

MASTER'S THESIS

A Generic Event Logger For Real-Time, Web-Based
Applications

Timelines, Selection & Intelligent Replay

Supervisor: Prof. Dr. Anton Eliens

Second Reader: Prof. Dr. Spyros Voulgaris

Leonidas Diamantis

l.diamantis@student.vu.nl

VU University 2014

Abstract

This thesis explores the possibilities of logging, selecting and intelligently replaying events in a generic manner. More specifically the thesis studies the potential of the Web platform for tracking the interaction of one or multiple users on a set of real-time, collaborative Web-applications, and projecting the outcomes on what we defined as the “R³ paradigm”, consisting of three aspects: i) record, by the means of a system that can capture the direct interaction as well as the context reaction within the application, ii) replay, using mechanisms to facilitate the re-occurrence of the application usage and re-stimulating to behave, and iii) reflect, by using the replayed behavior to draw conclusions and possibly shape the decisions of a user or administrator in a future usage. The research focused on defining an application which expresses a real life metaphor to be used as a proof of concept for the logging system. After building a theoretical framework and illustrating the metaphor, the thesis presents an architecture for a logging system which performs the basic operations, and reflects upon its efficiency & performance on various scenarios & versions of the concept application. The thesis continues by exploring and suggesting ideas & insights for improvement and natural integration of such a system in a distributed environment. The thesis concludes with ideas for further development and ways to take advantage of the results produced by such a system, followed by general conclusions.

Preface

This thesis is the product of almost one year of personal research & practical work, driven by a strong interest and deep inspiration, developed during my master studies, in generic event logging & intelligent replay in a distributed fashion. Motivated by serious games, computational analysis and artistic expression, focusing on real-time Web-based applications and aiming for results that would stimulate vastly diverse interpretations ranging from statistical analysis to sociology, art & philosophy. During my research and implementation of the logging system, I came across many big challenges ranging within the complete spectrum, from strict technical obstacles to conceptual gaps that needed to be bridged. With extra inspiration, I ultimately became able to put the pieces together and build a working prototype which would log & replay events that occur by using a real-time & collaborative concept Web-application I also built.

This study tries to balance between theoretical and applied research, attempting a systematic approach towards generic & flexible discrete time manipulation in the Web. I think that the results of this research provide very interesting insights and adequate motivation to move towards further study, for which there is still a lot of room left.

Leonidas Diamantis

Acknowledgments

This document is the result of combined research, considerations and efforts to create a system that uses technology to produce multi-dimensional results, from analytical insights to artistic performances. A system that allows users to “use the past in order to change the future”. During the implementation process there were many scientific, technical and moral aspects of the topic that account for, and quite a few decisions - both conceptual and technical - had to be made. And I honestly feel very lucky that I had a number of people who helped me in so many ways throughout this so challenging endeavor.

First and foremost, I would like to thank my supervisor, Prof. Dr. Anton Eliens, for his guidance & support, for all the times he inspired me to continue and all patience he showed with my technical dilemmas. I would also like to express my gratitude to Prof. Dr. Spyros Voulgaris for being an excellent academic mentor for me, for agreeing to be the second reader of this document and for the invaluable advice he always offered, especially during the kickstart and the design phase of my Master's Thesis. Furthermore I would like to thank my friends and colleagues Chronis Kalaitidis, Andreas Manios, Nikos Poullos, Lucas Sakizloglou & Bogdan Banu for sharing their comments and ideas with me during our conversations about this project. Last but not least, I want to thank Konstantina Mavridou for believing in me and for dedicating so much of her time & energy to morally support me along the way. I will forever be grateful for this.

Chapter 1: Introduction	8
Motivations	8
Research Questions, Scope & Approach	9
Outline	10
Chapter 2: Background	11
Definition Of The Event	11
System Overview	12
Sketch Of Scenarios	14
Single versus multi-modality	15
Uncritical versus critical replay	16
Application & Metaphor: A Simple Voting System	19
Chapter 3: Technical Architecture	21
Technological Overview	21
Multi-Modality & Collaboration	25
Event distribution	26
The cross-domain AJAX problem	28
Maintaining a global state in the server	29
Chapter 4: Performance & Efficiency Considerations	30
Scenario: High Event Density	30
Scenario: Geographically Distributed Interaction	33
Distribution of services	34
Distribution of data	35
Synchronization	36

Chapter 5: Towards A Fully Generic Event Logger	39
Extending The Event Format For Complex Event Structures	39
Introducing Reactive Programming Principles	40
Smart Loggers	41
Chapter 6: Future Work & Conclusions	43
Improving The Logger	43
Making Use Of Event Streams	43
Analysis/visualization tools	44
Manipulation of time	44
Expanding the users' conscience	44
General Conclusions	45
Appendix A: Key Operations Code Snippets	46
Client-Side Event Registration & Triggering	46
Enabling CORS In The Middleware	47
Server-Side Client-To-Socket Binding & Event Distribution	47
Appendix B: Journal	48
The Beginning	48
Pencil & Paper	49
Bibliography	50
References	50

Chapter 1: Introduction

The concept of posterior analysis of system behaviors is essential in the world of science, inspiring and concerning multiple fields from Mathematics to Computer Science and Business Intelligence, stimulating people to consider & invent advanced techniques and methods which would help them analyze the behavior of certain systems and draw useful conclusions about how they (mis)behave.

The whole process of such analysis boils down to one main task: extracting information from a dataset with data specific to the system's execution, and performing different kinds of analyses to determine how that system behaved, based on some criteria. Obviously, the nature of this dataset as in what types of data it carries will produce different kinds of conclusions at different depths. With the most common approach (from a Computer Science perspective) on this matter being that such information set would mainly include raw, computationally produced data (e.g the sum of two numbers for an addition application) the conclusion-making process is bound to follow specific analysis patterns and produce conclusions following respective patterns as well.

The system built for this thesis comes to offer an alternative approach to the one described above, in respect to the creation of the dataset: for our scope, we focus on capturing the actual interaction between user and application and creating a dataset which, instead of the data that the application produces, consists of meta-information about how has the application been made use of, from the user's standpoint. And with the (selective) replay features our system implements, the conclusion-making process for the application's behavior could be done by examining exactly what kind of usage made the system behave in the way it did.

Motivations

The research carried out for this Thesis was motivated by two main domains: computational analysis & the digital arts.

While considering the scope & approach for this project, we considered various use-case scenarios and types of systems which would render such a logging system useful. We considered collaborative applications which impose distortions in the communication between participants and various distributed systems which have the tendency to fail on arbitrary moments and modes. For such scenarios, we thought that such a system would make “the building blocks of determining what went wrong in a distributed system”, from a

user's perspective, as the information in which the logger is interested is the actual usage that made the application behave in the (possibly faulty) way that it did.

As a counterpart to the strict, Computer Science driven domain, we also considered a more abstract platform where a generic event logging system would bring some value. Inspired from time & synchronization oriented artistic performances, we gained motivation to create a system which would help “use the past to shape the future”, in the sense that things that happened could (selectively) re-happen, possibly extended or manipulated with an artistic twist.

Research Questions, Scope & Approach

The ways we chose to approach this problem are mostly based in exploration. First, we considered a number of real-life scenarios and metaphors where we believe that such a system would be useful and capable of producing valuable results. We created a number of simple applications to represent these scenarios and metaphors, and built the logging system around them, always keeping in mind that it has to be generic so we tried not to bind the system to the applications in ways that it would become an application-specific plugin. Our purpose was to test the flexibility of the platform (the Web), the levels of abstraction and generic representation that such a system can achieve and the degree of how can such a system stand in a distributed environment.

