events of

each child taking a turn riding a pony. There are a total of 27 possible configurations, and the limit of two ponies rules out just one of these, the one in which all three children are in the intermediate state of riding a pony. Geometrically we have the surface of a cube, whose dimension is the number of ponies. For want of a pony the interior is lost. We may write this process in 3-2 logic as $\neg(\hat{a} \wedge \hat{b} \wedge \hat{c})$. Pushing the negation down expands this to $\bar{a} \vee a \vee \bar{b} \vee b \vee \bar{c} \vee c$ expressing that the process must confine its activities to the six faces of the cube.

No change is needed to the definitions of the four process algebra operations in order to accommodate this new intermediate state, since only sequence and choice assume anything about K , and they look only for initial and final states; $\lrcorner$ is neither. However the insertion of that additional state permits a reasonable notion of “step” of a computation to be defined as already discussed. And it considerably increases the number of process states for multi-event processes, illustrated in Table 2 with the same examples as before.

a	a $\begin{bmatrix} 0\lrcorner2 \end{bmatrix}$	$ba + ca$	$\begin{bmatrix} a \begin{bmatrix} 000\lrcorner100\lrcorner1 \end{bmatrix} \\ b \begin{bmatrix} 0\lrcorner1110000 \end{bmatrix} \\ c \begin{bmatrix} 00000\lrcorner111 \end{bmatrix} \end{bmatrix}$
$a b$	$\begin{bmatrix} a \begin{bmatrix} 0\lrcorner10\lrcorner10\lrcorner1 \end{bmatrix} \\ b \begin{bmatrix} 000\lrcorner\lrcorner\lrcorner111 \end{bmatrix} \end{bmatrix}$	$a(b + c)$	$\begin{bmatrix} a \begin{bmatrix} 0\lrcorner11111 \end{bmatrix} \\ b \begin{bmatrix} 000\lrcorner100 \end{bmatrix} \\ c \begin{bmatrix} 00000\lrcorner1 \end{bmatrix} \end{bmatrix}$
ab	$\begin{bmatrix} a \begin{bmatrix} 0\lrcorner111 \end{bmatrix} \\ b \begin{bmatrix} 000\lrcorner1 \end{bmatrix} \end{bmatrix}$	$ab + ac$	$\begin{bmatrix} a \begin{bmatrix} 0\lrcorner11111 \end{bmatrix} \\ b \begin{bmatrix} 000\lrcorner100 \end{bmatrix} \\ c \begin{bmatrix} 00000\lrcorner1 \end{bmatrix} \end{bmatrix}$
$a + b$	$\begin{bmatrix} a \begin{bmatrix} 0\lrcorner100 \end{bmatrix} \\ b \begin{bmatrix} 000\lrcorner1 \end{bmatrix} \end{bmatrix}$	$ab + ba$	$\begin{bmatrix} a \begin{bmatrix} 0\lrcorner1010\lrcorner1 \end{bmatrix} \\ b \begin{bmatrix} 000\lrcorner\lrcorner\lrcorner111 \end{bmatrix} \end{bmatrix}$
$a (b + c)$	$\begin{bmatrix} a \begin{bmatrix} 0\lrcorner10\lrcorner10\lrcorner10\lrcorner10\lrcorner1 \end{bmatrix} \\ b \begin{bmatrix} 000\lrcorner\lrcorner\lrcorner\lrcorner111000000 \end{bmatrix} \\ c \begin{bmatrix} 0000000000\lrcorner\lrcorner\lrcorner\lrcorner111 \end{bmatrix} \end{bmatrix}$	$ab \otimes cd$	$\begin{bmatrix} ac \begin{bmatrix} 0\lrcorner111111111111 \end{bmatrix} \\ ad \begin{bmatrix} 00000\lrcorner\lrcorner\lrcorner\lrcorner11111 \end{bmatrix} \\ bc \begin{bmatrix} 000\lrcorner10\lrcorner10\lrcorner111 \end{bmatrix} \\ bd \begin{bmatrix} 000000000000\lrcorner1 \end{bmatrix} \end{bmatrix}$
$(b + c)a$	$\begin{bmatrix} a \begin{bmatrix} 000\lrcorner100\lrcorner1 \end{bmatrix} \\ b \begin{bmatrix} 0\lrcorner1110000 \end{bmatrix} \\ c \begin{bmatrix} 00000\lrcorner111 \end{bmatrix} \end{bmatrix}$	$(a + b) \otimes (c + d)$	$\begin{bmatrix} ac \begin{bmatrix} 0000000\lrcorner10000\lrcorner\lrcorner\lrcorner11 \end{bmatrix} \\ ad \begin{bmatrix} 00000\lrcorner100\lrcorner\lrcorner\lrcorner110000 \end{bmatrix} \\ bc \begin{bmatrix} 000\lrcorner10000\lrcorner1\lrcorner10000 \end{bmatrix} \\ bd \begin{bmatrix} 0\lrcorner10000000000\lrcorner1\lrcorner1 \end{bmatrix} \end{bmatrix}$

Table 2. 3-valued semantics of concurrency

The transitional state distinguishes $a||b$ from $ab + ba$ according respectively to whether or not state $\lrcorner\lrcorner$ is present. We still have $a(b + c) = ab + ac$ however.

The thirteen states of $ab \otimes cd$ constitute the thirteen elements of Allen’s basic interval algebra (All84; Pra00). And the states $0\lrcorner\lrcorner0$ and $\lrcorner00\lrcorner$ of $(a + b) \otimes (c + d)$ represent both pigeons in the process of simultaneously entering their respective holes.

2.11. Step and Run

Triadic event structures permit a notion of computation that is both richer and more refined than for their dyadic counterparts. We define a **step** to be a pair (x, y) of distinct elements of 3^A (not necessarily in X) such that every coordinate of (x, y) is one of 0, 0, 0, $\lrcorner$, $\lrcorner$, $\lrcorner$, $\lrcorner$, 1, or 1, 1, i.e. equal or adjacent in Figure 1(b). (Thus $(00, 0\lrcorner)$, $(00, \lrcorner\lrcorner)$, $(\lrcorner\lrcorner, 11)$, $(0\lrcorner, 01)$, and $(0\lrcorner, \lrcorner1)$ are all steps, but not $(0\lrcorner, 0\lrcorner)$, $(0\lrcorner, 00)$, or $(00, 01)$.) The **step graph** of (A, X) is the graph with vertex set X and edge set all steps (x, y) with $x, y \in X$. A **run** of an acyclic HDA (A, X) is a path in its step graph.

The example steps $(00, \lrcorner\lrcorner)$ and $(0\lrcorner, \lrcorner1)$ suggest the following notions illustrating the richness possible with triadic event structure computations. A step (x, z) (and hence any run containing that step) **bypasses** a state y when (x, y) and (y, z) are both steps. A run **synchronizes** events a and b when $a \equiv b$ at every state of the run; such a synchronizing

run thus bypasses all states in which a and b have “drifted out of synch.” Event a *segues* into event b in a run when that run includes at least one state with $a = \top$, $b = 0$ and another with $a = 1$, $b = \top$ but bypasses all states in which $a = 1$, $b = 0$. Seguing can be thought of as a staggered form of synchronization: the termination of a is synchronized with the starting of b .

Disallowing $(0, 1)$ as a step makes concurrent computational complexity a meaningful notion for triadic event structures. In the dyadic case the transitions were just inclusions and hence closed under composition, in which case only sequential computational complexity makes sense, where one counts the total number of events performed without acknowledging the performance benefit of concurrency. (One could allow credit for independence of participating events, but what is the right notion of independence here, and how should partial independence be weighted?) In the triadic case steps are in general not closed under composition (bypassing provides only “local” composition) and hence more than one step may be needed to accomplish a given run. In particular ab takes at least three steps: a starts, then b starts (with a possible step in between in which a stops before b starts), and finally b stops; in general $a_1a_2 \dots a_n$ takes $n + 1$ steps when all intermediate quiescent states are bypassed and $2n$ steps when they aren’t. However $a_1||a_2|| \dots ||a_n$ can be accomplished in two steps, start and stop, provided they all stay in synch, though it can also be as slow as $a_1a_2 \dots a_n$ when there is no synchronization at all (the run confines itself to the edges of 3^A thereby losing the performance benefit of the higher dimensional cells.)

These notions of step and run can be described very efficiently for all four atomic automata in Figure 1. Treat all of them as reflexive graphs whose non-identity edges are drawn explicitly, oriented to point upwards, and whose identity edges are left undrawn as understood. Do not however view these figures as the Hasse diagram of a poset, i.e. do not take transitive closure. This is unavoidable for 1(a) and 1(c) because there are no paths of length 2 needed to violate transitivity and so they are automatically Hasse diagrams, but for 1(b) and 1(d) it takes two steps to get from 0 to 1.

We can in this way succinctly define the step graph on A to be the irreflexive hull of the product graph formed by multiplying A copies of the reflexive graph K , that is, K^A less the identity steps, which we disallowed only to avoid meaninglessly long runs. Whereas the step graphs produced in this way from the graphs in 1(a) and 1(c) are transitive (since the product of transitive graphs is transitive), those from 1(b) and 1(d) are not transitive and contain shortest-paths of length up to $2|A|$.

A run is then a path in the step graph K^A , and a run of a process is a maximal path in X understood now as the intersection of the step graph with the states of the process. The advantage of this viewpoint is its independence of the details of K , which can be added to as needed to accommodate notions of behavior beyond the scope of the present paper such as exceptions and aborting.

3. Results

In this section we state and prove two elementary results about triadic event structures or acyclic HDAs, concerning respectively the necessity of three values for distinguishing

concurrent from interleaving behavior (their sufficiency being implicit in our discussion of their use), and the basis problem for a concrete language of the associated three-valued logic. We then reassign the transitional state to a different role, that of *cancellation*, as a semantics for branching time. We conclude with a 4-valued logic combining concurrency and branching time, having as its four values *before*, *during*, *after*, and *cancel*, coded as respectively 0, $\sqcup$, 1, and $\times$.

