

Distributed Computing Column 69

Proving PACELC and Concurrent Computing Summer School

Jennifer L. Welch
Department of Computer Science and Engineering
Texas A&M University, College Station, TX 77843-3112, USA
`welch@cse.tamu.edu`



The current column is devoted to concurrency from a pedagogical perspective. The first contribution, from Wojciech Golab, revisits Brewer’s CAP principle (“consistency (C), availability (A), and partition-tolerance (P) are not simultaneously achievable”) for large-scale distributed systems and Abadi’s extension called PACELC (“if P then A or C, else L or C”). The paper presents an alternative proof of the CAP theorem using a new latency lower bound for the asynchronous model, and then provides a formal interpretation of PACELC.

The second contribution, from Petr Kuznetsov, is a review of the first summer school on Practice and Theory of Concurrent Computing, which was held in summer 2017 in Saint Petersburg, Russia. The article provides overviews of the topics covered in the summer school and ends with some first-hand evaluations from teachers and students.

Many thanks to Wojciech and Petr for their contributions!

Call for contributions: I welcome suggestions for material to include in this column, including news, reviews, open problems, tutorials and surveys, either exposing the community to new and interesting topics, or providing new insight on well-studied topics by organizing them in new ways.

Proving PACELC

Wojciech Golab
University of Waterloo
Canada
wgolab@uwaterloo.ca



Abstract

Scalable distributed systems face inherent trade-offs arising from the relatively high cost of exchanging information between computing nodes. Brewer's CAP (Consistency, Availability, Partition-Tolerance) principle states that when communication becomes impossible between isolated parts of the system (i.e., the network is partitioned), then the system must give up either safety (i.e., sometimes return an incorrect result) or liveness (i.e., sometimes fail to produce a result). Abadi generalized Brewer's principle by defining the PACELC (if Partition then Availability or Consistency, Else Latency or Consistency) formulation, which captures the observation that the trade-off between safety and liveness is often made in practice even while the network is reliable. Building on Gilbert and Lynch's formal proof of the CAP principle, this paper presents a formal treatment of Abadi's formulation and connects this result to a body of prior work on latency bounds for distributed objects.

1 Introduction

Designers and implementers of distributed systems suffer many headaches over failures, concurrency, and also physical limits related to the exchange of information between components. Notably, the propagation delay for communication between data centers grows linearly with distance due to the finite propagation speed of light, which makes it difficult to build systems that are scalable in a geographic sense, and yet maintain low latency for practical workloads. To make matters worse, failures of wide-area network links can partition a system into components that can communicate internally but cannot communicate with each other. In this scenario, many systems are unable to fulfill their goals, which can be categorized broadly as ensuring *safety* (e.g., never producing incorrect outputs) and *liveness* (e.g., producing outputs eventually for all requests). Brewer's *CAP principle* summarizes this observation by stating that the combination of Consistency (safety), Availability (liveness), and Partition tolerance are not achievable simultaneously.

Following Brewer's keynote speech at PODC 2000 [3], the CAP Principle became the subject of lively discussion, raising questions such as how to define consistency and availability precisely,

and how to categorize existing systems according to which correctness property they sacrifice to overcome the conjectured impossibility. This has led to some confusion, for example the “two out of three” interpretation of CAP, which treats C, A and P symmetrically and suggests that every system can be classified as AP, CP, or CA. In fact some systems (e.g., Apache Cassandra) can be tuned to provide either AP, CP, or none of the above. Moreover, the interpretation of CA (i.e., consistent and available but not partition tolerant) is questionable because lacking P seems to imply that either C or A is lost in the event of a partition, unless perhaps the system is not distributed to begin with, in which case it tolerates partitions trivially. Abadi re-visited the CAP principle by raising two technical points in his 2012 article [1]: (i) no trade-off is necessary at times when the network is reliable, meaning that an AP or CP system may in fact provide both C and A most of the time; and (ii) many practical systems sacrifice C to reduce latency (L) irrespective of network failures. This observation is known as Abadi’s PACELC (“pass-elk”) formulation: if Partition then Availability or Consistency, Else Latency or Consistency. This formulation distinguishes P from A and C, thus discouraging the “two out of three” interpretation, and also separates the inherent C-A trade-off during a network partition from the voluntary L-C trade-off exercised during failure-free operation.¹