The two main research questions of this thesis are:

- How and in what depth can the information of user interaction be captured as a generic entity?
- What reliable way is there for the logged user interaction to be replayed and produce a consistent behavior of the system?

For the scope of this project, even though we provide insights on how could this information be used later on, we do not focus on the conclusion-making process as much as creating the platform for logging such interaction events & producing such datasets, and ultimately offering features which would re-produce the occurrence of these events in a flexible way. Furthermore, the platform of choice for this thesis is the Web, focusing on Web-based applications which could be used by one or multiple users, possibly in a real-time manner. The reasons for forming the scope as such are the following:

- The Web has evolved in ways that made it a platform for purely distributed systems.

- The feature sets of Web applications and the services they provide become richer and more popular every day, making the Web the prominent platform for most types of Internet usage.
- There has been a considerable amount of research for capture/replay techniques for stand-alone applications (mostly focused on Java & JVM-based technologies [1]), yet there is not much research done in that respect for the ever-growing Web platform.

Thus, it is reasonable to say that the scope we define covers some cutting-edge, real-world systems with a highly distributed nature in all senses and allows us to touch upon fresh ground to research and conclude with interesting insights for the future.

Outline

Chapter 2 aims to position the topic of this project in a theoretical space by discussing the major components of the research and by elaborating on the different scenarios which were taken into account for designing and implementing the logging system.

Chapter 3 focuses on the technical aspect of the project, explaining the architecture design of the system, the technologies which were used and the reasons why those technologies were selected and examining the techniques that were implemented on each phase and use case. Additionally, it discusses issues that were raised by the distributed nature of the system & the application and ways they can be tackled.

In Chapter 4, we look upon some efficiency considerations with regards to the performance of the system in various cases. Chapter 5 provides a series of insights for creating a more generic and performant event logging system, by examining specific areas in which it can be improved.

The thesis ends with Chapter 6 where we present a list of ideas for future work regarding the use of such a system's results and datasets, along with certain features that would transform it into a generic, domain-agnostic, cutting edge tool. Furthermore, we reflect upon the contributions of this research, followed by our final conclusions.

Chapter 2: Background

In this chapter, we try to provide a theoretical background about the system that was built for this project. We explain the working principles of the major components of the logger, define the basic entities of interest and sketch the high level scenarios & criteria we decided to investigate. We also give an overview of the application(s) we built for exploring the features of the logging system, and attempt to enrich the application concepts with some metaphor(s) from real life, allowing the reader to place the event logger in a context more tangible than pure, abstract academic research.

Definition Of The Event

An event, as we decide to define it in the context of this project, is the product of any action or interaction which will provide an application with enough input to stimulate it to behave in the manner it was designed for. For this project, the event is the quintessential piece of information, as it is what we are interested in logging and reproducing while in the replay phase. Through streams of events, or *timelines* as they are also interchangeably referred to in this document, we get the picture of “*what happened*” while the application was used.

Since the event logger operates in the Web realm, we expect the events we are dealing with to belong to a -broad, yet apparent- scope. For example, we know a-priori that the typical target element of an (inter)action occurrence will most probably be a DOM element, and that the primitive types of the events themselves (as objects or entities) will follow the computational scheme & representation prototype of a Web-interpreted language (e.g. JavaScript). In a later chapter we discuss concepts and cases where the event object could be extended for different (more advanced) kinds of logging and replay, but for the sake of consistency with the design principles of the logging system, at this stage we will focus on the main characteristics of the event object which enable the logger to perform some meaningful work.

The main characteristics (or properties, since we are treating the event as an object) of an event that we are interested in are the following:

- **Type**, which indicates whether the event is the press of a keyboard button, a mouse click or movement, a native JavaScript event or any other event type that is relevant to the application and the logger.

- **Timestamp**, is the real UNIX timestamp of when the event occurred, typically measured at the source.
- **Relative Timestamp**, indicating the time difference between the occurrence of the current event and the previous one.
- **Source**, and identifier of the party that initiated the event occurrence.
- **Target**, which in turn is an object describing the (DOM) element of the application upon which the (inter)action took place. The target's properties are:
 - **ID**, the HTML id of the element (if it exists), given that it is a DOM object.
 - **Class**, the list of HTML classes of the element (if they exist), given that it is a DOM object.
 - **Tag-name**, the type of the HTML target element (div, form, etc).
 - **Text-content**, the inner content of the HTML target element, if it contains any.
 - **Type**, the value of the HTML *type* attribute of the target element, if it contains such an attribute (this attribute is usually a property of HTML form elements).
 - **Outer-HTML**, which is the source HTML code which describes the rendered target element.

Many of the event properties are optional, especially the ones which refer to the target. We want that for two main reasons: (1) the Web is a dynamic environment, and Web development is quite flexible regarding the definition of the elements of a web application, thus we want to be able to identify a target with as minimal amount of information as possible, and (2) we want the logger to treat event objects in the most generic way that it can, and that can be achieved by enabling a more loose event object structure.

System Overview

In order to provide with a high level overview of the logger, we want to list the main components which resemble it, as well as the phases in which its flow takes place.

The logging system integrates and functions in three main phases:

- **configuration**, where the nature of the application (multi/single modal), the types of events to be logged and the ways the replay is performed are defined,

- **event logging**, where the actual logging is registered and communicated from the application to the logger, and
- **event replay**, where an event stream is retrieved from the logging system to the application and the its behavior is replayed by the means of triggering the interaction events in that stream.

The aforementioned phases define a number of data flows between the application and the various components which resemble the logger. Specifically these components are:

- A set of **capture & replay mechanisms**, parts of which (typically concerning the types of events and the modality configuration of the application) are produced during the configuration phase. These mechanisms exist in a *client-side* context, a concept which is discussed extensively further on in this document.
- A **middleware layer**, which acts as a gateway of computation & logic between the application (specifically the capture & replay mechanisms) and the back-end of the logging system. Depending on the requirements and the scenario, this layer can be *fat* or *thin*, with regards to how rich functionality it has and how much computation and communication it needs to perform in the flow.
- A **API & storage server**, which exposes a set of web services to the application and/or the middleware for all relevant actions regarding storing or retrieving events & event streams, as well as maintaining states about the stage of the logging or retrieval/replay phase and generally performing global computations. Furthermore, the server maintains a datastore with all the relevant information about the interaction events and the participant(s) of the application.

[Figures 1](#) & [2](#) demonstrate the architecture and data flow in the logging phase and the replay phase respectively. We will revisit each component from a technical perspective in Chapter 3.

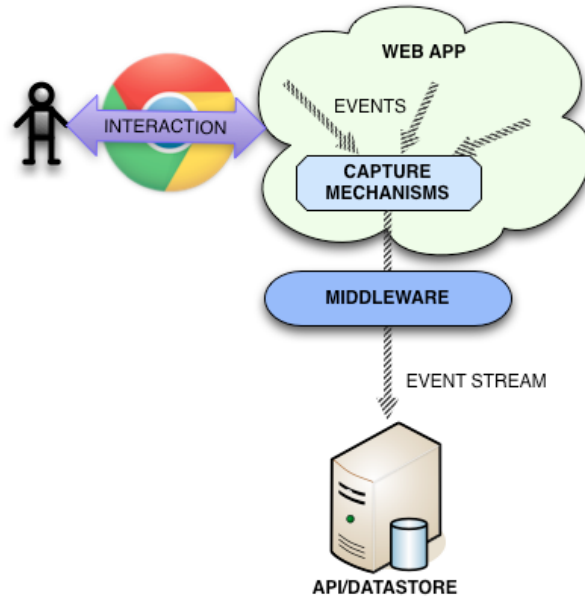


Figure 1: The flow of event logging.