3.1. Necessity of Three-valued Event Structures

The above connection between triadic event structures and acyclic HDAs (Pra00) demonstrates the sufficiency of three values, leaving open whether there might be some other approach that did not require leaving two-valued event structures. We felt intuitively that the passage to three values was not capricious but necessary, but made no attempt to bring that intuition closer to the surface. Here we give a precise sense in which this passage is necessary.

Although a single bit can't express much, sufficiently many bits can express any picture, concept, or fact that can be found in cyberspace or on a hard drive. Since any bit string of length 2^n can be expressed as the truth table of a Boolean formula in n atoms, any negative result about the expressive power of Boolean formulas must necessarily depend on assumptions about how the formulas are used. And if some obstacle is overcome by passing from two values to three, those assumptions must include accounting carefully for every bit or it will be possible to refute the result by allowing two bits (binary digits) to simulate a trit (ternary digit).

We make the following assumptions. First, we assume that each event is associated to a single atom.

This assumption is popularly violated by the encoding of a “long” event a as a pair of “short” events, namely the start and end of a , which admits the possibility that a has started but not yet ended. To avoid the opposite possibility one adds the restriction that if a has ended then it must have (previously) started. This eliminates one of the four assignments or states possible for two atoms, which boils this encoding down to exactly the three-valued case.

Three-valued logic can thus be understood as a two-bit start-end encoding of events in which one of the four configurations of two bits is disallowed. Viewing this situation as three values for one atom instead of two values for each of two atoms has several benefits. It simplifies language by halving the number of Boolean atoms. It simplifies logic by replacing the global axiom start-precedes-end with a cardinality adjustment to the set of truth values. And it exposes the underlying geometry of concurrency by turning disagreement of start and end into dimension of state: dimension 0 for agreement (with *before* the event has neither started nor ended while with *after* the event has both started and ended), dimension 1 for disagreement (with *during* the event has started but not ended).

The complete independence of $a||b$ makes it a concurrent activity. The implicit mutual exclusion of $ab + ba$ is naturally understood as a choice followed by sequential behavior. Event structures would not appear able to draw this distinction.

However the one-atom-one-event assumption has a loophole: if we can't duplicate atoms, let's duplicate the events themselves. Treat the events in the ab choice as being distinct from those in the ba choice and call the latter $b'a'$ to make this distinction explicit. Then $ab + b'a'$ contains four events while its automaton contains five states, $\{\}, \{a\}, \{a, b\}, \{a'\},$ and $\{a', b'\}$.

A drawback of this approach to distinguishing the concurrent behavior of $a||b$ from the interleaving behavior of $ab + ba$ is that labeling now becomes necessary if any connection between a and a' is to be made. Otherwise the claimed distinction is really between $a||b$ and $ab + cd$, a distinction not between concurrency and interleaving but simply between activities performing different tasks.

But labeling used in this way conflicts with the use of labeling to indicate when two unambiguously distinct events perform the same action. For example aa consists of two consecutive occurrences of the action a , whereas the two appearances of a in $ab + ba$ are of the same occurrence in different temporal relationships to b . The labeling mechanism should be reserved for truly distinct occurrences of the same action.

Another drawback is that the automaton cannot forget the order in which ab or ba was made. Ideally there would be a state $\{a, b\}$ reached by performing a and b sequentially in either order and then forgetting the order as being immaterial. Treating as distinct the events associated with each order necessarily creates a permanent record of the choice because the two branches of the choice then arrive at distinct states.

With these concerns in mind we make a second assumption, that the distinction between interleaving and concurrency already arises for the unlabeled case of event structures. With this assumption we can no longer meaningfully clone an event of an event structure as above.

The obvious correspondences famously confuse $a||b$ with $ab + ba$. This however leaves open the possibility that event structures as a two-valued logic of event occurrences have some other interpretation that makes the desired distinction. Our theorem is the stronger result that concurrency-interleaving confusion must occur no matter how event structures are interpreted, subject only to the above two assumptions, that events have only two values, and that identity across choices is maintained extensionally at the event level (as opposed to intensionally with the aid of labels).

The first result of this paper comes in two parts, for prime coherent event structures and for general event structures.

Theorem 8. Prime coherent event structures cannot distinguish interleaving from concurrent behavior.

Proof. It suffices to consider $ab + ba$ and $a||b$ themselves. Both evidently call for symmetric event structures, but there are only two such on two events, independence (no forbidden configurations) and conflict ($\{a, b\}$ forbidden). But we can hardly say a and b are performed concurrently if one of them does not happen, and in any event conflict is assigned a separate meaning. This leaves independence as the only possibility. It can serve to denote either concurrent behavior or interleaving or both, but it cannot distinguish the two notions. $\square$

Now event structures are of limited expressive power compared to general Boolean logic. The question then arises as to whether the latter has sufficient additional room to permit the concurrent behavior $a||b$ to be represented by some formula that is available by virtue of representing no interleaving behavior.

Here a reasonable candidate presents itself, namely $a \equiv b$, which was unavailable with prime coherent event structures. This describes so-called *step* concurrency in which concurrent events occur synchronously, starting and ending together. (This is a different notion of step from the 0-to- $\top$, $\top$ -to-1 notion of step used in this paper since $\top$ does not enter into step concurrency.) We now must confront what we expect of $a||b$. If we are satisfied with step concurrency then since logical formulas cannot confuse synchronized activity with interleaving activity, it follows that concurrent activity cannot be confused with interleaving either.

But if $a||b$ is understood as permitting a to precede b , meaning that there can be a state in which a has happened and b has not, then it must be distinct from $a \equiv b$, as the latter insists on simultaneity, allowing only the state in which neither are done and the state in which both are done.

Without making any commitment to one interpretation or another, we can nevertheless ask whether it is possible to distinguish between the interleaving behavior $ab+ba$ (atomic mutual exclusion), $a \equiv b$ (synchronization), and $a||b$ (independence). We show here that there are not enough dyadic event structures to draw all three distinctions: one of them must be either omitted or identified with one of the other two.

Theorem 9. Dyadic event structures cannot distinguish between all three of interleaving, synchronous, and independent behavior.

Proof. As before we take as our counterexample the case of two events where we wish to draw these distinctions. We have already argued that concurrent behavior must accept the configuration $a = b = 1$ in order that both events happen. Dual reasoning forces acceptance of $a = b = 0$, since without that configuration one of a or b cannot happen during the behavior because “happening” entails passing from state 0 to state 1. This leaves only two symmetric Boolean operations, *true* and $a \equiv b$, not enough for a three-way distinction. $\square$

Synchronous behavior should of course be associated with $a \equiv b$, and the interleaving behavior $ab+ba$ clearly should not. We have the choice of identifying independent behavior with either synchronous behavior or interleaving behavior, but there are not enough event structures to keep all three separate. If it is acceptable to identify concurrent with synchronized behavior then two-valued logic can satisfactorily realize $||$ as $\equiv$. If not then two values are not enough, i.e. three values are necessary, and we have the necessary justification for the 3-2 logic of concurrency.

3.2. A Basis for 3-2 Logic

The passage from two logical values to three is invariably assumed to apply to both the inputs and outputs of logical operations. This follows a long tradition of simplicity

through homogeneity: n -ary mathematical operations on a set K are customarily taken to be of type $K^n \rightarrow K$, even for such operations as division whose type would be more appropriately $K \times K^{\neq 0} \rightarrow K$.

For three-valued logic this tradition suggests taking the appropriate abstract Boolean formulas to be some if not all of the A -ary operations on 3, i.e. functions $3^A \rightarrow 3$. And indeed this has been the standard practice in all explorations of three-valued logic we are aware of, with the semantics of a binary operation such as $a \rightarrow b$ being given as a 3×3 truth table with entries drawn from 3.

However natural this homogeneity may seem, the set 3^{3^A} is considerably larger than needed to describe and reason about concurrent behavior. For that purpose, and arguably for other situations calling for a logic with three-valued atoms, the smaller set 2^{3^A} , i.e. functions $3^A \rightarrow 2$, is adequate. And at least as importantly, it fits perfectly with acyclic HDAs, triadic event structures, and triadic Chu spaces.

Extending the practice of referring to logics whose operations are of type k^{k^A} as k -valued logics, we shall refer to logics whose operations are of type k^{j^A} as j - k -valued logics, or just j - k logics, with $j = 3, k = 2$ for the case at hand.

In practice one works not with the whole of k^{j^A} but rather with just a few basic operations from which the remainder are then obtained by composition. For $j = k = 2$, ordinary Boolean logic, Scheffer stroke or NAND forms a complete basis for A nonempty, though the traditionally preferred basis is $\wedge, \vee, \neg$, or $\wedge, \vee, \neg$ when A is infinite.

In fact this preferred basis remains complete even under the restriction that $\neg$ be applied only to atoms. This follows from the completeness of the unrestricted basis (however proved) by the method of pushing negations down to the atoms with the help of the De Morgan laws $\overline{P \vee Q} = \overline{P} \wedge \overline{Q}$ and its dual.

However the completeness of this restricted case can also be proved *ab initio* by expressing the operation $f : 2^A \rightarrow 2$ as a disjunction $f = \bigvee_{f(t)=1} g_t$ each of whose disjuncts denotes an operation $g_t : 2^A \rightarrow 2$ which is 1 at exactly one A -tuple t of 2^A . Each disjunct g_t is in turn expressed as a conjunction of literals (possibly negated atoms), one for each $a \in A$, taking the sign of the a -th literal to be positive if $t_a = 1$ and otherwise negative. In fact this is an excellent way of proving completeness of the unrestricted basis (and circumvents the De Morgan laws into the bargain.)

An equivalent syntax for literals writes a as $a = 1$ and $\bar{a}$ as $a = 0$. For fixed a and i , the formula $a = i$ is made an element of $2^A \rightarrow 2$ by defining it as $\lambda t. t_a = i$, the input to which is an A -tuple t whose a -th component t_a is to be compared with i . The literals then consist of all such elements of $2^A \rightarrow 2$, forming a set isomorphic to $A \times 2$.