In parallel with efforts by practitioners to finesse the interpretation of CAP and related trade-offs, the theory community has sought to formalize these observations as mathematical facts. Two years following Brewer’s keynote, Gilbert and Lynch [5] brought rigor to the discussion by formalizing CAP as the impossibility of simulating a read/write register in a message passing system that simultaneously guarantees Lamport’s atomicity property [7] (consistency) and eventual termination (availability) in the presence of a network partition (arbitrary message loss). This result is commonly referred to as the *CAP theorem*, and is distinguished from Brewer’s informal and intuitively appealing conjecture. Naturally, the proof of the CAP theorem also validates PAC, or the first half of Abadi’s PACELC formalism.

Building on the formal model adopted by Gilbert and Lynch [5], this paper aims to present a rigorous treatment of Abadi’s PACELC formulation by applying an alternative proof technique based on latency bounds for shared objects. Specifically, the paper makes the following contributions:

- Section 3 discusses known latency bounds for shared objects in partly synchronous systems [8, 2], and proves an analogous bound for the asynchronous model.
- Section 4 establishes a connection between latency bounds for shared objects and CAP-related trade-offs by using the lower bound established in Section 3 to derive an alternative proof of the CAP theorem.
- Section 5 states a formal interpretation of Abadi’s PACELC formulation in terms of the results presented in Sections 3 and 4.

2 Formal Model

I/O automata and their composition Similarly to [5], this paper adopts the asynchronous system model formalized by Lynch in Chapter 8 of [9]. There are n reliable processes that communicate using point-to-point FIFO (first-in first-out) communication channels. Two varieties of such

¹ Abadi explains that latency is “arguably the same thing” as availability since a network that loses messages is indistinguishable from one that delays message delivery indefinitely [1]. Thus, L is in some sense synonymous with A.

channels are considered in different parts of this paper: reliable channels that may delay messages but guarantee eventual delivery, and unreliable channels that may drop message entirely. Both processes and channels are modeled as *I/O (input/output) automata*, and their composition is an automaton A representing a system that simulates a single read/write register. Process automata are denoted $P_1, \dots, P_n$, and the channel automaton by which P_i sends messages to P_j , $i \neq j$ is denoted $C_{i,j}$. Processes interact with channels using *send* and *receive* actions on messages. Processes also support two types of output actions, *invoke* and *respond*, by which they initiate and (respectively) compute the result of an operation on the simulated read/write register.

Executions and traces The behavior of the system in a given run is modeled as an *execution*, which is an alternating sequence of states and actions, beginning with a start state determined by the initial value of the simulated register. A *trace* corresponding to an execution α , denoted $trace(\alpha)$, is the subsequence of *invoke* and *respond* actions (i.e., external actions) in α .² The projection of an execution α (respectively trace β) onto the actions of a particular process P_i is denoted $\alpha|P_i$ (respectively $\beta|P_i$). We assume that every execution is *well-formed*, meaning that each process executes an alternating sequence of *invoke* and *respond* actions, starting with *invoke*. The *invoke* action is enabled for each process in every start state, and also in every state where the last output action of the process was a *respond*.

Fairness and timing An execution is *fair* if every process or channel automaton that is enabled to execute an action eventually either executes this action or ceases to be enabled to execute it. In this context, fairness means that every message sent is eventually either dropped or received, and every process eventually invokes another read or write operation if it has computed the response of its previous operation. Thus, a fair execution may in principle have a finite trace if the protocol becomes stuck with no actions enabled. Executions are *timed*, meaning that each event (occurrence of an action) is tagged with a timestamp from a global clock.³ This makes it possible quantify the time taken for a channel to deliver a message in an execution (time of *receive* minus time of *send*, or else ∞ if *receive* does not occur), or the latency of a read or write operation (time of *respond* minus time of *invoke*, or else ∞ if *respond* does not occur).

Correctness properties An execution α of the system automaton A is *consistent* if $trace(\alpha)$ satisfies Lamport's *atomicity* property for read/write registers [7] (a special case of Herlihy and Wing's *linearizability* property [6]), whose formalization is discussed in Chapter 13 of [9]. Quoting [5], atomicity is explained informally as follows:

Under this consistency guarantee, there must exist a total order on all operations such that each operation looks as if it were completed at a single instant.