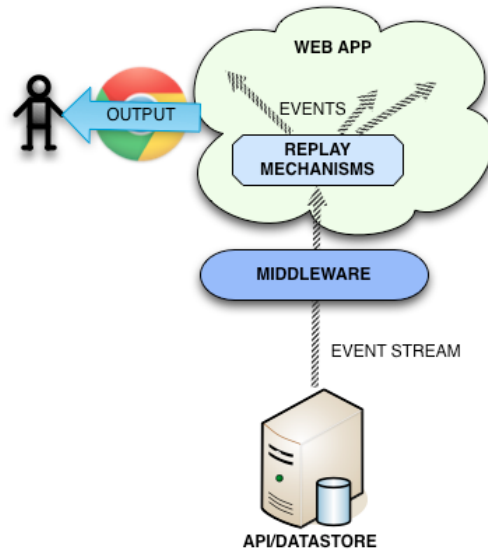


Figure 2: The flow of event replay.

Sketch Of Scenarios

In this section, we examine the cases which affect the design, implementation & feature set of the logging system as well as the application(s) under which it will operate. Each scenario is defined by applying a set of criteria of how should the application work, and in

turn it shapes the design process of the system. At this stage, we are interested in two main criteria: the application modality and how critical is the event replay operation.

With the term modality, we refer to the level of collaborative quality in the application usage: by single-modality, we imply that the application is used by one user at a time, whereas a multi-modal application is one that can be used by multiple users simultaneously (e.g a multi-player game or a voting system).

Considering the critical nature of the replay, we mainly refer to how strongly are the event occurrences coupled with each other, either in a causal or a temporal manner. This section elaborates on how each of the cases above affects the approach & risks of designing and implementing the logging & replay features of the system.

The application we built to implement & study the logger was created in versions in which we combined a number of variations between critical & uncritical replay cases of single and multi-user modes.

Single versus multi-modality

For this project, we are generally interested in applications which encourage the user to interact with the interface, because this is the main type of information we want to capture: events of (inter)action that will make the application expose its behavior. As a rough overview of our thoughts regarding application scope, our focus included applications which make the user participate in some sort of a “dialogue”, either with the application components or with other users or participants (e.g Web-based games, chat applications, voting applications, art-driven Web-based installations, etc).

In our context, a single-modal application is one where the aforementioned “dialogue” takes place between a single user and the application’s components, for example a single-player game or a drawing application. In this scenario, we have the a-priori knowledge that the timelines logged by the system originate from one source only, since there is only one user who interacts with the application. This knowledge brings specific implications when designing the logger and shapes the system in the following ways:

- When storing an event, the identity of the source is not critical because there is only one source and, therefore, only one recipient of a timeline to be replayed.
- The event logging & timeline retrieval API does not have to consider distribution of events and timelines to multiple recipients.

- The middleware layer between the application and the storage/API server needs not be aware of any peers, because there are not any.

Regarding the scenario of a multi-modal application, the fundamental difference with the single-modality case in the application level is that the input -and therefore the event streams- originate from multiple sources, often occurring simultaneously (e.g. multi-player games or collaborative art performance tools). This fundamental difference adds more depth to the design of the logger, bringing a number of intricate matters to be handled by the system.

Firstly, we know that in the multi-modal scenario the source where the events originate from matters, because the timeline can consist of events that come from different sources. The identity of the application users or participants (whether it is explicit or implicit -a concept we are discussing in detail in Chapter 3) matters and needs to be taken into account both in the phase of logging the events as well as in the phase of replaying a timeline back to the application.

The second main concern for the logger in this scenario is the temporal relationships of the event occurrences and their synchronization. Since we have multiple sources interacting simultaneously and in real-time, the logger will most probably have to process events which originate from machines which do not have fully synchronized system clocks and behave under different network schemes (which adds a variety of communication delays in the whole process), so computing the relative times of event occurrences in a timeline of such a diverse nature of origins is not a trivial task. We thus had to consider various mechanisms for maintaining a synchronized environment, also depending on how critical the replay is (more on this concept in the next section), keeping some state infrastructure on the server side and ensuring that a multi-modal timeline reliably represents the original chain of interactions in the application.

Uncritical versus critical replay

As we denoted already, the most essential feature of the event logger is its replay functionality, by which the user(s) of a given application will be allowed to fetch the stream of events they created with their usage and make them happen again in the application environment. In this feature, we wanted to provide the user(s) with certain degrees of flexibility, such as the ability to selectively replay a subset of events from a given timeline or create multiple timelines with different properties. To achieve such flexibility, we used techniques which applied both in logging-time and in the selection phase of the event stream, and we

discuss their technical details in a later chapter of this document. Similarly to the modality factor we discussed before, in this case we had to take the replay criticality into account and design our replay techniques based on it. In order to decide on our approach, the basic question we had to answer for each scenario was: *How important is a **correct** behavior of the application while/after replaying a timeline?*

At this point, it might be necessary to put the term **correct** behavior into context. When a user interacts with an application, they obviously expect to get some results back. The nature of these results is vastly variable; they can be computationally generated, multi-media output, a networked invocation or an interaction with another application or user, just to name a few examples. For the scope of our project, we see correctness as the combination of the application-specific set of results and the expectations of the user as to what is the minimum information they need to get back from the application. For example, in a color-changing application, if the user commands the application to change the color from grey to green and all they want to see is a color changing, without caring about whether the changed color is green, then in the case that the application changes the color to red is still considered to be a correct behavior (even though probably something is going wrong in the application internals), from the user's perspective. This approach might seem unconventional, but if we consider our initial motivations and the fact that one of our targets for this project is the possible artistic expression through real-time computer systems (in which case randomness is a usual means for sparkling artistic glory), it becomes valid because it widens the horizons for interpreting the timelines as well as implementing the logging and replay functionalities of the system. Of course, if the application architecture dictates that there should be a deterministic relationship between usage and result output for its function to be sensible & correct, then the definition of correctness takes its traditional deterministic form and the user's expectations should not differ from the application-set paradigm of correct behavior.

We use the above note about correct application behavior to speak about the criticality of the replay functionality. In our system, we look at the importance of a correct application behavior when replaying the events that happened during usage. And in the scenario where the replay is not critical, we imply that the system behavior (as it is formed by the replayed timeline) can be acceptable even if it does not match the event stream from the original usage in a one-to-one basis. Such a concept refers to applications where there are no causal dependencies between interaction events, regardless of the application modality and multi-user support. Furthermore, uncritical replay applies to systems where there are no temporal dependencies in their usage, and if there are they do not matter when replay-

ing a timeline. Thus, in the case of uncritical replay we most likely speak about applications through which the user(s) aim to produce non-deterministic results, and possibly some artistic expression through replaying their interactions.

The most profound implication brought by uncritical replay when designing the logging system is the fact that we do not have to perform extensive “event housekeeping”, as we do not care much about the ordering of the events in time, the synchronization of interaction steps or sometimes even the source of the events, in the multi-modal case. Thus, the logging and replay phases are fairly simple and they require minimal specification. Additionally, the event structure as an entity becomes somewhat loose, in the sense that it needs to contain less meta-properties as there is a limited amount of information that the logger needs to know about the event objects in order to perform the logging & replay functions.