This technique extends immediately to j -2 logic by allowing i to range over j values instead of 2. The set of such literals is then isomorphic to $A \times j$, with the obvious choice of name for the (a, i) -th literal being simply $a = i$. (Although we have been writing 3 for $\{0, \perp, 1\}$, in the present context it is more natural to take j to be $\{0, 1, \dots, j-1\}$.)

Theorem 10. The operations $\wedge, \vee$, and the literals $a = i$ form a complete basis for j -2 logic.

Proof. Every A -ary j -2 operation $f : j^A \rightarrow 2$ is evidently representable as

$$f = \bigvee_{f(t)=1} \bigwedge_{a \in A} a = t_a.$$

□

For any given $t \in j^A$, the function $\bigwedge_{a \in A} a = t_a$ from j^A to 2 is 1 at t and 0 elsewhere. $\bigvee_{f(t)=1}$ enumerates those t at which f is 1. When $j = 2$ this construction forms the complete disjunctive normal form (DNF) of f , where “complete” means that every $a \in A$ appears with the appropriate sign in every needed disjunct. Evidently there is nothing special about the choice of 2 for j in this notion of normal form.

For 3-2 logic it suffices to expand any complete basis L for 2-2 (ordinary Boolean) logic to $L_\perp$ by adjoining the unary operation $\hat{a}$ understood as the literal $a = \perp$ (reverting to $3 = \{0, \perp, 1\}$). $L_\perp$ thus has three forms of literal: negative $\bar{a}$, **transitional** $\hat{a}$, and positive a .

The simplicity of this expansion is one of the benefits of 3-2 logic over homogeneous 3-3 logic, which requires operations other than ordinary conjunction and disjunction in order to generate all three output values.

A more practical basis allows literals to be combined with other Boolean operations besides $\bigwedge$ and $\bigvee$, allowing atomic mutual exclusion $ab + ba$ for example to be expressed as $\neg(\hat{a} \wedge \hat{b})$. Pushing this negation down to the literals blows this formula up to $\bar{a} \vee a \vee \bar{b} \vee b$, a blow-up not encountered with 2-2 logic. This blown-up formula, which geometrically confines the process to the four sides of the $ab + ba$ square, makes clear why atomic mutual exclusion is not distinguishable from $a||b$ using dyadic event structures: it expresses a dyadic tautology but not a triadic one, whereas $a||b$ is both a dyadic and triadic tautology.

Though we shall not need it here, Chu spaces being adequately served by j -2 logic, the above construction for the j -2 case generalizes to j - k with the reinterpretation of $\bigwedge$ and $\bigvee$ as max and min and the addition of suitable constants and a selection operation.

Theorem 11. The operations $\bigwedge$ understood as max, $\bigvee$ as min, selection $m * n$ defined by $m * 1 = m$, $m * n = 0$ for $n \neq 1$, the literals $a = i$, and the constants $0, 1, \dots, k - 1$ form a complete basis for j - k logic.

Proof. Represent f as

$$\bigvee_{f(t)=m} m * \bigwedge_{a \in A} a = t_a.$$

□

3.3. Branching Time: 3-2 Logic for the Unlabeled Case

The basic example  of early branching appears to require event labels to relate the two transitions leaving the start state. Now with $ab + ba$, the two occurrences of a were understood as being the same event in two alternatives, and likewise those of b . In the present situation we have two occurrences of a in $ab + ac$ that we would like to understand as being the *same* event in two alternatives differing only in their choice of whether b or c happens. Making this connection with labels again conflicts with the use of labels to account for the string **aa** as two *distinct* occurrences of the same action **a**, thus entailing two events, the same argument we made when considering concurrency. We were able to

make the two occurrences of a in $ab + ba$ the same event via 3-valued event structures; can we do the same here?

We shall show how this can be done with 3-2 logic. We neglect concurrency for the time being.

We take the three event values to be *before*, *after*, and *cancel*, which we code as respectively 0, 1, and $\times$ as in Figure 1(c), with respective literals $\bar{a}$, a , and $\emptyset$ defining the language $L_\times$ by analogy with $L_\sqcup$. Each event starts out in state 0, and passes to one of 1 or $\times$ depending respectively on whether it happens or is cancelled. Thus 0 is the initial state of an event and $\times$ and 1 are its two final states. Since event structures of this kind are structurally different from acyclic HDAs or triadic event structures we shall call them **cancelling event structures**.

A process is in its initial state when all its events are in the initial state, as with the 2-valued case. However we can now say that it is in a final state when each of its events is in one of the two possible final states. Without the $\times$ state we had no comparably simple and local test of finality: an event could be in state 0 either because it had been cancelled or because it had not yet got around to starting. The best we could do in that case was to declare a state to be final if it was not included in any larger state. Thus one benefit of $\times$ is that it permits a local definition of final state.

The process algebra operations for cancelling event structures are as for dyadic event structures with the following changes.

We first define $\check{\mathcal{A}}$ more locally as $\mathcal{A}$ with the added requirement that every event be either completed (in state 1) or cancelled (in state $\times$). This is equivalent to saying that $\check{\mathcal{A}}$ is obtained from $\mathcal{A}$ by removing all states in which at least one event is 0 (or $\sqcup$ when that is included). We can now simplify $\check{\mathcal{A}}$ in the definition of sequence to $\check{[A]}$, defined as the process on A consisting of all possible final states (of which there 2^A since every event is optionally cancellable). (The square brackets in $\check{[A]}$ indicate that $\check{}$ is applied to each member of its set argument, with the results then conjoined.) This simplification is possible because the definition of sequence independently imposes the constraint of being a state of $\mathcal{A}$. (The second-order notion of final state rules out this simplification, as does the popular method of explicitly labeling every state according to whether it is final as done with traditional automata theory.)

Sequence. Now we could keep our original definition of sequence $\mathcal{A}\mathcal{B}$, modified to $\mathcal{A} \wedge \mathcal{B} \wedge (B = 0 \vee A = \check{})$ with the above trick for termination, but then we still have only transitivity of sequence instead of idempotence. Instead we shall free up $\mathcal{A}\mathcal{B}$ a little in the case that there are any shared events, but not as much as with $B - A = 0 \vee \check{[A]}$.

We define $\mathcal{A}\mathcal{B}$ to be $\mathcal{A} \wedge \mathcal{B} \wedge (B - A = 0 \vee \check{[A]}) \wedge (B = 0 \vee \check{[A - B]})$.

A little algebra shows this to be equivalent to $\mathcal{A} \wedge \mathcal{B} \wedge (B = 0 \vee \check{[A]} \vee (B - A = 0 \wedge \check{[A - B]}))$, i.e. the original definition weakened to liberate events common to $\mathcal{A}$ and $\mathcal{B}$ provided the rest of $\mathcal{A}$ has finished and the rest of $\mathcal{B}$ has not yet started.

For A and B disjoint we have the ordinary notion of sequence. For $A = B$ however we have $\mathcal{A}\mathcal{A} = \mathcal{A} \wedge \mathcal{A} = \mathcal{A}$ (since all events are common and therefore liberated) and therefore we have idempotence. But the symmetry of the definition also overcomes the counterexample $a(ba) \neq (ab)a$ we encountered with only $B - A = 0 \vee \check{[A]}$; both sides

now denote $a \equiv b$. A little more algebra demonstrates associativity (exercise). As always sequence is not at all commutative.

Thus we have the surprising result that sequence behaves conventionally with disjoint arguments (the normal case in automata theory) but nevertheless is an idempotent operation!

The operation is now sufficiently rational to allow us to define a new internal implication which we write as $\mathcal{A} \rightarrow \mathcal{B}$, defined to be the process algebra operation $\mathcal{B}\mathcal{A}$. (To avoid conflict with material (Boolean) implication the latter can be written $\mathcal{A} \supset \mathcal{B}$.)

We did not dare to explore the consequences of this notion based on the definition of sequence we were using for dyadic event structures because of its awkward notion of termination. This implication depends in an essential way on having cancellation as an alternative termination state to 1. However in the absence of conflict, i.e. when termination requires all events to happen, $\mathcal{A}\mathcal{B}$ simplifies to just $\mathcal{A}||\mathcal{B}$ with the additional requirement that every event of A precede every event of B , which can be written

$$\mathcal{A} \wedge \mathcal{B} \wedge \bigwedge_{a \in A, b \in B} (b \rightarrow a),$$

(making this case of $\mathcal{A}\mathcal{B}$ clearly associative). This simpler notion can itself be taken as a novel internal implication $\mathcal{B} \rightarrow \mathcal{A}$ for Boolean logic. The difference between this simpler notion of implication and $\mathcal{A}\mathcal{B}$ in the case of cancelling event structures must be kept in mind when $\mathcal{A}$ contains any conflict $\times$, since then not every event of $\mathcal{A}$ can happen and so in this simplified $\mathcal{B} \rightarrow \mathcal{A}$, $\mathcal{B}$ never gets to start (every satisfying assignment assigns 0 to all of B not in A) yet does get to start in $\mathcal{A}\mathcal{B}$ provided $\mathcal{A}$ has at least one final state. This simpler notion also gives a direction to start out from in the above exercise for associativity of sequence.

Choice. In the dyadic case ($K = \{0, 1\}$) we had defined $\mathcal{A} + \mathcal{B}$ as

$$A \cup B = 0 \vee (\mathcal{A} \neq 0 \wedge B - A = 0) \vee (\mathcal{B} \neq 0 \wedge A - B = 0).$$

The significance of $B - A = 0$ in dyadic choice is that when $\mathcal{A}$ is chosen over $\mathcal{B}$, those events of B not participating in $\mathcal{A}$ can never pass to 1 (or even to $\perp$ when we reintroduce it). The difficulty here is that one cannot deduce this “never” from inspection of individual states: one must consider the whole process.