For the impossibility results presented in this paper, it suffices to adopt a weaker notion of consistency based on Lamport's *regularity* property, which is easier to formalize. In the single-writer

² The *send* action is an output action of each process automaton, and an internal action of the composed automaton A . This is accomplished by hiding *send* actions in A .

³ The global clock is introduced to simplify analysis, and in this version of the model processes do not have access to the clock.

case, it states that a read must return either the value assigned by the last write preceding⁴ it in the execution, or the value assigned by some write that is concurrent⁵ with the read.

An execution α of the system automaton A is *available* if for every process P_i , any invocation action of P_i is eventually followed by a *respond* action of P_i (i.e., every operation invoked eventually produces a response).

An execution α of the system automaton A is *partition-free* if for every message m sent, the *send* action for m is eventually followed by a *receive* action for m (i.e., all messages sent are delivered eventually).⁶

3 Latency bounds

Prior work on simulating read/write registers in a message passing model has established bounds on operation latencies. Informally speaking, these results observe that $r + w \geq d$ where r and w are upper bounds on the latencies of reads and writes, respectively, and d is a lower bound on the network delay. This point was first proved by Lipton and Sandberg for the coherent random access machine (CRAM) model [8], and then formalized and strengthened by Attiya and Welch for sequential consistency [2]. Both results assume partly synchronous models, and therefore neither can be applied directly in this paper because the worst-case latencies of reads and writes are unbounded in the model defined in Section 2 due to asynchrony. In fact, the upper bounds r and w do not exist if one considers all possible executions of a system, or even all fair executions. This statement remains true even if message delays are constant (i.e., messages are delivered and processed in a timely manner) because the processes are asynchronous. For example, a process that is enabled to send a message may take an arbitrarily long time to transfer that message to a communication channel.

The known lower bound on worst-case operation latency can be recast in the asynchronous model as a lower bound over a special subset of executions. As stated in Theorem 3.1, the lower bound is asserted universally for all executions in the special subset, and implies that operation latency greater than half of the minimum message delay is inherent in the protocol rather following from asynchrony alone.

Theorem 3.1. *Let A be an automaton that simulates a read/write register initialized to value v_0 in the asynchronous system model with at least two processes. Suppose that every execution of A is consistent. Let α be any execution of A that is available, comprises a write by some process P_W of some value $v_1 \neq v_0$ and a read by some other process P_R , and where the two operations are concurrent. Let r and w denote the latencies of the read and write in α , respectively, and let $d > 0$ be a lower bound on the message delay. Then $r + w \geq d$.*

Proof. Since every execution of A is assumed to be consistent, it follows that the read returns either v_0 or v_1 . Therefore, the following case analysis is exhaustive.

Case 1: The read returns v_1 .

First, note that the *invoke* action of P_W 's write causally precedes⁷ the *respond* action of P_R 's read,

⁴ Operation op_1 *precedes* operation op_2 if op_1 has a *respond* action and its timestamp is less than the timestamp of the *invoke* action of op_2 .

⁵ Operation op_1 is *concurrent* with operation op_2 if neither op_1 precedes op_2 nor op_2 precedes op_1 .

⁶ It is assumed without loss of generality that all messages sent are distinct.

⁷ The “causally precedes” relation is the transitive closure of two rules: action a causally precedes action b if a occurs before b at the same process, or if a sends a message that is received in b .

which implies that P_W communicates with P_R either directly or indirectly (i.e., by way of one or more other processes) in α . This is because α is otherwise indistinguishable to P_R and P_W from an execution where the events are shifted so that P_R 's read responds before P_W 's write begins (i.e., v_1 is read before v_1 is written), which would contradict the assumption that all executions of A are consistent. The causal relationship implies that α contains a set of messages $m_1, m_2, \dots, m_k$ for some $k \geq 1$, such that P_W sends m_1 during its write, P_R receives m_k during its read, and for $1 \leq i < k$ the recipient of m_i sends m_{i+1} after receiving m_i . Such a scenario is illustrated in Figure 1 in the special case where $k = 1$. Since the write begins before m_1 is sent and the read finishes after m_k is received, it follows that the *invoke* action of P_W 's write is separated from the *respond* action of P_R 's read by k message delays or kd time. Moreover, since the two operations are assumed to be concurrent, it also follows that the sum of their latencies is at least kd . Since $k \geq 1$, this implies $r + w \geq d$, as required.