The landscape changes dramatically -as it is expected- in the case where replay has a critical nature. By critical replay we mean that the replayed event stream should reflect the exact original timeline which was produced by usage, or it should at least follow a set of strict criteria to achieve the desired (and inherently correct) behavior of the application during replay.

Typically, critical replay is required by applications which produce interaction timelines with events that are causally -and thus also temporally, from a timeline perspective- dependent on each other. It is important for such applications that the relationship of “*event B happened because event A happened*” is maintained while replaying a timeline, in order for the application to behave correctly, and this means that the logging system has to keep the correct ordering of the event occurrences into account. Additionally, since from a timeline perspective (and for the scope of this project) we want to assume that causally dependent events are also temporally dependent (rephrasing the causal concept to “*event A happened first and event B happened second*” indicates that events A & B are tied in time), the logging system needs to carefully handle the time of each event occurrence and order them accordingly in the timeline to be passed to the replay phase. With all the above in mind, and especially if we consider the scenario of a multi-user web-application with critical replay required, it is quite clear that achieving consistent logging and replay of events which occur in a distributed context and emerge from multiple sources with variable network and otherwise existent overhead is not a trivial issue. Chapter 3 discusses how we decided to tackle these issues in detail.

Application & Metaphor: A Simple Voting System

The application we built to test, prove & improve the features of the logging system had to comply with a set of criteria:

- It had to be a Web-based application.
- It should allow a single user to use it, as well as multiple users simultaneously.
- It needed to behave in a real-time fashion, with regards to the output the user(s) would receive as a product of their interaction.
- It had to represent a real-life metaphor, to assist our research and further argumentation regarding the usefulness of the logger.

For the purposes of this thesis, and given the fact that the application is merely a “play-ground” we used to base our case studies on and not the objective of the project, we wanted to create an environment that is minimal in terms of its visual interface & feature set, yet functional enough so we can observe the behavior of the logging system, as well as a representative example of a hypothetical real-life use case scenario.

The application we decided to build therefore consists of a fairly simple interface and functionality, and it supposedly represents a primitive, color based voting system, with which the user can express their stance on a given subject in three different levels, and each level is represented by a color that gets registered after their voting action and is shown as output on a square area, as [Figure 3](#) illustrates. The user is able to act on the controls of the application indefinitely, and each time the color they choose is painted on the square area in real time.

The application was built in two versions, one where only one user can vote at a time and their results only appear to them, and one where there is collaboration involved and multiple users can vote simultaneously, and each vote appears in all participants' output square areas.

The way with which the user can interact in this application is by clicking one of the color labeled buttons above the square area. Thus, the event logging configuration we applied was tracking the click events on these buttons, together with the complete context of interest and event-specific properties we discussed in an earlier section of this chapter. In Chapter 3 we provide a detailed technical description over what technologies we used as

building blocks to assemble the application's architecture, communication scheme & interaction with the logger.

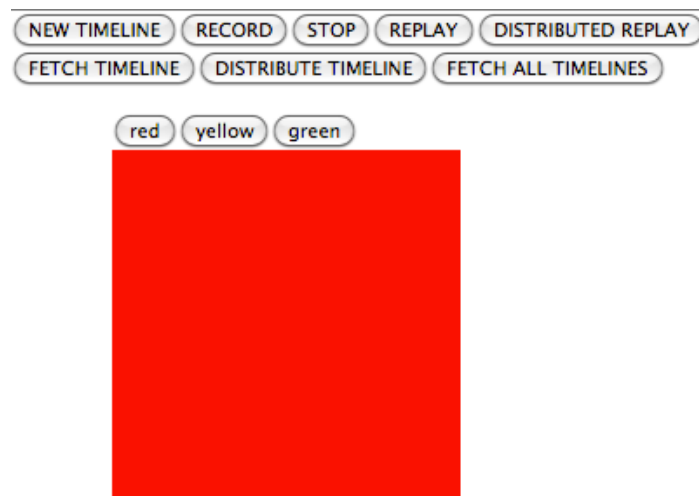


Figure 3: The interface of the application we built to test the logging system. The controls at the top of the window are application-independent and are used to manage the logger's features.

Chapter 3: Technical Architecture

In this chapter, we take an in-depth look over the building blocks of the logger and the application it operates upon, from a technical perspective. We discuss about specific technology selection & design decisions that were made and the reasons behind them, explain the execution & communication flow of the system and describe techniques we used handle & tackle issues that are brought up by the distributed environment which we worked on.

Technological Overview

Looking at the pace with which the Web is growing as a platform, one can say that in building a Web-based application, there is such thing as “too many options” throughout the whole stack of its architectural components. There are literally hundreds of frameworks, libraries, helper systems & meta-systems, spanning from the front-end to the back-end and in-between, all of them being different with each other and all claiming that they serve a specific purpose that others do not. Thus, the big challenge for a developer when starting to build a Web-based system is to align the requirements of their system with a -as small as possible- set of programming tools which will help intuitively visualize and construct the architectural components of the application. It is of big importance that this set should consist of blocks which harmonically interoperate with each other in programmability, data representation & communication aspects. As always, there is the overhead price to be paid in each case, but the main goal is that one selects the set of technologies which cover the application requirements in the most complete way possible. Following a bottom-up (from low-level to high-level components) approach, we will now present the technologies we chose for this project and explain the rationale behind each decision.

From a coding standpoint, we wanted to use a limited amount of programming languages and data representation schemes in order to achieve a natural cooperation among all the components of the logger and the application. And for that matter we were in luck, because as the Web-based technologies expanded, one programming language grew mature in order to support all the new features of the Web as a platform, and this language is -of course- JavaScript. Nowadays, JavaScript-based technologies have reached stages which allow developers to use only one language and build a full-fledged system and all its components, from complex server-side modules to cutting-edge front-end representation environments. Thus, our language of choice for this project was JavaScript, as it supports a vast amount of tools which cover the complete scope for this project. Moreover, to

achieve maximum compliancy in data representation, the model we chose is the JavaScript Object Notation (JSON) model, which is a text format for the serialization of structured data, and it derives from the object literals of JavaScript. JSON defines a small set of formatting rules for the portable representation of structured data [2] which makes it simple, lightweight, extremely easy to manipulate with JavaScript, and therefore a perfect candidate for becoming -if it is not there already- a standard for data interchange.

For a data storage engine, we wanted a system which would fulfill a set of general requirements:

- It should provide an easy-to-use & intuitive JavaScript interface.
- It should allow storage & management of flexible data structures, which may not follow a specific schema.
- It should perform efficiently when managing store & retrieve actions of big amounts of data.

On the database field, the big debate concerning the decision of which database engine to use is the one of “SQL versus NoSQL”, with the former referring to the classic Relational DataBase Management Systems (RDBMS) and the latter to the non-relational database engines which, although they have existed since the late 1960s, are actively being developed with new types of NoSQL database systems appearing and gaining market traction. The primary advantage of NoSQL databases is that, unlike relational databases, they handle unstructured data such as word-processing files, e-mail, multimedia, and social media efficiently [3]. A primary disadvantage of NoSQL in comparison to relational databases is the fact that NoSQL database engines, in order to achieve efficiency and speed when handling massive amounts of data, tend to “sacrifice” reliability. Relational databases natively support ACID (atomicity, consistency, isolation, durability), while NoSQL databases don't. NoSQL databases thus don't natively offer the degree of reliability that ACID provides. If users want NoSQL databases to apply ACID restraints to a data set, they must perform additional programming [4].