The $\times$ value lets us express the “never” part of this explicitly in a single state: a cancelled event is one that can never move towards 1. Thus when the process enters a state in which $\mathcal{A}$ is chosen, the events of $\mathcal{B}$ not shared with $\mathcal{A}$ are all explicitly cancelled *in that state* instead of merely being left unstated, and vice versa when $\mathcal{B}$ is chosen. Once cancelled an event stays cancelled because there is no transition from $\times$ to any other value of K in the primitive event automata of Figures 1(c) and 1(d).

We therefore substitute $B - A = \times$ for $B - A = 0$ and $\mathcal{A} \not\leq \times$ for $\mathcal{A} \neq 0$ (to ensure that some event of A starts as soon as the choice is made) in the dyadic definition of choice to yield the definition

$$A \cup B = 0 \vee (\mathcal{A} \not\leq \times \wedge B - A = \times) \vee (\mathcal{B} \not\leq \times \wedge A - B = \times).$$

(However naturally this definition may appear to have flowed from the dyadic defini-

tion, it was suggested to us by P S Thiagarajan as an alternative to the definition we had previously derived in exactly the same way from an equivalent definition of dyadic choice as $A \cup B = 0 \vee (\mathcal{A} \wedge B - A = 0) \vee (\mathcal{B} \wedge A - B = 0)$. When $B - A = \times$ is substituted here for $B - A = 0$ we obtain a notion of choice in which the cancellations are performed *before* starting the selected process. Unfortunately our definition fails associativity, witnessed by the state $\times 00$ which is present in $a + (b + c)$ but absent from $(a + b) + c$. In the Thiagarajan definition the chosen process starts up simultaneously with the cancellation of all rejected events, *including those cancelled in the selected process at its outset*. Hence when c is chosen in $a + (b + c)$, a and b are cancelled simultaneously, precluding state $\times 00$. Viewed more abstractly, the disjuncts in our definition of dyadic choice were not disjoint; substituting $B - A = \times$ for $B - A = 0$ preserves associativity provided the disjuncts are disjoint.)

Other notions of choice. An alternative definition of $\mathcal{A} + \mathcal{B}$ allows the chosen process to start without waiting for the unchosen events to be cancelled. They would eventually be cancelled due to the termination requirement that every event either terminate or be cancelled. This however furnishes $ab + ac$ with a state 100, giving it the possibility of deferring the choice of b or c to the same time $a(b + c)$ chooses. The two processes are nevertheless distinguished by the *possibility* in $ab + ac$ of cancelling an event before a has happened. With this definition we have $\mathcal{A}(\mathcal{B} + \mathcal{C}) \vdash \mathcal{A}\mathcal{B} + \mathcal{A}\mathcal{C}$.

However $\mathcal{A}\mathcal{B} + \mathcal{A}\mathcal{C} \vdash \mathcal{A}(\mathcal{B} + \mathcal{C})$ makes more intuitive sense, meaning that $\mathcal{A}\mathcal{B} + \mathcal{A}\mathcal{C}$ is more constrained (has fewer options) than $\mathcal{A}(\mathcal{B} + \mathcal{C})$. This is achievable as follows. Whenever a run of $\mathcal{A} + \mathcal{B}$ (as defined by Thiagarajan) passes through x and then y , such that a is 0 in x and $\times$ in y , we adjoin to the states of $\mathcal{A} + \mathcal{B}$ the additional state x' derived from x by changing b from 0 to $\times$. This definition satisfies $\mathcal{A}\mathcal{B} + \mathcal{A}\mathcal{C} \vdash \mathcal{A}(\mathcal{B} + \mathcal{C})$ while retaining the crucial distinction that $a(b + c)$ contains 100 and so *may* perform a before cancelling either b or c , unlike $ab + ac$.

All definitions of choice proposed above, buggy or not, are easily seen to be consistent with dyadic choice in the sense that they reduce to it under the identification $\times = 0$.

The Thiagarajan definition of choice being expressible in closed form, unlike the two alternatives proposed above making $\mathcal{A}\mathcal{B} + \mathcal{A}\mathcal{C}$ and $\mathcal{A}(\mathcal{B} + \mathcal{C})$ comparable (at least as we have expressed them), we take it as the meaning of choice in what follows.

The atomic process a by convention happens, i.e. may not be cancelled when standing alone, and therefore by default is permitted only the states 0 and 1. Thus the logical form of a is no longer *true* but instead $\neg \phi$, equivalent to $\bar{a} \vee a$. This also implies that even though $\check{\mathcal{A}}$ now allows $\times$ in final states, $\check{a}$ still means $a = 1$.

Substituting $\mathcal{A} = a$ and $\mathcal{B} = b$ in the preferred definition of choice, $a + b$ translates into logic as

$$(\bar{a} \wedge \bar{b}) \vee (a \not\leq 0 \wedge \not{b}) \vee (b \not\leq 0 \wedge \phi).$$

This asserts that $a + b$ has a zero state and that when it is not in that state one of a or b is cancelled and the other has started (which in the absence of $\lrcorner$ means it is done).

That is, $a + b = \begin{smallmatrix} a & \begin{smallmatrix} 0 & 1 & \times \\ 0 & \times & 1 \end{smallmatrix} \\ b & \end{smallmatrix}$

Event a can be turned into a cancellable event, by adding the empty process $0^\S$ having no events and one state, yielding the process $a + 0$ having one event and all three possible states. For $a + 0$ to take the 0 branch means that a is cancelled and nothing else happens. Hence the equation $\mathcal{A} + \mathcal{A} = \mathcal{A} + 0$ holding for dyadic event structures fails for the triadic case because with $\mathcal{A} + \mathcal{A}$ the two branches have the same set of events and hence no events need cancel (unless they already did so in $\mathcal{A}$). We do however have idempotence $\mathcal{A} + \mathcal{A} = \mathcal{A}$ when $\mathcal{A}$ has the all-zero state.

So states 0 and $\times$ participate in the definitions of both sequence and choice while 1 enters only in sequence. Concurrency and orthocurrence remain independent of any particular elements of K .

With the new atomic state $\times$ our examples change as per Table 3 below.

a	$a \begin{bmatrix} 01 \\ \end{bmatrix}$	$ba + ca$	$\begin{bmatrix} a \begin{bmatrix} 00101 \\ \end{bmatrix} \\ b \begin{bmatrix} 011\times\times \\ \end{bmatrix} \\ c \begin{bmatrix} 0\times\times11 \\ \end{bmatrix} \end{bmatrix}$
$a b$	$\begin{bmatrix} a \begin{bmatrix} 0101 \\ \end{bmatrix} \\ b \begin{bmatrix} 0011 \\ \end{bmatrix} \end{bmatrix}$	$a(b + c)$	$\begin{bmatrix} a \begin{bmatrix} 0111 \\ \end{bmatrix} \\ b \begin{bmatrix} 001\times \\ \end{bmatrix} \\ c \begin{bmatrix} 00\times1 \\ \end{bmatrix} \end{bmatrix}$
ab	$\begin{bmatrix} a \begin{bmatrix} 011 \\ \end{bmatrix} \\ b \begin{bmatrix} 001 \\ \end{bmatrix} \end{bmatrix}$	$ab + ac$	$\begin{bmatrix} a \begin{bmatrix} 01111 \\ \end{bmatrix} \\ b \begin{bmatrix} 001\times\times \\ \end{bmatrix} \\ c \begin{bmatrix} 0\times\times01 \\ \end{bmatrix} \end{bmatrix}$
$a + b$	$\begin{bmatrix} a \begin{bmatrix} 01\times \\ \end{bmatrix} \\ b \begin{bmatrix} 0\times1 \\ \end{bmatrix} \end{bmatrix}$	$ab + ba$	$\begin{bmatrix} a \begin{bmatrix} 0101 \\ \end{bmatrix} \\ b \begin{bmatrix} 0011 \\ \end{bmatrix} \end{bmatrix}$
$a (b + c)$	$\begin{bmatrix} a \begin{bmatrix} 010101 \\ \end{bmatrix} \\ b \begin{bmatrix} 0011\times\times \\ \end{bmatrix} \\ c \begin{bmatrix} 00\times\times11 \\ \end{bmatrix} \end{bmatrix}$	$ab \otimes cd$	$\begin{bmatrix} ac \begin{bmatrix} 011111 \\ \end{bmatrix} \\ ad \begin{bmatrix} 000111 \\ \end{bmatrix} \\ bc \begin{bmatrix} 001011 \\ \end{bmatrix} \\ bd \begin{bmatrix} 000001 \\ \end{bmatrix} \end{bmatrix}$
$(b + c)a$	$\begin{bmatrix} a \begin{bmatrix} 00101 \\ \end{bmatrix} \\ b \begin{bmatrix} 011\times\times \\ \end{bmatrix} \\ c \begin{bmatrix} 0\times\times11 \\ \end{bmatrix} \end{bmatrix}$	$(a + b) \otimes (c + d)$	$\begin{bmatrix} ac \begin{bmatrix} 01\times \\ \end{bmatrix} \\ ad \begin{bmatrix} 0\times1 \\ \end{bmatrix} \\ bc \begin{bmatrix} 0\times1 \\ \end{bmatrix} \\ bd \begin{bmatrix} 01\times \\ \end{bmatrix} \end{bmatrix}$

Table 3. 3-valued semantics of branching time

The process $a + b$ starts out in the 00 start, then one of a or b is cancelled and the other proceeds. An alternative semantics would be to omit the initial 00 state and make the choice of a or b at the outset by starting out in one of two initial states, either $0\times$ or $\times0$. The role of the common initial 00 state is to allow $a + b$ to defer the choice as long as there is other work to be done beforehand, as in $c(a + b)$. It also helps ensure that every process built with these connectives from atoms has a unique initial state.

In $a(b + c)$ we see the significance of the 00 state for $b + c$, which permits a to finish before choosing one of b or c . In $ab + ac$ by contrast, a cannot proceed until the choice of b or c has been made, and in that respect differs from $a(b + c)$.