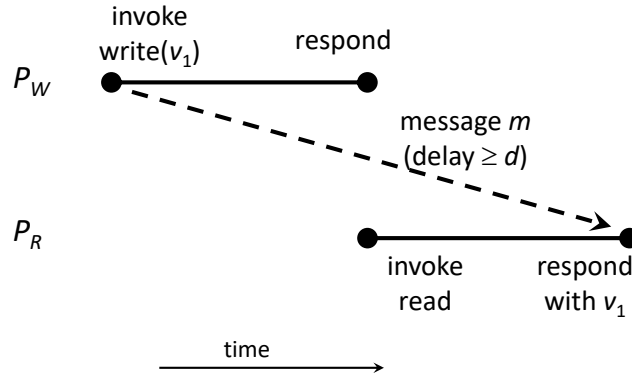


Figure 1: Execution α in the proof of Theorem 3.1.

Case 2: The read returns v_0 .

The analysis is similar to Case 1. First, note that the *invoke* action of P_R 's read causally precedes the *respond* action of P_W 's write, which implies that P_R communicates with P_W either directly or indirectly in α . This is because α is otherwise indistinguishable to P_R and P_W from an execution where the events are shifted so that P_W 's write responds before P_R 's read begins (i.e., v_0 is read after v_1 is written), which would contradict the assumption that all executions of A are consistent. The causal relationship implies that α contains a set of messages $m_1, m_2, \dots, m_k$ for some $k \geq 1$, such that P_R sends m_1 during its read, P_W receives m_k during its write, and for $1 \leq i < k$ the recipient of m_i sends m_{i+1} after receiving m_i . Since the read begins before m_1 is sent and the write finishes after m_k is received, it follows that the *invoke* action of P_R 's read is separated from the *respond* action of P_W 's write by k message delays or kd time. Moreover, since the two operations are assumed to be concurrent, it also follows that the sum of their latencies is at least kd . Since $k \geq 1$, this implies $r + w \geq d$, as required. $\square$

The proof of Theorem 3.1 considers two operations on a single read/write register, as opposed to [8, 2] where a weaker four operations on two registers are considered. This is a consequence of the different interpretations of consistency: this paper deals with atomicity and regularity, which assume that *invoke* and *respond* actions are totally ordered; [8, 2] deal with sequential consistency, which assumes that such actions are only partially ordered (by program order and the “reads-from”

relation).

4 From latency bounds to CAP

The CAP principle in the context of the model from Section 2 is stated formally in Theorem 4.1 below, which is modeled after Theorem 1 in [5].

Theorem 4.1 (CAP). *Let A be an automaton that simulates a read/write register in the asynchronous system model with at least two processes. Then A cannot satisfy both of the following properties in every fair execution α (including executions that are not partition-free):*

- α is consistent
- α is available

Gilbert and Lynch prove their version of Theorem 4.1 by contradiction, supposing initially that A ensures both consistency and availability. They construct an execution α involving at least two processes, initially partitioned into two disjoint groups $\{G_1, G_2\}$ that are unable to communicate with each other due to dropped messages. Letting v_0 denote the initial value of the read/write register, some process in G_1 is chosen to invoke a write operation that assigns a new value $v_1 \neq v_0$. Since A ensures that α is available, even if it is not partition-free, this write produces a response eventually. Next, a process in G_2 is chosen to invoke a read operation, which once again produces a response eventually. However, since there is no communication between processes in G_1 and processes in G_2 , α is indistinguishable to processes in G_2 from an execution where the write never occurs. Therefore, the read must return v_0 instead of v_1 , which contradicts the assumption that α is consistent in addition to being available.

Alternatively, Theorem 4.1 can be proved using the latency bound from Section 3. Roughly speaking, the proof argues that if a system ensures consistency then operation latencies grow with message delays, and hence operations cannot terminate eventually (i.e., system cannot ensure availability) if the network is partitioned.