Coming back to our requirements, when picking a database engine for our project, we knew a priori that we were going to work with event objects, which are loosely structured and will most probably arrive at the back-end in high incoming rates, and will be retrieved in big chunks. Therefore, we decided that a risk on database reliability was a risk we wanted to take and a trade-off we could pay by enforcing reliability programmatically in the server -also having in mind that for an academic project such as the present thesis, strict

reliability on an engine level was not the main purpose-, while having a database that can efficiently store and retrieve big amounts of data which do not necessarily follow a rigid, pre-set schema. Thus, we chose to use a NoSQL system.

As it is expected, there are numerous NoSQL database engines out there, each built with a specific purpose in mind and optimized towards solving a class of problems, always following the non-relational paradigm. There are many ways to classify all the NoSQL systems that exist, depending on the most critical characteristic that reflects upon one's application requirements (e.g data model, query model, consistency characteristics, API support etc). For our project, the taxonomy we found most appropriate to decide upon which NoSQL system to use is the data model. Typically, there are four general data model classes of NoSQL systems (which contain sub-classes that follow the same principle and differ in specifics):

- **Key-value stores:** From a data model perspective, key-value stores are the most basic type of NoSQL database. Every item in the database is stored as an attribute name, or key, together with its value.
- **Document stores:** Whereas relational databases store data in rows and columns, document databases store data in documents. Documents typically follow a representation model akin to JSON, and can be seen as sets of key-value pairs of a dynamic structure.
- **Column family stores:** Such database engines use a sparse, distributed multi-dimensional sorted map to store data. Each record can vary in the number of columns that are stored, and columns can be nested inside other columns called super columns. Columns can be grouped together for access in column families, or columns can be spread across multiple column families. Data is retrieved by primary key per column family.
- **Graph databases:** These engines use graph structures with nodes, edges and properties to represent data. In essence, data is modeled as a network of relationships between specific elements.

We found the document store model most appropriate for the purposes of this project, since we immediately envisioned the event structure being compliant to a document-based structure inside a database. The specific system we chose to work with is MongoDB, with the rationale being that this engine is already quite mature, very efficient with operations upon big amounts of data, takes consistency & scalability considerations into

account by offering different types of addressing the aforementioned matters and provides a native JavaScript interface for database management, which fulfills the first of our database requirements mentioned above.

Maintaining the principle of reduced context switching between client & server-side development from a language standpoint, we chose Node.js for a back-end engine/framework. Node.js is a platform built on [Chrome's JavaScript runtime](#) for easily building fast, scalable network applications. It uses an event-driven, non-blocking I/O model that makes it lightweight and efficient, perfect for data-intensive real-time applications that run across distributed devices [5].

An interesting fact about how Node.js is implemented is that it does not use multithreading to execute tasks which need to run concurrently. It uses the event-driven model instead. One could think of the Node server process as a single-threaded daemon that embeds the JavaScript engine to support customization. This is different from most eventing systems for other programming languages, which come in the form of libraries: Node supports the eventing model at the language level [6].

Node's eventful nature of handling I/O & network operations helped create a conceptually consistent environment regarding the operations we wanted to perform between client & server via the API we created. Moreover, Node's interaction with MongoDB is almost native, which made the server-side back-end and data storage infrastructure cleaner & more intuitive, thus it was natural to choose it as the back-end platform for the logger.

For communication between server & client-side -and, for that matter, the middleware layer whenever applicable- we used two methods. The first one is the server-side exposure of a RESTful API for logging, retrieving & handling events that is consumed by the client-side via the classic AJAX protocol. Beneath this communication scheme, we created a lower level infrastructure based on Web-sockets, for direct and real-time communication among multiple peers at once. The technology we used to setup this infrastructure is the Socket.io engine, which is built with the mechanics of Node.js in mind and enables real-time bidirectional event-based communication [7]. We made extensive use of this infrastructure in the multi-modality scenario, in ways which are explained in detail later on in this chapter.

On the client-side, we chose the jQuery library for creating bindings between the application elements which are being tracked for interaction events, as well as for consuming the API exposed by the server. Although the usage of jQuery can result in unmaintainable code, one must admit that it offers straightforward mechanisms for recognizing DOM

events & handling them with standard JavaScript callback actions, and additionally provides a simple and fairly robust asynchronous AJAX wrapper of the standard XMLHttpRequest. Thus, for the purposes of this project and its beginning stage, jQuery offered an adequate set of features to perform all the event-related client-side operations.

[Figure 4](#) illustrates how all the aforementioned technologies are laid-out as the architectural components of the logging system. Later on in this chapter, we describe how we used these components to handle specific technical challenges.

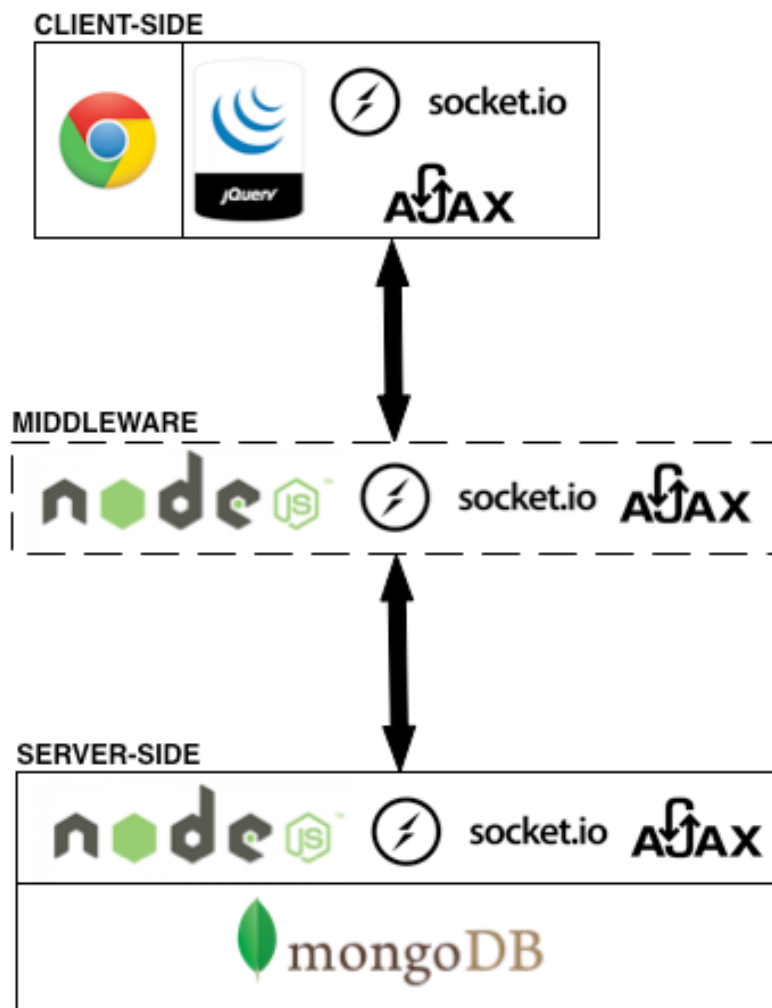


Figure 4: All the technologies used for the logger, laid-out as architectural components.

Multi-Modality & Collaboration

Although the single-modal version of the application (one user, monolithic timelines & fairly trivial replay) was a good starting point of the project, our real interest always resided on

the multi-modal version, having the application create an ad-hoc & distributed environment for the logger to operate under. In order to setup this scenario, we used a backend-as-a-service platform called Firebase to achieve real-time broadcast [8] of every vote to all the participants, connected to the application via a Angular.js [9] binding library called AngularFire [10]. All computations, communication & vote synchronization are done in the front-end, and the back-end of the application is a thin Node.js server which is acting as a simple Web-server, serving the pages, styles & scripts of the application. [Figure 5](#) sketches the high-level architecture of the application in this multi-modal version.