In the absence of $\lrcorner$, $ab + ba$ remains equal to $a||b$. Choice cannot distinguish them because even though $ab + ba$ must make a choice no event is thereby cancelled. It follows that $\times$ can only distinguish early from late choice when that choice affects *which* events get done. More subtle influences, e.g. on order of events, cannot be distinguished by $\times$. In particular $a(bc + cb)$ and $abc + acb$ remain the same, namely $\begin{bmatrix} a \begin{bmatrix} 01111 \\ \end{bmatrix} \\ b \begin{bmatrix} 00101 \\ \end{bmatrix} \\ c \begin{bmatrix} 00011 \\ \end{bmatrix} \end{bmatrix}$.

The $\times$ state does not enter with any of concurrency, sequence, or orthocurrence, which give the same results as with Table 1.

$(a + b) \otimes (c + d)$ is now described in more detail than in the previous two tables.

[§] 0 denotes the zeroary Boolean operation (constant) tt , having no events and one state; the Boolean constant ff is the logical form of the process $\bullet$ with no events and no states.

In particular we can see that when (a, c) is cancelled, so is (b, d) : no pigeon may start to enter its hole until it has been determined which pigeon-hole pairs have been ruled out. But the Thiagarajan definition of choice calls for the pigeons to enter the holes simultaneously with the cancellations.

The following observation concerning $\check{\mathcal{A}}$, describing $\mathcal{A}$ at its termination (i.e. the deletion from $\mathcal{A}$ of all states in which any event is still 0 without however changing the event set A), allows for a tidier account of pigeons if not trains. This theorem depends on the first-order definition of $\check{\mathcal{A}}$ made possible by $\times$; we have previously seen pigeonhole counterexamples $(a||b$ vs. $0 + a + b + a||b)$ to this theorem with our $K = \{0, 1\}$ definition of $\check{\mathcal{A}}$.

Theorem 12. $\check{(\mathcal{A} \otimes \mathcal{B})} = \check{\mathcal{A}} \otimes \check{\mathcal{B}}$.

Proof. $\check{(\mathcal{A} \otimes \mathcal{B})} \vdash \check{\mathcal{A}} \otimes \check{\mathcal{B}}$ because every state of $\check{(\mathcal{A} \otimes \mathcal{B})}$ is free of zeroes and hence can be built solely from zero-free states of $\mathcal{A}$ and $\mathcal{B}$, i.e. from states of $\check{\mathcal{A}}$ and $\check{\mathcal{B}}$. Conversely $\check{\mathcal{A}} \otimes \check{\mathcal{B}} \vdash \check{(\mathcal{A} \otimes \mathcal{B})}$ because $\check{\mathcal{A}} \otimes \check{\mathcal{B}} \vdash \mathcal{A} \otimes \mathcal{B}$ and no state of $\mathcal{A} \otimes \mathcal{B}$ with a 0 can be produced using only states of $\check{\mathcal{A}}$ and $\check{\mathcal{B}}$. $\square$

In particular $\check{(a + b)} \otimes \check{(c + d)}$ more efficiently describes the two ways of putting two pigeons in two holes. The case of three pigeons and two holes, with up to two pigeons per hole, has six solutions, which we might expect to be described by $\check{((a||b) + (b||c) + (c||a))} \otimes \check{(c + d)}$. The relevant calculations for these two examples are as follows.

$$\begin{array}{c} a \\ \times \\ b \end{array} \begin{array}{c} 1 \\ \times \\ 1 \end{array} \otimes \begin{array}{c} c \\ \times \\ d \end{array} \begin{array}{c} 1 \\ \times \\ 1 \end{array} = \begin{array}{c} ac \\ ad \\ bc \\ bd \end{array} \begin{array}{c} 1 \\ \times \\ 1 \end{array} \quad \begin{array}{c} a \\ \times \\ b \\ \times \\ c \end{array} \begin{array}{c} 1 \\ \times \\ 1 \\ \times \\ 1 \end{array} \otimes \begin{array}{c} d \\ \times \\ e \end{array} \begin{array}{c} 1 \\ \times \\ 1 \end{array} = \begin{array}{c} ad \\ ae \\ bd \\ be \\ cd \\ ce \end{array}$$

While $\check{(a + b)} \otimes \check{(c + d)}$ is as expected, $\check{((a||b) + (b||c) + (c||a))} \otimes \check{(c + d)}$ has no states! The problem is that the new definition of final state means that $((a||b) + (b||c) + (c||a)) \otimes (c + d)$ has no final states. This comes about because in the hole with only one pigeon, we always have a second pigeon in state 0 when finality demands it be in state $\times$. We have been sloppy in writing $(a||b) + (b||c) + (c||a)$ when what we really meant was $a + b + c + (a||b) + (b||c) + (c||a)$, the choice of one pigeon or two. Back when we were not distinguishing 0 from $\times$ there was no difference. Now that we are making the distinction we have to count pigeons more accurately. Adding in the $a + b + c$ possibility gives the expected six solutions, as follows.

$$\begin{array}{c} a \\ \times \\ b \\ \times \\ c \end{array} \begin{array}{c} 1 \\ \times \\ 1 \\ \times \\ 1 \end{array} \otimes \begin{array}{c} d \\ \times \\ e \end{array} \begin{array}{c} 1 \\ \times \\ 1 \end{array} = \begin{array}{c} ad \\ ae \\ bd \\ be \\ cd \\ ce \end{array} \begin{array}{c} 1 \\ \times \\ 1 \\ \times \\ 1 \end{array}$$

A similar calculation shows that the example $(a||b) \otimes (c + d)$ of up to two pigeons in each of two holes must be changed to $(0 + a + b + a||b) \otimes (c + d)$ where 0 is the consistent (one-state) zero-event process.

This points up yet another problem with the intensional (pre-cancellation) definition of final state, one that admittedly only arises via orthocurrence with the four-operator process algebra presented here but that could easily arise when the enthusiastic process algebraist replaces orthocurrence by some other operator more to their liking. The

problem is that a state that should not be final (e.g. because some programmer pact has decreed states of its ilk to be temporarily inconsistent system states) might under the old definition of finality be unjustly promoted to the rank of final state by the disappearance of certain crucial later states that the pact does recognize as consistent. This might fool the process whose turn is next into forging ahead when it is dangerous to do so. The second-order definition of finality can be trickier to administer than the first-order one made possible by explicit cancellation.

On the face of it, we appear to have used the same 3-2 logic for branching time as for concurrency. However there is a crucial difference in how the third value enters. For concurrency we identified a neglected intermediate state through which each event passes on its way from *before* to *after*. This additional state makes the natural relationship between 2 and 3 one of inclusion: 2 is a subset of 3. This is reflected in the passage back to 2-valued semantics of concurrency, namely by *deleting* those states containing any occurrence of $\perp$, the process by which we recover Table 1 from Table 2 (and Table 3 from Table 4, yet to come).

For branching time however we *refined* the *before* state according to whether it was willing to pass to the *after* state. This relationship between 2 and 3 thus makes 2 a quotient of 3 resulting from the identification $\times = 0$. Performing this quotient by replacing $\times$ by 0 turns Table 3 into Table 1 (and Table 4 into Table 2) not by deleting states but by *identifying* those states that previously differed only by having $\times$ in place of 0. One side effect of this quotient is to identify $a(b+c)$ and $ab+ac$, which the $\times$ state nicely distinguishes without however disturbing the generally accepted equality between $(b+c)a$ and $ba+ca$, in both of which the choice is made at the beginning.

Despite this basic difference between these two uses of 3-2 logic, we adopt the same logic for both. The only difference is in how we think of the third literal, whether as transition or cancellation.

3.4. Infinite Processes

In the case of an infinite process, one with infinitely many events, it is reasonable to suppose at least for an effective process that each of its final states have only finitely many events in the 1 state. For example the sequential process $\mathbf{a}^*$ which performs finitely many instances of action $\mathbf{a}$ and then halts is described as the process on the infinite set $A = \{a_0, a_1, a_2, \dots\}$ that changes each of $a_0, a_1, \dots$ in turn from 0 to 1, and then at some finite stage changes the remaining events to $\times$. The allowed nonfinal states are therefore all sequences of the form 1^*0^ω , while the allowed final states are all sequences of the form $1^*\times^\omega$, nonfinal and final being the only two possibilities. Going by our definition of step, any nonfinal state has a step to a final state which changes a finite initial segment of the trailing 0^ω to 1's (the stragglers) and the rest to $\times$. We revisit this example after bringing in $\perp$ to join $\times$, whose principle benefit is a more rational notion of step.

The process $(\mathbf{ab})^*$ is treated as for $\mathbf{a}^*$ with two differences. First, we replace the set of final states with sequences of the form $(11)^*\times^\omega$, i.e. only an even number of 1s is allowed in final states. Second, we now need a labeling function, but this is a triviality: just label

all the events alternately **a** and **b**. From this point of view labeling is always a triviality, all the interesting structure resides in the unlabeled part of the theory.

The process $(\mathbf{a} + \mathbf{b})^*$ has an event set better described as $A = \{a_0, a_1, \dots\} \times \{a, b\} = \{(a_0, a), (a_0, b), (a_1, a), (a_1, b), \dots\}$. Every state has an initial segment of length $k \geq 0$ in which every pair of events $(a_i, a), (a_i, b)$ for $i < k$ are either $1\times$ or $\times 1$ according to whether a or b respectively was chosen at the i -th stage. Immediately after this segment there is an optional $(a_k, a), (a_k, b)$ pair which is either $0\times$ or $\times 0$, indicating that one of a or b has been chosen but has not yet happened. Past that point there are two possibilities: either every pair is 00 ; or, provided every chosen event has happened (no $0\times$ or $\times 0$ pair), every pair is $\times\times$, the latter constituting the final states. These are all the states of $(\mathbf{a} + \mathbf{b})^*$. The operational interpretation should be clear, with the same caveat as before that any state has a single step to a final state, dealt with as before by $\sqcap$. A variant of this process would allow the tail to be all $\times\times$ even in the presence of one $0\times$ or $\times 0$ pair, corresponding to the process deciding it is time to halt without thereby entering a final state (the last chosen event still needs to be done, though without $\sqcap$ the concluding step would have the opportunity to cancel the chosen event instead of performing it).