Alternative proof of Theorem 4.1. Let A be an automaton that simulates a read/write register in the asynchronous system model with at least two processes. Suppose for contradiction that A ensures that every fair execution is both consistent and available, even if the execution is not partition-free. Let P_W and P_R be distinct processes, and suppose that the network drops all messages. There exists a fair execution α_1 of A where the initial value of the register is v_0 , then P_W writes $v_1 \neq v_0$, then P_R immediately reads the register (i.e., P_R 's *invoke* action is consecutive with P_W 's *respond* action) and produces a response. Since A ensures consistency even if the execution is not partition-free, this implies that P_R 's read returns v_1 and not v_0 . Now let α_2 be the prefix of α_1 ending in the state immediately following the read's response. Since α_2 is indistinguishable to all processes from an execution where the messages are merely delayed and not dropped, it is possible to extend α_2 to a finite partition-free execution α_3 by delivering all sent messages eventually (after the response of the read), without introducing any additional read or write operations. Now construct α_4 from α_3 by swapping the relative order of P_R 's *invoke* action and P_W 's *respond* action, which preserves the property that the execution is both consistent and available, and also makes Theorem 3.1 applicable. Let w and r denote the latencies of the write and read, respectively, in α_4 . Suppose without loss of generality that the message delay d in α_4 is constant and greater than

$r + w$, which ensures that no message sent by P_W after starting its write can influence the outcome of P_R 's read operation. This scenario is illustrated in Figure 2 in the simplified case when P_W and P_R are the only two processes in the system. Then α_4 contradicts Theorem 3.1 since this execution is both available and consistent with $d > r + w$. $\square$

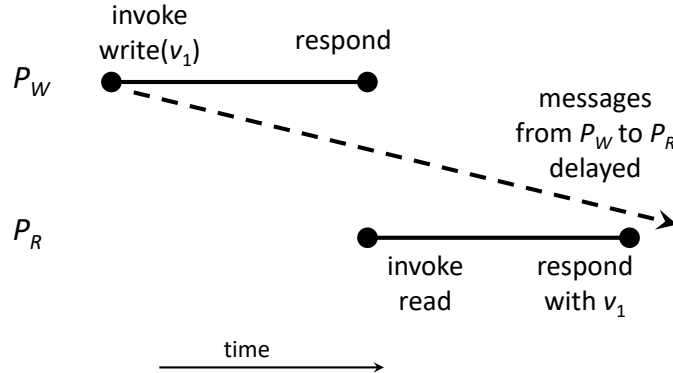


Figure 2: Execution α_4 in the alternative proof of Theorem 4.1.

5 Formal interpretation of PACELC

The conjunction of Theorem 4.1 and Theorem 3.1, both of which are proved in this paper using latency arguments in an asynchronous model, constitutes a formal statement of Abadi's PACELC formulation. Theorem 4.1 implies that for executions that are fair and not partition-free, the system cannot always guarantee both consistency and availability: *if Partition then Availability or Consistency*. On the other hand, Theorem 3.1 implies that for executions that are partition-free, the system cannot always guarantee both consistency and operation latency less than half of the minimum message delay, irrespective of asynchrony (i.e., even if message delays are constant and processing delays⁸ are zero): *Else Consistency or Latency*.

Attiya and Welch [2] proved that the latency lower bound $r + w \geq d$ stated in Theorem 3.1 is tight in a partially synchronous model where processes have access to local clocks that can be used as timers, and where message delays are constant. Specifically, if message delays are exactly d (which implies partition-freedom), then there exists a protocol that guarantees atomic consistency and where either reads are instantaneous and writes have latency d , or reads have latency d and writes are instantaneous. Such protocols maintain a copy of the register's state at each process, and use timed delays to compensate for message delays. For example, in the instantaneous read protocol, a read operation returns the local copy of the state without any network communication, whereas a write operation first broadcasts the new value to all other processes, then sleeps for d time, and finally updates its local state. A process updates its local copy of the state instantaneously upon receiving the broadcast value.

⁸ In the asynchronous model with timed executions, one can define processing delay as the time between when a *send*, *receive*, or *respond* action is enabled and when that action is executed.

Practical distributed storage systems such as Amazon’s Dynamo [4] and its open-source derivatives are designed to operate in a failure-prone environment, and therefore rely on explicit acknowledgments rather than timed delays to ensure delivery (i.e., receipt and processing) of messages between processes. As a result, these systems exhibit operation latencies exceeding the lower bound in Theorem 3.1 by a factor of at least two even in executions where message delays are exactly d . For example, a quorum-replicated system such as Amazon’s Dynamo [4] can be configured for local reading but then writing requires a full network round trip, or $2d$ time, to ensure Lamport’s regularity property [7]. This is accomplished using full replication and a read-one, write-all quorum configuration. Operation latency is increased further if the system is configured to tolerate server failures, for example by using majority quorums, in which case both reads and writes require at least $2d$ time.