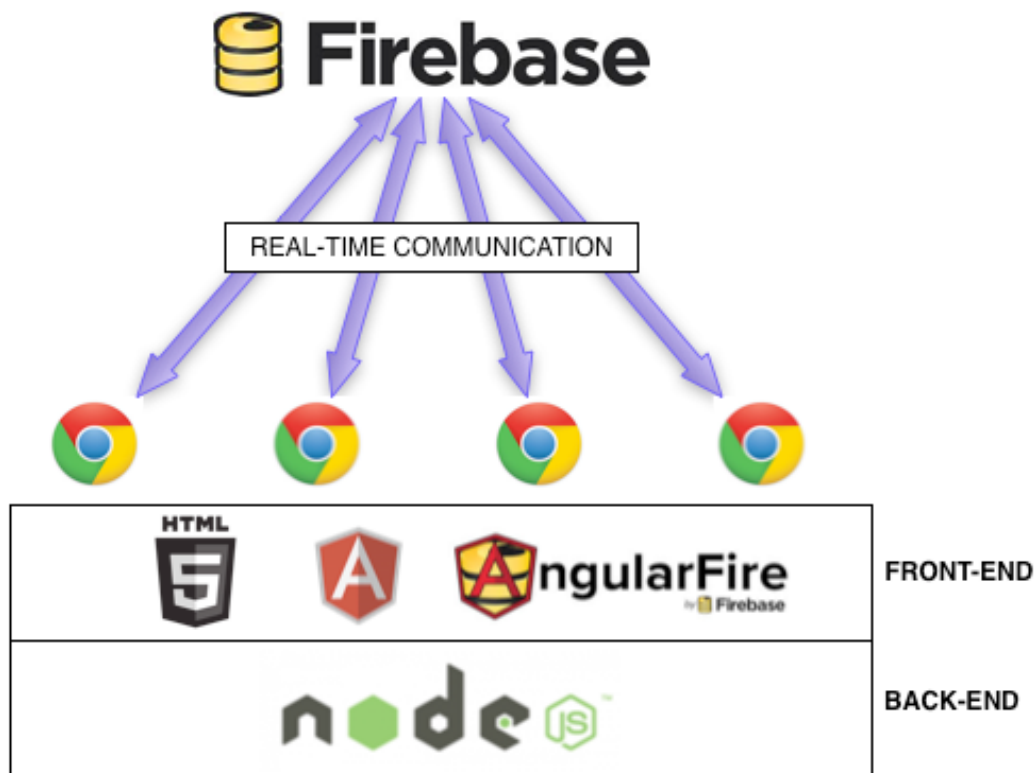


Figure 5: The high-level architecture of the multi-modal, collaborative version of the voting Web-application. All communication is synchronized & broadcasted real-time by the Firebase back-end.

Event distribution

Comparing the multi-modal version of the voting application to the single version of it, we notice two main differences in the way they function, both of which have implications on fundamental procedures of the logger. First is the existence of client identity, which the logger needs to take into account both while tracking/storing events, as well as when

events need to be replayed, since now the sources are multiple and dynamically added to (or removed from) the participant group. This identity information, whether it is explicit (user-defined credentials from the application) or implicit (a unique identifier that is automatically generated) needs to be reflected on the event stream. In other words, the identifier of the source needs to be carried along with the event information that this source produced.

In the concept application we built, the explicit identity of the source is not a requirement, meaning that, from a UX perspective, a user can place their vote in an anonymous manner. However, given that it was crucial for the correct behavior of the event track & replay to define & maintain such identity, we did so in an implicit way. After the connection establishment between client & server, the former generates a random identifier which it saves in a cookie and communicates it to the server, and from that point forward this identifier is sent together with every event that the client registers to the logger.

The second property of the multi-modal scenario which requires special handling by the logging system is the very distributed nature of the application, and it mainly resides on the replay phase. After a timeline has been completely logged and is requested to be replayed, that request can be done by the means of an API call to the back-end server. The logging system is designed in such way that replay request only needs to happen once, and as a result the timeline should be distributed to all the participants, with two important requirements:

- each participant should only receive & replay their own events to ensure the correct behavior of the application, and
- since the replay request only happens once by (any) one user, the rest of the participants should be able to have their respective events pushed their way, without explicitly asking for them.

The mechanism to address timeline distribution is facilitated in a lower communication level than AJAX, with the help of WebSockets. More specifically, once the API request for replaying a timeline hits the server, the respective events get retrieved from the database and, via the aggregation mechanisms of MongoDB [11], they are grouped by client identifier and pushed to their respective socket. On the client-side, code for handling that specific socket event gets triggered, and the logger proceeds with queuing the events to be fired. The queuing process takes each event's relative timestamp into account, which indicates how much time should pass before each event gets triggered, and depends on

when did the previous event occur. Since each event stream moves along a universal time signature, assuming (for the sake of simplicity) that there are no delays between client-side & server-side one can be fairly sure that the events will get triggered by all participants in a synchronized fashion, and the universal time signature that exists in the relative timestamp scheme of all events will produce the correct application behavior. Later on in this thesis, we consider and reflect upon scenarios of the not-so-perfect -and very real- world of communication with delays and other inefficiency threats.

The cross-domain AJAX problem

Looking back at the initial, single-modal application we started this project with, the communication scenery comprised of layers which were all local to each other, in the sense that application, middleware & back-end all belonged to the same codebase, therefore resembling an environment of absolute communication locality. All kinds of client-side scripting & communication (including remote API & WebSocket calls) would proceed & complete normally, without any security objection by the HTTP (in application-layer) or TCP (in transport-layer) protocols.

However, now that the layers are separated from each other existing under different domains and/or port numbers, client-side communication is no longer local. This change of scenery brings the result of all client-side API & WebSocket calls getting restrained by the HTTP Same-Origin Policy [12] and thus being rejected as non-secure communication.

Such issue revealed a major flaw in the communication scheme & architecture of the logger, and had to be tackled. In order to address this issue, we enabled cross-domain AJAX & WebSocket requests by using the Cross-Origin Resource Sharing (CORS) mechanism [13]. CORS requires specific HTTP response headers to be set server-side, by which the server indicates what types of cross-origin operations they permit, and in what depth. For the operations required by the logging system, we had to set three response headers:

- the **Access-Control-Allow-Origin**, indicating which specific domain and/or port is allowed to make requests,
- the **Access-Control-Allow-Methods**, which indicates the HTTP methods that can be used during the actual request, and
- the **Access-Control-Allow-Headers**, specifying the HTTP header that a cross-origin request can contain.

The CORS headers were set both in the middleware and the back-end layers, since the communication flow would pass from client-side to middleware to the back-end in the same, AJAX-based or WebSocket, manner.

Maintaining a global state in the server

Every system which is operating under such an ad-hoc environment must fulfill the imperative requirement of maintaining some sort of a global state structure. In Computer Science, the term *global state has been used* for a big variety of concepts and from many different standpoints, with perhaps the most profound one referring to the global execution state of processes of distributed systems [14].