3.5. Combining Branching Time and Concurrency

We now have two applications of 3-2 logic, one to concurrency and one to branching time. However their third value has a different meaning in each, so we cannot use them together. Instead we combine them by passing to 4-2 logic, in which *before* and *after* are common to all fragments but *during* and *instead-of* are distinguished as separate states originating in respectively the concurrency and branching time worlds. The result is the primitive event automaton of Figure 1(d).

The transitions of our four-state automaton (assumed to be oriented upwards, making 0 the start state) induce transitions between states of a multi-event process coordinate-wise: one state vector can pass to another just when the passage is permitted in each coordinate. Thus there is a transition from 00 to $0\times$, and one from 00 to $\sqcap\times$ expressing the simultaneous stepping of two events one of which is cancelled, but no transition from $0\times$ to $0\sqcap$ because an event may not start once it is cancelled.

It is tempting to allow a transition from $\sqcap$ to $\times$ as well, but the semantics of cancellation should forbid starting in order to avoid unwanted side effects. The passage from $\sqcap$ to an alternative to 1 should instead be to a state called *abort*, a fifth state with the implication that the aborted event may have had some effect before its untimely end. Adding this state to the above setup should require no modification other than to include the abort state among the final states of an event. An alternative here would be to treat it instead as an error state, in which case a state vector would be an error state when at least one event had aborted. We leave this to future treatments.

Theorem 10 ensures that 4-2 logic has a complete basis just as for 3-2 logic, namely the infinitary monotone connectives $\bigwedge$ and $\bigvee$ (and any other Boolean operations that might be convenient to include, which may be viewed as either expansions to the signature or abbreviations) together with the four literals $\bar{a}$, $\hat{a}$, a , and $\not a$ constituting the language $L_{\sqcap\times}$.

In this combined semantics, the definitions of the process algebra connectives given for branching time require no change. The only other change needed is that the atomic event a is defined as $a \begin{smallmatrix} 0 \\ \neg \end{smallmatrix}$ as for the 3-2 logic of concurrency. (Thus a process with an event in state $\neg$ that is not permitted to pass to state 1 cannot terminate; that event is deemed to be in a blocked state. One might suppose it to be instead in a divergent state, but on the reasonable and popular assumption that any atomic event may on closer inspection turn out to be compound (refinement), blocking and divergence must be indistinguishable, to which anyone who has patiently awaited a web page can attest.)

The same behaviors we considered before expand as per Table 4 below.

a	$a \begin{smallmatrix} 0 \\ \neg \end{smallmatrix}$	$ba + ca$	$\begin{smallmatrix} a & 000\neg 100\neg 1 \\ b & 0\neg 111\times\times\times\times \\ c & 0\times\times\times\neg 111 \end{smallmatrix}$
$a b$	$\begin{smallmatrix} a & 0\neg 10\neg 10\neg 1 \\ b & 000\neg\neg\neg 111 \end{smallmatrix}$	$a(b + c)$	$\begin{smallmatrix} a & 0\neg 11111 \\ b & 000\neg 1\times\times\times \\ c & 000\times\times\neg 1 \end{smallmatrix}$
ab	$\begin{smallmatrix} a & 0\neg 111 \\ b & 000\neg 1 \end{smallmatrix}$	$ab + ac$	$\begin{smallmatrix} a & 0\neg 111\neg 111 \\ b & 000\neg 1\times\times\times\times \\ c & 0\times\times\times 00\neg 1 \end{smallmatrix}$
$a + b$	$\begin{smallmatrix} a & 0\neg 1\times\times \\ b & 0\times\times\neg 1 \end{smallmatrix}$	$ab + ba$	$\begin{smallmatrix} a & 0\neg 1010\neg 1 \\ b & 000\neg\neg 111 \end{smallmatrix}$
$a (b + c)$	$\begin{smallmatrix} a & 0\neg 10\neg 10\neg 10\neg 10\neg 1 \\ b & 000\neg\neg\neg 111\times\times\times\times\times \\ c & 000\times\times\times\times\times\neg\neg\neg 111 \end{smallmatrix}$	$ab \otimes cd$	$\begin{smallmatrix} ac & 0\neg 111111111111 \\ ad & 00000\neg\neg\neg 11111 \\ bc & 000\neg 10\neg 10\neg 111 \\ bd & 0000000000\neg 1 \end{smallmatrix}$
$(b + c)a$	$\begin{smallmatrix} a & 000\neg 100\neg 1 \\ b & 0\neg 111\times\times\times\times \\ c & 0\times\times\times\neg 111 \end{smallmatrix}$	$(a + b) \otimes (c + d)$	$\begin{smallmatrix} ac & 0\neg 111\times\times\times\times \\ ad & 0\times\times\times\neg\neg 11 \\ bc & 0\times\times\times\neg 1\neg 1 \\ bd & 0\neg 1\neg 1\times\times\times\times \end{smallmatrix}$

Table 4. 4-valued process algebra

Note that $ab + ba$ coincided with $a||b$ in the case of branching time semantics but not when concurrency semantics was added. We still have $(b + c)a = ba + ca$ as expected, but not $a(b + c) = ab + ac$.

The principle example in this paper of an infinite process, $\mathbf{a}^*$ in Section 3.4, had the property with $\times$ but without $\neg$ that every nonfinal state had a step to a final state in which finitely many stragglers finished and the rest are cancelled. This situation, which clearly does not capture the essence of $\mathbf{a}^*$, would not have arisen in any reasonable sequential model. It arose in our $\neg$ -less concurrent model for lack of any reasonable way to acknowledge the performance contribution of intermediate degrees of independence.

We capitalize on $\neg$ by adding to the states of $\mathbf{a}^*$ a third class of states, those of the form $1^* \neg 0^\omega$ in which one event is in progress. With the more reasonable step semantics we use with $\neg$ it is still possible for any stragglers to make the jump to the 1 state as part of the step from a penultimate to a final state, but every 0 in a penultimate state must change to $\times$ in the final step because the semantics of a step no longer allows a direct leap from 0 to 1. Hence stragglers must have already started by the time of the final jump, but now there is only room for at most one such straggler. Whether there is a straggler depends on whether the nonfinal state from which that last step is taken has zero or one $\neg$'s, i.e. is of dimension zero or one.

3.6. HDAX Geometry

Although the geometric view of an acyclic HDA (A, X) with $X \subseteq 3^A$ understood as a sculpted cube is clear enough, how should a subset of 4^A , or HDAX[¶] as we shall call it, be understood geometrically?

Given an HDAX state x which is $\times$ at a subset $B \subseteq A$ (i.e. all events of B are cancelled in x), we propose to view x as a cancellation-free state of 3^{A-B} . An event that is cancelled in the course of a run thus drops out of sight altogether, projecting the computation as a whole onto a lower-dimensional space by projecting out the axis of the cancelled event. This projection does not change the state of the run at the time of projection but merely removes some of the possible future directions.

Consider for example $(a||b) + a$, which is $\begin{smallmatrix} a & \begin{bmatrix} 0\textcolor{blue}{\neg}10\textcolor{blue}{\neg}10\textcolor{blue}{\neg}1\textcolor{blue}{\neg}1 \\ 000\textcolor{blue}{\neg}\textcolor{blue}{\neg}\textcolor{blue}{\neg}\textcolor{blue}{\neg}111\times\times \end{bmatrix} \\ b & \end{smallmatrix}$. In the beginning its geometric meaning is indistinguishable from that of $a||b$, namely a square. Taking the second option of the choice by stepping from state 00 to state $\textcolor{blue}{\neg}0$ however projects the square onto the a axis, thereby removing b 's opportunity to start. (It is always the future that is projected onto the present, never the past onto the present because one can only cancel an event that has not yet started.) Note that a run of this process can take the step $(\textcolor{blue}{\neg}0, \textcolor{blue}{\neg}\times)$, i.e. the existence of state $0\times$ does not preclude the possibility that this state could be bypassed, allowing the collapse to happen while a is in progress, or even, via the step $(10, 1\times)$, after it terminates. Contrast this with $ab + ac$, in which no run can initiate a without first choosing between b and c . The reason $(a||b) + a$ need not commit to a branch before starting a is that the right-hand alternative a contains no cancellable states!

A run that terminates does so at the final (all-ones) corner of a cube of dimension the number of events that happened in the course of the computation. Normally this number will be finite, even if the computation began with infinitely many dimensions. Structure in the final cube, in the form of omitted cells, constitutes the obstacles the run had to negotiate; with the full cube present the run could have been as short as two steps, but the more typical case will omit many cells and the computation will accordingly have been obliged to take much longer.

This dynamic picture of HDAX behavior could certainly be made static, but we do not see any way of doing this that is sufficiently geometrically intuitive as to recommend it over the dynamic picture.

As for other structures modeled as Chu spaces, HDAX morphisms are obtained simply as the morphisms of **Chu**₄, the category of Chu spaces over a 4-letter alphabet. A Chu morphism can be understood as a continuous function, in that the inverse image of every state of its target is required to be a state of its source, where the inverse image of $y : B \rightarrow K$ under $f : A \rightarrow B$ is defined as the composite $y \circ f : A \rightarrow K$. The category of biextensional Chu spaces, extensional Chu spaces that are also separable (every pair of distinct events is distinguished by some state), over any given K is $*$ -autonomous in the sense of Barr (Bar79), i.e. a self-dual symmetric monoidal closed category and as such a model of linear logic (Gir87).