6 Conclusion

This paper presented both an alternative proof of the CAP principle and a formal treatment of Abadi’s PACELC formulation based on the inherent trade-off between operation latency and network delay. These results complement and extend the CAP theorem of Gilbert and Lynch, which was published prior to Abadi’s article, and draw a precise connection between CAP-related trade-offs and latency bounds for shared objects.

References

- [1] D. Abadi. Consistency tradeoffs in modern distributed database system design: CAP is only part of the story. *IEEE Computer*, 45(2):37–42, 2012.
- [2] H. Attiya and J. L. Welch. Sequential consistency versus linearizability. *ACM Trans. Comput. Syst.*, 12(2):91–122, 1994.
- [3] E. A. Brewer. Towards robust distributed systems. In *Proc. of the 19th ACM Symposium on Principles of Distributed Computing (PODC)*, page 7, 2000.
- [4] G. DeCandia, D. Hastorun, M. Jampani, G. Kakulapati, A. Lakshman, A. Pilchin, S. Sivasubramanian, P. Voshall, and W. Vogels. Dynamo: Amazon’s highly available key-value store. In *Proc. of the 21st ACM Symposium on Operating System Principles (SOSP)*, pages 205–220, October 2007.
- [5] S. Gilbert and N. A. Lynch. Brewer’s conjecture and the feasibility of consistent, available, partition-tolerant web services. *SIGACT News*, 33(2):51–59, 2002.
- [6] M. Herlihy and J. M. Wing. Linearizability: A correctness condition for concurrent objects. *ACM Transactions on Programming Languages and Systems*, 12(3):463–492, July 1990.
- [7] L. Lamport. On interprocess communication, Part I: Basic formalism and Part II: Algorithms. *Distributed Computing*, 1(2):77–101, June 1986.
- [8] R. J. Lipton and J. Sandberg. PRAM: A scalable shared memory. Technical Report CS-TR-180-88, Princeton University, 1988.

- [9] N. Lynch. *Distributed Algorithms*. Morgan Kaufman, 1996.

The First Summer School on Practice and Theory of Concurrent Computing SPTCC 2017

Petr Kuznetsov
Télécom ParisTech and Université Paris Saclay
France
`petr.kuznetsov@telecom-paristech.fr`



Nowadays concurrency is everywhere. Be it a mainstream multi-core machine, a computing cluster, or a large-scale distributed service, a modern computing system involves multiple processes that concurrently perform independent computations and communicate to synchronize their activities. Understanding concurrency is therefore getting essential in both practice and research in computer science.

The first summer school on Practice and Theory of Concurrent Computing took place on July 3-7, 2017 in Saint-Petersburg, Russia. The school was hosted by ITMO university, and financially supported by DevExperts, Yandex, Télécom ParisTech, and ANR-DGF DISCMAT project.



1 Program

The goal of the school was to give both an introduction to the field of concurrency and a glimpse of the state of the art. We interleaved classes on the fundamentals (wait-freedom and linearizability, lock-based and nonblocking synchronization, universal constructions) with systems-oriented classes (transactional data types, distributed recommenders, relaxed data structures and memory management).

- The course “Locking, from Traditional to Modern” given by Nir Shavit served as an introduction to the very basics of synchronization in computing, starting from simple Test-and-Set locks and proceeding to increasingly more sophisticated locking techniques. The class explored the efficiency features of various lock implementation, such as cache locality and scalability, and concluded with a discussion of prominent client-server style locking, such as Owiki locks and flat combining.
- Danny Hendler gave the course “Lock-free concurrent data structures” discussing concurrent algorithms that do not use locks. Such *lock-free* algorithms provide more resilience with respect to asynchronous conditions than more conventional, lock-based algorithms. The course covered lock-free algorithms for several concurrent data-structures, with a focus on algorithmic techniques, such as elimination, that can be used for devising efficient lock-free implementations.



- “Wait-free computing for dummies”, the course given by Rachid Guerraoui, provided a study of *wait-free* algorithms that enable the highest degree of robustness to asynchrony and failures. The course recalled the fundamental impossibility and universality results, and covered implementations of read-write shared memory, snapshots and counters.
- Michel Raynal gave a comprehensive survey of *universal constructions*, algorithms that build a distributed implementation of any object provided its sequential specification, with a focus on basic concepts and mechanisms these constructions rely on.