For this project, by the term “global state” we refer to the global communication information of the application participants, and the global time signature of the event stream the system is logging. To store and maintain such information, we defined a **globalState** object in the back-end, which in the current design has three attributes:

- A client list, which is initially empty and is essentially an associative array and holds key-value pairs of which the key is the client identifier, and the value is the socket identifier for that client. We need to store this information because the socket identifier is a highly dynamic value which changes every time a specific client disconnects from the socket or reconnects to it after a possible network failure, and if this identifier is lost it becomes impossible for the logger to know where to forward the respective events this client generated, to be replayed.
- The last logged timestamp, which is initially equal to zero and makes the logger aware of the temporal status of the participants with regards to their real clocks.
- The last logged relative timestamp, which is initially equal to zero and is used by the back-end server to compute the temporal relationships of incoming events, using the relative timestamp of a previous event to calculate the offset of occurrence of the current event it is processing.

Maintaining such information in a central (or semi-central) fashion brings a considerable degree of simplicity in event processing & calculations, and also increases a fair amount of accuracy when time & synchronization decisions have to be made server-side.

Chapter 4: Performance & Efficiency Considerations

With the analysis of the previous chapter, we solidly established the fact that the logging system itself as well as the applications it aims to act upon belong to the group of the so-called distributed systems. And in such systems, characteristics such as scale, transmission rates and spatial distribution are most likely to manifest anomalies of various types, affecting performance & efficiency and leading systems to deviate from their design and expected behavior.

In this chapter, we will examine the expectations of how the logger should perform under a series of real-life scenarios with properties which are likely to pose threats on the system's performance, and try to identify the possible points of failure within its architectural components. We will focus on two main aspects that can be proven to be soft-spots of the system: scalability, both in terms of work load & input density as well as geographical distribution, and synchronization, which is an important metric regarding the efficiency of the logger, given the fact that we want to log, process & replay real-time events.

Scenario: High Event Density

Let us consider the following scenarios:

- The participants of a large-scale live concert use the collaborative voting system described in the previous chapter to express how appealing the performance of the artist(s) is to them, or how adequate is the quality of the audio-visual experience at the place where they sit/stand, all in real time. The information generated by each user's interaction is processed by a (possibly) centralized service, followed by a multi-cast type of propagation of this information to the other participants, perhaps based on a "focus-group" paradigm (e.g. participants share their impressions on sound quality with the ones located at the same general area as themselves).
- A number of people participate in a collaborative, multi-player Web-based game which is operated by the means of an advanced motion-capture controller device, such as Microsoft's Kinect sensor.

From the event logger perspective, both the aforementioned scenarios have one main characteristic in common, and that is the vast amounts of events being generated by the participants over relatively small time intervals. For the sake of context, it is appropriate to indicate that the Kinect sensor captures depth and color images simultaneously at a frame rate of up to 30 fps and the integration of depth and color data results in a colored point

cloud that contains about 300,000 points in every frame [15]. And if the logging system operates with a high event resolution, one can easily notice the increased work load for the system in terms of datastore operations (database writes & reads). Furthermore, in both the above scenarios we have some pieces of sensitive information carried along with each event, with the identity of the source (or the source group) being the most prominent. Such event properties bring additional complexity during the event processing by the logger and, with the event density being so high, the processing overhead could increase exponentially.

In front of such landscape, the logging system is challenged to be capable to perform well with increasing work load, or achieve work-load scalability, and thus be able to easily expand and contract its resource pool to accommodate heavier or lighter loads or number of inputs [16].

The impact of this scenario would primarily become most visible in the back-end of the logger's architecture, bringing immense load in the back-end server and especially database, where all the generated events need to be stored and, later on, retrieved. With the current configuration of the system's components, this block in the architecture needs to be thoroughly protected because it currently suggests a single point of failure, which is a hazardous property in any distributed system. The first (and the relatively simplest) measure to be taken for this scenario is to scale the back-end server in a vertical manner, that is to say to add resources to the node, typically involving the addition of CPUs or memory to it. This technique would increase the capabilities of the back-end & storage server for handling bigger loads in terms of processing, storage and memory allocation for the dynamic data-stores. The technique of horizontal scaling could be applied as well in this case, meaning that the processing & storing phases could be separated and being performed from two different servers, instead of having only one machine doing both operations. However, even in this case the database server should be a sufficiently powerful machine, especially with regards to memory size and disk space, if it has to store such huge amounts of data.

Furthermore, the table & document structure that is currently defined for event storage in our MongoDB instance would have to be modified in order to sustain vast amounts of data. Since the structure currently dictates that each user has their own document in the table, and all the events they produce are appended in this document (the MySQL equivalent would be similar to a row in a table in which columns are dynamically added), runs the risk of quickly exceed MongoDB's document size limit, which currently is equal to 16MB.

With such vast event volumes, one can safely assume that this limit will indeed be surpassed and from that point on, actions need to be taken from the programmer's perspective. Two approaches are identified:

- introduce an algorithm to wrap the document into a new one once the size limit is exceeded, or
- re-structure the table structure so that users and event streams are connected and related with each other.

The first alternative follows a simple approach which trusts the existing structure, however it has the long-term disadvantage of adding extra computation effort in the server for wrapping the documents, while maintaining consistency around the table and gathering them back again when they are being retrieved. On the other hand, the second technique requires more mental effort for coming up with a solution which re-defines the correlations within the database and takes advantage of the efficient and performant tools of the NoSQL system (which requires taking distance from the traditional way of thinking in the world of relational databases). However, albeit a fairly intensive mental task, our expectation is that such a technique would benefit the capabilities of the database when handling vast loads of data, both upon storage and retrieval. Regarding data consistency, one must in this case trust on the eventually consistent nature of NoSQL, meaning that transactions will not be present (which is for the best in this scenario, since the extra computation transactions need would be prohibitive on such data loads).

From the communication & synchronization aspect, given that we want to protect the (current) single point of failure of the architecture, there are a few techniques which can be applied, namely the middleware of the logger can be extended in such ways that it relieves some of the load from the back-end. More specifically, since each middleware node is dedicated to a small number of clients (if not only one), a layer of caching can be implemented on such a node, using some smart caching algorithm to store the generated events in a fast and light-weight, key-value based database system such as Redis and periodically synchronizing the cached events with the back-end server. The periodical communication between the caching modules of the middleware and the back-end server could be done either by having the back-end poll the middleware in a manner of status maintenance, or it could happen in the opposite manner, with the middleware nodes pushing the data to the server either via a RESTful API call or via WebSockets. This way, storing events in the central datastore happens in a less rapid rate and retrieval may not even require communication with the back-end server.

Regarding temporal consistency between the clients and the events they generate, the middleware can also contribute on this front. In co-ordination with the back-end server, an instance of the Berkeley synchronization algorithm would notify the middleware nodes about the time in the server, thus creating a state of time offset on each node. This way, every middleware node would be synchronized with the back-end server's master clock, and they would calculate the relative timestamps of each event based on their respective offset which will be derived from that master clock.

Scenario: Geographically Distributed Interaction

In the previous section, we reflected upon the logging system's scalability & synchronization consistency potential against scenarios which had the specific characteristic of high density of events coming in the logger as input, creating an increased work load along the complete architecture stack. In this section, we will examine how the logger would accommodate and maintain the same qualities in cases of a different nature. We will look at scenarios which introduce performance & efficiency bottlenecks caused by common characteristics of real-life networking: the wide geographical diaspora of nodes & clients which immediately implies differences in transmission rates, network delays & possible transmission failures.

Let us consider the following scenarios:

- A real-time, Web-based, pan-European polling system, in which citizens of the whole continent are called to place their stance on certain matters concerning the European Community.
- A collaborative music synthesis Web-application, where users are located around the world and each of them controls virtual instrument, having their contribution add up to a general musical piece which is being synthesized in an ad-hoc manner.
- An augmented reality multi-player Web game, where each player takes a turn and "shows" a landmark with a cinematographic identity from the place they are currently located in the world, and the rest of the players try to guess on-the-fly what its identity is.