Counterintuitively **Chu**₄ is defined independently of any structure on $K = 4$, which

¶ Alternative reading: Hitherto, During, After, $\times$.

is taken to be just a set, and yet the Chu morphisms respect the structure on processes implied by the above transitions on $\{0, \sqcup, 1, \times\}$. The key lies in the seemingly clairvoyant ability of $\mathbf{Chu}_K$'s morphisms to respect structure imposed *a posteriori* on K , a phenomenon we have explained elsewhere (Pra99). The secret is that every possible structure on K lifts in a uniform and natural way to structure on objects of $\mathbf{Chu}_K$ in a way that ensures that the (preordained) morphisms respect precisely that structure, no more and no less. An early example of the phenomenon is Lafont and Streicher's observation (LS91) that the Chu morphisms between vector spaces represented as Chu spaces (which they call games) over the *set* (not field) of complex numbers are exactly the linear functions between those vector spaces, despite the lack of any reference to complex arithmetic in the definition of that category of Chu spaces. Slightly less magical (since the Chu morphisms are defined as continuous functions), they also observe that the Chu morphisms between Chu spaces over $\{0, 1\}$ representing topological spaces are exactly the continuous functions between those spaces.

3.7. Comparison with Rodriguez-Anger Branching Time

All four structures of Figure 1, along with two others called partially ordered time and relativistic time, have their counterparts in Rodriguez and Anger's branching time variant (RA01) of Allen interval algebras (All84). The elements of an interval algebra are the possible relationships between the two intervals as they slide past each other on parallel tracks. The relationships, of which there are 13, are traditionally derived from the three binary relations $<$, $=$, and $>$ between two endpoints one from each interval. To analyze the case where the tracks diverge at some point, Rodriguez and Anger apply a fourth (symmetric) binary relation $\parallel$ of incomparability, holding just when the comparison is between points on distinct arms of the diverged tracks, and combine it with the other three relations to obtain the 19 interval relationships possible with diverging tracks.

This variant of interval algebra continues an ongoing study by Rodriguez and Anger of such variants based on the introduction of additional binary relations such as $\parallel$. Referring to the original interval algebra as characterizing linear time, they treat partially ordered time and relativistic time using respectively 4 and 6 binary relations, obtaining respectively 29 and 82 interval relationships (AR91; RA93b; RA93a). Partially ordered time corresponds to tracks that diverge only temporarily, with the same notion $\parallel$ of incomparability as above, while relativistic time is the same with two additional relations describing one endpoint at the point of divergence or convergence with the other on the diverged section, corresponding to the boundary of the relativistic light cone.

Elsewhere (Pra00) we recast the relations $<$, $=$, and $>$ between pairs a, b as atomic states 0 , $\sqcup$, and $\parallel$ of such pairs construed as events to allow the relation algebra techniques typically used to study interval algebras to be replaced by direct application of the orthocurrence operator, thereby correlating the Allen configurations to the 13 states of the orthocurrence $ab \otimes cd$ as in Table 2 above. We then interpreted the additional

^{||} Our earlier HDA papers coded these three values as 0,1,2.

relationships as additional event states and represented the two intervals as processes in such a way that their orthocurrence had respectively the above 29 and 82 states.

In (RA01) Rodriguez and Anger point out that time is symmetric for all three of linear time, partially ordered time, and relativistic time, and argue that this symmetry masks a need to reverse one of the arguments to orthocurrence when applied to interval algebras. To demonstrate this they break the symmetry of partially ordered time by replacing temporary divergence of tracks by permanent divergence (no rejoining), naming the resulting notion of time branching time. Orthocurrence applied as in the other examples now gives the wrong answer, but time-reversing one of the intervals (by exchanging their endpoints and also exchanging $<$ and $>$ in the matrix) does give the correct answer.

While these names for various kinds of time are appropriate within the setting of interval algebras, the question naturally arises as to whether they have the same meaning as used elsewhere. In particular can this interval-algebra notion of branching time distinguish $a(b+c)$ from $ab+ac$?

The role of the interval in the interval algebra analysis of time is to represent sequencing, as evident in the successful use of $ab \otimes cd$ in modelling interval algebra for the three previous notions of time. Now when an interval slides off on a different track from the other interval, the permanent incomparability of its leading endpoint eventually results in permanent incomparability of its trailing endpoint. It follows that if $||$ is used to distinguish $a(b+c)$ from $ab+ac$ by treating it as the state of the unchosen event of $b+c$ in the manner of cancellation, any event following the unchosen one must itself enter state $||$. But this would mean that d in $a(b+c)d$ could not happen. More generally, any process containing a choice would block the first time it made any choice.

This establishes that $||$ cannot be used like $\times$ if it is to pass the litmus test for branching time of distinguishing $a(b+c)$ from $ab+ac$. This does not rule out the possibility of drawing this distinction using $||$ in some other role, e.g. one that views the one instance of a in $a(b+c)$ as being comparable with itself while the two copies of a in $ab+ac$ are incomparable (by virtue of having decided differently between b and c). Such a role would then fully justify the name “branching time” for this variant of interval algebra. We do not see however how to do this in the unlabeled event structure framework as used in the present paper.

This difference between $||$ and $\times$ emphasizes the versatility of orthocurrence, a neutral operator that respects the different meanings of these states to give correct results for processes using either one (when applied correctly of course).

It is worth pointing out an alternative analysis of branching-time interval algebra in which orthocurrence need not reverse one of its arguments, consistent with not doing so for any other application of orthocurrence. The need for reversal arises when the processes ab and cd in $ab \otimes cd$ are treated as having the same orientation, following nearly two decades of interval algebra tradition. However inspection of the trains-and-stations scenario reveals that if ab is an eastbound system of trains, placing train a to the east of train b , then stations must be encountered in the order west to east, placing station c west of station d in cd .

To capture divergence we can replace the stationary stations c, d by a westbound pair of trains c, d moving along a second parallel track, with train c ahead of and hence west

of train d . Assuming the tracks diverge when traveling east, the westbound trains reverse the order in which $||$ is encountered: instead of blocking they unblock. The proper analysis of orthocurrence $\mathcal{A} \otimes \mathcal{B}$ in this case is then seen to be that $\mathcal{B}$ is simply a different system of trains from $\mathcal{A}$, experiencing scenery isomorphic to $\mathcal{A}$'s but in the reverse order. In this analysis the need to reverse one argument to orthocurrence becomes a property not of orthocurrence but of its constituent processes, much as the abstract group theoretic analysis of the integers views subtraction $i - j$ as merely a convenient shorthand for $i + (-j)$. Both notations are legitimate and interchangeable.

4. Concluding Reflections

Our earlier aspersions on homogeneity in type notwithstanding, there may possibly be some merit in the homogeneous formulation of 3-valued logic, i.e. the type $3^A \rightarrow 3$, for either concurrency or branching time alone, or conceivably even the homogeneous formulation $4^A \rightarrow 4$ of 4-valued logic when concurrency and branching time are combined (but this starts to look unwieldy). There might for example be some justification for expanding the number of information distinctions to match that of the number of temporal distinctions. Scott domains or information systems leap immediately to mind here, specifically **Bool** consisting of the two truth values tt and ff together with $\perp$ satisfying $\perp \sqsubseteq tt$ and $\perp \sqsubseteq ff$. Standardly interpreted, information systems live on the automaton side of the schedule-automaton or time-information duality, in that their elements are conventionally understood as states ordered by information content via $\sqsubseteq$, with the bottom element $\perp$ constituting the state of utter ignorance.

A domain-oriented model would allow three states of knowledge about HDA cells: present, absent, and undecided, corresponding to the values tt , ff , and $\perp$ of the 3-element CPO **Bool**. Such a model would lend itself to the treatment of *discovery* of HDAs, that is, the automatic production of concurrent automata via a suitable recursive procedure. The bottom state is total ignorance about all cells. The maximal states are the HDAs towards which such a procedure would progress starting from ignorance.

Is the homogeneity of the resulting $3^A \rightarrow 3$ type structural or merely superficial? In particular does 3 have a natural structure here counterpart to the rich Boolean structure of 2? And can the three temporal values *before*, *during*, and *after* be linked in some way to the three information values tt , ff , and *undecided*, e.g. in terms of triadic Chu spaces?

Elsewhere (CCMP91) we pointed out that there are just two commutative idempotent 3-element quantales, each corresponding to a natural use for 3-valued time, with the one that is a Heyting algebra, which we called **3**, being the appropriate one for causal-accidental time and the other, **3'**, the one for strict-nonstrict or before-during-after time. Is **3'** the appropriate quantale for the 3-valued CPO for the type **Bool**? The only connection we see is that the so-called face “lattice” of **3'** as a 1-cell (made an actual lattice by adjoining a bottom element of dimension -1) is the order dual of the CPO **Bool**. However we do not see any good structural connection between face lattice structure and quantale structure, suggesting that cardinality may be the only connection between the two threes in this approach to making sense of $3^A \rightarrow 3$. We leave this homogeneity question as an open problem possibly bearing further investigation.

A far more important problem is the comprehensive treatment of fixed points, solutions of $\mathcal{A} = \tau\mathcal{A}$ (as well as prefixed points $\tau\mathcal{A} \vdash \mathcal{A}$ and postfix points $\mathcal{A} \vdash \tau\mathcal{A}$), whether least, greatest, optimal, etc. Here the natural choice of language to start out with for τ is (in our biased view of course) the four basic process algebra operations defined in this paper. But process algebra operations are limited only by the imagination of process algebraists and the tolerance of users for languages allowed to grow like Topsy, leaving the subject of fixpoints for HDAXs dauntingly wide open.