- Maurice Herlihy offered the course “Transactional Memory”, devoted to a fascinating programming abstraction that intends to enable a simple programming interface and to efficiently exploit the increasing computing parallelism provided by modern machines. The course gave a survey of the area, a discussion of existing algorithms and open research questions.
- Liuba Shrira gave the course “Implementation techniques for libraries of transactional concurrent data types” that addressed the relation between the *semantics* of a concurrent data type and the *efficiency* of its implementations. It is often argued that transactional memory is inherently subject to fundamental performance limitations caused by its very nature of being universal. The course explored the advantages of knowing the semantics of a data type in the transactional context and discussed mechanisms of tracking conflicts between abstract operations.



- The course “Recommenders and distributed machine learning: algorithms and systems” given by Anne-Marie Kermarrec discussed the concept of *recommendation* and surveyed various algorithmic techniques to implement it. The focus here was on distributed recommenders based on *collaborative filtering*.
- Roman Elizarov, a leading software engineer at JetBrains and a tutor at ITMO University, gave the course “Lock-Free Algorithms for Kotlin Coroutines”. The course gave an overview of lock-free data structures, such as doubly-linked lists and multiword compare-and-swap (CASN), that can be used as synchronization abstractions for Kotlin language coroutines.



- The course “Relaxed concurrent data structures” given by Dan Alistarh is devoted to the emerging class of data structures that provide higher degree of parallelization at the expense of weaker consistency guarantees.
- Erez Petrank proposed the course “Memory management for concurrent data structures”, containing a comprehensive introduction into concurrent *garbage collection*. A particular attention was devoted to the open problem of *lock-free* memory management. Then the course discussed specific memory managers that are meant to provide support for lock-free data structures without foiling their lock-free guarantees.

The school ended with a visit to the Saint-Petersburg office of Yandex, where Ivan Puzyrevskij gave us a very interesting insight on how the Yandex data infrastructure is organized (apparently an application-specific variant of Paxos is implemented from scratch!).

Links to the school material (slides, exercises and videos) are available at <http://neerc.ifmo.ru/sptcc/courses.html>. A photo album can be found here.



2 Outcomes

I believe that this was the first time an event of this scale devoted to a topic in distributed computing took place in Russia, and I am particularly happy that it happened in my beloved city of Saint-Petersburg. It is safe to say that the school was a success, primarily thanks to the quality of the courses and the engagement of the audience. It appears that both teachers and attendees enjoyed it, which gives aspirations for future editions.

Last but not least, I would like to express my gratitude to the people without whom this school would never take place. Vitaly Aksenov, a doctoral student at ITMO and INRIA, Paris, was invaluable in the organization. Daria Kozlova, Daria Yakovleva and Vladimir Ulyantsev were extremely helpful on the ITMO side. We gratefully acknowledge the video and photo support generously provided by Alexey Fyodorov and the JUG.ru team. Very special thanks should go to Mikhail Babushkin, DevExperts, for helping on the financial side, and to Roman Elizarov for inspiring the very idea of running such a school in Saint-Petersburg.



3 Personal impressions

We conclude with several personal accounts from teachers and attendees.

I gave two lectures on emerging synchronization techniques, covering both hardware and software aspects. A number of students asked questions, both during the lectures and after class. I was impressed by the quality of the questions. The students were very well-prepared, and had studied the material very carefully. (In fact, some of them pointed out typos and small errors in my presentation.)

Above and beyond the technical aspects of the course, the ability for students to network with a larger community is an important contribution of the school. Outside lectures, I spoke to a number of students about their research, giving them advice on which topics were timely, how to pursue specific topics, and which conferences might be appropriate for their papers. I think some of these students felt isolated, and seemed very happy to discover they could connect to a larger community of researchers. In the weeks since the class ended, I have corresponded with several of the students, commenting on technical issues, and offering advice on how to keep up with the current state of research in their areas.

— Maurice Herlihy, Brown University, USA



I participated as a lecturer in the THE FIRST SUMMER SCHOOL ON PRACTICE AND THEORY OF CONCURRENT COMPUTING. I enjoyed the experience very much. The crowd was highly engaged and attentive. There were lines of students waiting to ask questions in between the two lectures and after the second lectures. The questions reflected knowledge and seriousness. Some of the attendants asked me about joining my group to do research together, which seemed like a very strong positive feedback. The organization was excellent and high quality videos of the lectures are now available online. Last, I did not expect to enjoy the location so much. Saint Petersburg is a great city and I recommend it to all.