References

- J.F. Allen. Towards a general theory of action and time. *Artificial Intelligence*, 23:123–154, 1984.
- F. D. Anger and R. V. Rodriguez. Time, tense, and relativity revisited. In B. Bouchon-Meunier, R. R. Yager, and L. A. Zadeh, editors, *Uncertainty in Knowledge Bases: Proc. of the 3rd International Conference on Information Processing and Management of Uncertainty in Knowledge-Based Systems, IPMU'90*, pages 286–295. Springer, Berlin, Heidelberg, 1991.
- M. Barr. **-Autonomous categories*, volume 752 of *Lecture Notes in Mathematics*. Springer-Verlag, 1979.
- M. Barr. **-Autonomous categories and linear logic*. *Math Structures in Comp. Sci.*, 1(2):159–178, 1991.
- S.D. Brookes, C.A.R. Hoare, and A.D. Roscoe. A theory of communicating sequential processes. *Journal of the ACM*, 31(3):560–599, 1984.
- R. Buckland and M. Johnson. Echidna: A system for manipulating explicit choice higher dimensional automata. In *AMAST'96: Fifth Int. Conf. on Algebraic Methodology and Software Technology*, Munich, 1996.
- J.A. Bergstra and J.W. Klop. Process algebra for synchronous communication. *Information and Control*, 60:109–137, 1984.
- J.A. Bergstra, A. Ponse, and S.A. Smolka (editors). *Handbook of Process Algebra*. Elsevier (North-Holland), 2000.
- J.C.M. Baeten and W.P. Weijland. *Process Algebra*. Cambridge University Press, 1990.
- R.T. Casley, R.F. Crew, J. Meseguer, and V.R. Pratt. Temporal structures. *Math. Structures in Comp. Sci.*, 1(2):179–213, July 1991.
- L. Fajstrup, E. Goubault, and M. Raussen. Detecting deadlocks in concurrent systems. In *Proc. of CONCUR'98*, volume 1466 of *Lecture Notes in Computer Science*, pages 332–347. Springer-Verlag, 1998.
- E. Goubault and R. Cridlig. Semantics and analysis of Linda-based languages. In *Proc. 3rd Int. Workshop on Static Analysis*, volume 724 of *Lecture Notes in Computer Science*, pages 72–86, Padova, 1993. Springer-Verlag.
- J.-Y. Girard. Linear logic. *Theoretical Computer Science*, 50:1–102, 1987.
- J.L. Gischer. The equational theory of pomsets. *Theoretical Computer Science*, 61:199–224, 1988.
- E. Goubault and T.P. Jensen. Homology of higher dimensional automata. In *Proc. of CONCUR'92*, volume 630 of *Lecture Notes in Computer Science*, pages 254–268, Stonybrook, New York, August 1992. Springer-Verlag.
- E. Goubault. Homology of higher-dimensional automata. In *Proc. of CONCUR'93*, volume 630 of *Lecture Notes in Computer Science*, pages 254–268, Stonybrook, New York, August 1993. Springer-Verlag.

- E. Goubault. *The Geometry of Concurrency*. PhD thesis, École Normale Supérieure, 1995.
- E. Goubault. Schedulers as abstract interpretations of hda. In *Proc. of PEPM'95*, La Jolla, June 1995. ACM Press.
- E. Goubault. Durations for truly-concurrent actions. In *Proceedings of ESOP'96*, pages 173–187. Springer-Verlag, 1996.
- E. Goubault. A semantic view on distributed computability and complexity. In *Proceedings of the 3rd Theory and Formal Methods Section Workshop*. Imperial College Press, 1996.
- E. Goubault (ed.). Geometry and concurrency. *Mathematical Structures in Computer Science, special issue*, 10(4):409–573 (7 papers), August 2000.
- V. Gupta and V.R. Pratt. Gates accept concurrent behavior. In *Proc. 34th Ann. IEEE Symp. on Foundations of Comp. Sci.*, pages 62–71, November 1993.
- J. Grabowski. On partial languages. *Fundamenta Informaticae*, IV.2:427–498, 1981.
- H.J. Genrich and P.S. Thiagarajan. A theory of bipolar synchronization schemes. *Theoretical Computer Science*, 30:241–318, 1984.
- J. Gunawardena. Homotopy and concurrency. *EATCS Bulletin* 54, pages 184–193, October 1994.
- V. Gupta. *Chu Spaces: A Model of Concurrency*. PhD thesis, Stanford University, September 1994. Tech. Report, available as <http://boole.stanford.edu/pub/gupthes.ps.gz>.
- R.J. van Glabbeek and F.W. Vaandrager. Petri net models for algebraic theories of concurrency. In *Proc. PARLE, II, LNCS 259*, pages 224–242. Springer-Verlag, 1987.
- P.R. Halmos. *Finite-Dimensional Vector Spaces*. Springer-Verlag, 1974.
- D. Harel, D. Kozen, and J. Tiuryn. *Dynamic Logic*. MIT Press, Boston, 2000.
- Y. Lafont and T. Streicher. Games semantics for linear logic. In *Proc. 6th Annual IEEE Symp. on Logic in Computer Science*, pages 43–49, Amsterdam, July 1991.
- A. Mazurkiewicz. Concurrent program schemes and their interpretations. Technical Report DAIMI Report PB-78, Aarhus University, Aarhus, 1977.
- R. Milner. *A Calculus of Communicating Systems*, volume 92 of *Lecture Notes in Computer Science*. Springer-Verlag, 1980.
- M. Nielsen, G. Plotkin, and G. Winskel. Petri nets, event structures, and domains, part I. *Theoretical Computer Science*, 13:85–108, 1981.
- C. Papadimitriou. *The Theory of Database Concurrency Control*. Computer Science Press, 1986.
- D. Park. Concurrency and automata on infinite sequences. In *Proc. Theoretical Computer Science*, volume 104 of *Lecture Notes in Computer Science*, pages 167–183. Springer-Verlag, 1981.
- C.A. Petri. Fundamentals of a theory of asynchronous information flow. In *Proc. IFIP Congress 62*, pages 386–390, Munich, 1962. North-Holland, Amsterdam.
- A. Pnueli. The temporal logic of programs. In *18th IEEE Symposium on Foundations of Computer Science*, pages 46–57, October 1977.
- G.D. Plotkin and V.R. Pratt. Teams can see pomsets. In *Proc. Workshop on Partial Order Models in Verification*. AMS, December 1996.
- V.R. Pratt. Semantical considerations on Floyd-Hoare logic. In *Proc. 17th Ann. IEEE Symp. on Foundations of Comp. Sci.*, pages 109–121, October 1976.
- V.R. Pratt. On the composition of processes. In *Proceedings of the Ninth Annual ACM Symposium on Principles of Programming Languages*, January 1982.
- V.R. Pratt. Some constructions for order-theoretic models of concurrency. In *Proc. Conf. on Logics of Programs*, volume 193 of *Lecture Notes in Computer Science*, pages 269–283, Brooklyn, 1985. Springer-Verlag.

- V.R. Pratt. Modeling concurrency with partial orders. *Int. J. of Parallel Programming*, 15(1):33–71, February 1986.
- V.R. Pratt. Modeling concurrency with geometry. In *Proc. 18th Ann. ACM Symposium on Principles of Programming Languages*, pages 311–322, January 1991.
- V.R. Pratt. Chu spaces: Notes for school on category theory and applications. Technical report, University of Coimbra, Coimbra, Portugal, July 1999. Manuscript available as <http://boole.stanford.edu/pub/coimbra.ps.gz>.
- V.R. Pratt. Higher dimensional automata revisited. *Math. Structures in Comp. Sci.*, 10:525–548, 2000.
- V.R. Pratt. Orthocurrence as both interaction and observation. In *Proc. Workshop on Spatial and Temporal Reasoning* (ed. R. Rodriguez and F. Anger), *IJCAI'01*, Seattle, August 2001.
- R. V. Rodriguez and F. D. Anger. An analysis of the temporal relations of intervals on relativistic space-time. In B. Bouchon-Meunier, L. Valverde, and R. R. Yager, editors, *IPMU'92: Advanced Methods in Artificial Intelligence - Proc. of the 4th International Conference on Information Processing and Management of Uncertainty in Knowledge-Based Systems*, pages 139–148. Springer, Berlin, Heidelberg, 1993.
- R. V. Rodriguez and F. D. Anger. Constraint propagation + relativistic time = more reliable concurrent programs. In *Proc. of the Sixth International Conference on Industrial and Engineering Applications of Artificial Intelligence and Expert Systems IEA/AIE-93*, pages 236–239, Edinburgh, Scotland, 1993.
- R. V. Rodriguez and F. D. Anger. Branching time via Chu spaces. In *Proc. Workshop on Spatial and Temporal Reasoning* (ed. R. Rodriguez and F. Anger), *IJCAI'01*, Seattle, August 2001.
- V. Sassone and G. L. Cattani. Higher-dimensional transition systems. In *Proceedings of LICS'96*, 1996.
- M. Shields. Deterministic asynchronous automata. In E.J. Neuhold and G. Chroust, editors, *Formal Models in Programming*. Elsevier Science Publishers, B.V. (North Holland), 1985.
- Y. Takayama. Extraction of concurrent processes from higher-dimensional automata. In *Proceedings of CAAP'96*, pages 72–85, 1996.
- R. van Glabbeek. Bisimulations for higher dimensional automata. Manuscript available as <http://theory.stanford.edu/~rvg/hda>, June 1991.
- R. van Glabbeek and G. Plotkin. Configuration structures. In *Logic in Computer Science*, pages 199–209. IEEE Computer Society, June 1995.
- G. Winskel. *Events in Computation*. PhD thesis, Dept. of Computer Science, University of Edinburgh, 1980.
- G. Winskel. Event structure semantics for CCS and related languages. In *Proc. 9th ICALP*, volume 140 of *Lecture Notes in Computer Science*, pages 561–576. Springer-Verlag, 1982.
- G. Winskel. Event structures. In *Petri Nets: Applications and Relationships to Other Models of Concurrency, Advances in Petri Nets 1986*, volume 255 of *Lecture Notes in Computer Science*, Bad-Honnef, September 1986. Springer-Verlag.
- G. Winskel. A category of labelled Petri nets and compositional proof system. In *Proc. 3rd Annual Symposium on Logic in Computer Science*, Edinburgh, 1988. Computer Society Press.
- G. Winskel. An introduction to event structures. In *Linear Time, Branching Time and Partial Order in Logics and Models for Concurrency, REX'88*, volume 354 of *Lecture Notes in Computer Science*, Noordwijkerhout, June 1988. Springer-Verlag.