— Erez Petrank, Technion, Israel



I had the pleasure of attending the first ever summer school on the theory and practice of concurrent programming. The school was very well organised, making normally painful university bureaucracy a breeze. The invited speakers were of the absolute top quality. The school size was small, about 100 people in attendance and single track so I didn't miss out on any lectures. I had the opportunity to sit down with the speakers themselves during the breaks and during the student sessions. The one to ones I had were very fruitful, leaving me with new ideas in my head and contacts for the future. Discussions surrounding the exercises were very interesting, with new insights on how other people solve problems and a lay of the land in the field itself. All in all it was the most academically enlightening event I have ever had the fortune of attending.

— Robert Kelly, Graduate Student, Maynooth University, Ireland



I attended the first summer school on practice and theory of concurrent computing on July 3rd to 7th at ITMO University in Saint Petersburg and I felt it was an extremely valuable learning opportunity. During the summer school we had the opportunity to see active researchers and practitioners in the field present on foundational and emerging topics, with topics ranging from practical applications to theoretical foundations.

Being able to interact with notable members of the research community, such as Maurice Herlihy, was a special opportunity as someone just beginning my research career. All the presenters were extremely knowledgeable and insightful, and graciously answered many student questions and inquires throughout. The mixture of both industry and academic presenters gave me unique insights into the different approaches and concerns of industry and academia toward problems in the field. I learned many valuable techniques and insights into concurrent computing during the week that I hope to bring back to my own research.

— David Tenty, Graduate Student, Ryerson University, Canada



I was participating in the SPTCC school, held on July 3-7 in the St. Petersburg, Russia. There were several reasons why I decided to participate. First, the invited professors. The school was the rare occasion to meet in person the most famous researchers in the field and to learn about the cutting-edge research topics. Second, while being a practitioner, I realize how important foundations of concurrent and distributed computing are. To build a truly reliable, predictable and scalable system, one must have a strong background in these areas. My objective for the school was to revisit and enrich my knowledge – and that was fully accomplished.

Also, I was delighted to see both attendees from the industry and the academia. Having a sustainable interest from both researchers and practitioners is important for the discipline to flourish. I met colleagues working on programming languages, tooling, high-frequency trading software and data infrastructure. Such a variety of applications re-emphasizes the importance of school topics. I hope and to see and to help the school return back next year.

— Ivan Puzyrevskij, Software Engineer, Yandex, Russia



At first I wasn't sure I would be able to participate, because my first baby had just been born and required much attention from us. However, the organizers managed to gather a respectable crowd of lecturers, so I couldn't let that opportunity slip — and wasn't disappointed. The lectures covered both the classical topics and the new material. I would highlight Nir Shavit's lecture as the best one about classic results, Dan Alistarh's as the best one about current research, and Roman Elizarov for the practitioner's point of view. Another highlight is the exercises and a (somewhat improvized) discussion of them at the end of the school, which was an excellent idea.

— Igor Baltiysky, Software Engineer, ALM Works, Russia

In the beginning of July 2017, the Summer School on Practice and Theory of Concurrent Computing (SPTCC) took place in Saint-Petersburg. I was very fortunate to attend it, as the curriculum consisted entirely of courses by brilliant lecturers, who initiated and explored many exciting topics in concurrent computing. Having spent some of my PhD on verification of concurrent algorithms, I was looking forward to getting a new perspective at the state of research in the field.

Overall, the lectures were a great crash course for beginning students and offered quite a bit of food for thought to more experienced ones. Outside of the lecture room, the attendees could meet fellow students and concurrency enthusiasts from Facebook, JetBrains, Yandex, Intel etc. Interacting with them was a great fun, which contributed to a friendly atmosphere by the end of the school, and, of course, the beautiful city of Saint-Petersburg also added a lot of charm to the event. All in all, I genuinely recommend everyone interested in concurrent computing to keep an eye on the future editions of SPTCC. Many thanks to Petr, Vitaly and the others from the organizing committee for making an amazing event.

— Artem Khyzha, Graduate Student, IMDEA Software Institute, Spain