The common characteristic of the scenarios above is that the communication range of the application escapes the local-area network scheme and moves to a wider, possibly world-wide distribution of clients & nodes. Such wide geographic distribution introduces the most common bottleneck for any system that operates in a distributed manner, since the network heterogeneity with regards to network infrastructure (bandwidth, latencies etc) along

the participants imposes performance & efficiency threats. In that respect, the logging system by design has a significant advantage and an equally significant disadvantage, which needs to be addressed if we want the logger to scale reliably from geographical aspect.

The advantage of the logging system resides on the asynchronous communication scheme it uses. It is known that one of the main reasons why it is currently hard to scale existing distributed systems that were designed for local-area networks is that they are based on synchronous communication. In this form of communication, a party requesting service, blocks until a reply is sent back. This approach generally works fine in LANs where communication between two machines is generally at worst a few hundred microseconds. However, in a wide-area system, one needs to take into account that interprocess communication may be quite a few orders of magnitude slower. The logging system was designed to use JavaScript along its complete stack, thus following the natural asynchronous operational schemes of the language. From I/O operations to network communication, there is no place where some layer of the system would block its execution and wait, unless a part of some communication algorithm would strictly dictate such behavior. Therefore, extending the logging system to operate under a wide-area application (similar to the multi-modality scenario we talked about in the previous chapter) would not threaten its performance in the synchronous/asynchronous communication respect.

However, there is one “flaw” in the current design of the logging system which would prevent it to efficiently scale geographically, and that is the centralized fashion in which the back-end & storage server operates, which resembles a scenery of both centralized services and centralized data. Typically, geographical scalability is strongly related to the problems of centralized solutions that hinder size scalability. If we have a system with many centralized components, it is clear that geographical scalability will be limited due to the performance and reliability problems resulting from wide-area communication. It is obviously imperative that such issues need to be tackled, and hereby we place some thoughts on possible solutions for addressing them.

Distribution of services

A centralized server that needs to serve requests at a large geographic scale, quite easily becomes the bottleneck when the requests increase and originate from a wide range of locations. And for the logging system, although it does not block communication with the asynchronous scheme it uses, the differences in network configuration of each client may cause possible anomalies in event processing.

The first -and almost inevitable- measure to address such an issue is to scale the back-end layer horizontally, by increasing the number of servers and expose the event API in a more “localized” fashion, allowing clients & middleware to communicate with the back-end that is located closer to them, thus decreasing the possibility of bad performance due to high network latencies. Additionally, and given that the applications we log are collaborative, the API servers could function as a type of peer-to-peer community, using slick communication & intra-synchronization techniques such as gossiping algorithms for spreading the system state information or voting algorithms for load management. The above technique could also be applied in a second degree of depth, by extending the middleware to carry specific parts of the API and take even bigger advantage of the geographic locality they possess in the general architecture.

The goals of this approach essentially aim at making the back-end system more flexible, distributing the load along a larger number of servers while maintaining consistency regarding the incoming event operation requests & ultimately achieving communication transparency on the client-side while logging as well as replaying events.

Distribution of data

As far as creating a possibly failure-prone distributed system, centralized data is essentially as bad as centralized services, for the obvious reason that all data-related requests (which for the logging system are approximately 95% of the total request volume) end up in one place, creating great work load and over-saturating the connection links. Being a visible threat to the capability of the logger to scale geographically, it is an issue that needs to be tackled as well.

In a similar fashion with the distribution of services, the first step towards addressing the issue would be to scale out and add more machines at the data storage layer. The immediate implication for the logger’s architecture is that now the back-end part (API and event processing) is physically separated from the data-storage part, which as a side effect relieves work load by having the back-end server(s) handling & processing the event requests & responses, and the database server(s) taking care of the storage, retrieval & data manipulation operations.

With multiple database servers available, there are a number of techniques one can use to improve performance & geographic scalability. The most common practice in this case (and the one that is proven to work with good results) is replication, which defines a master/slave relationship between the database servers with the master(s) being the one(s)

responsible to keep the slave(s) up-to-date with the data that have been inserted/updated/deleted from the database. Replication is a very smart technique and can very much improve the efficiency and performance of the logging system when operating under an application of a large geographical scale, with the possible overhead of having to ensure data consistency, namely that all the replicas will as soon as possible receive the latest changes in the data, regardless of where they happened. With our design decision for a database being MongoDB, a NoSQL system without transactions (the most common tool for orchestrating consistency algorithms), we will have to trust the so-called eventually-consistent model it works with, meaning that the data will be consistent at some point in time.

Moreover, given that currently we have read & write operations in the logger being finely independent (event logging only involves writing in the database, and event replay only involves reading from the database), one could define groups among the database servers, with one group receiving & performing all the write operations (probably the masters in the replication scheme) and a second group handling all the read operations (probably the slaves in the replication scheme). Of course, such groups would need to be defined with certain criteria in mind, most important of which is the locality requirement, which dictates that the distributed back-end servers we described in the previous section should be neighboring with database groups which contain both read-only & write-only, thus providing the clients with access of the complete feature set (event logging & event replay) from their neighborhood. Such technique is expected to actually assist in maintaining data consistency by introducing a strongly recognizable group of replication master servers, which could communicate in a fashion similar as the one we suggest for the back-end servers in the previous section as well as keeping their neighboring slave replicas updated. It does not bring any serious additional overhead to the overall replication mechanisms (the data need to be propagated to all the replicas anyway), instead we expect that it decreases the distribution of load at one more order of magnitude, while in the meantime it takes advantage of the locality created by this grouping.

[Figure 6](#) summarizes in a schematic manner what we discussed in this and the previous section.

Synchronization

For every system that operates in a wide geographical scale, managing time can -and, most likely, will- be a non-trivial problem. And it has been made quite clear in this thesis that for the logging system, working with real-time events, timing & synchronization plays a

crucial role. In all the application scenarios we mention in this section, the users can be located anywhere in the world, which implies that their systems will not be synchronized and their clocks will probably be set to a different timezone, with both these facts being reflected to the events they produce. Thus, the aim is for the logging system to store & retrieve the event streams in the right order by using a unique & reliable metric to determine their temporal relationships and calculate their relative timestamps.

Having set the scenery, our first intuition is that it would be too much overhead to force synchronization exclusively on the client-side, which instead needs to be achieved between the middleware & back-end layers, with the clients making as minimal an effort as possible.

Starting from the back-end servers, the last layer before forwarding the events for storage, absolute synchronization is crucial. The servers resembling this group could be synchronized with each other by using the Berkeley algorithm, with an elected server being the master (and getting synchronized by a universal clock via the NTP protocol) and the rest of the servers being the slaves and getting synchronized with the clock of the master.

After ensuring synchronization in the back-end, the servers would communicate their time (which will hopefully be the same along the back-end layer) with the middleware, and the middleware machines would then maintain the time offset based on their local clocks. Moving up to the client-side, a snapshot-based algorithm could be used in the sense of having a timer counting while events are being produced and communicating that counted time with each logged event to the middleware. At this stage, the middleware will have enough information from both the client-side and the server-side to compute the correct relative timestamp for each event, always based on the time offset of the back-end.

During replay, the relative timestamps of the event stream (and therefore the order & temporal relationships of the events) are already set, therefore the only synchronization aspect to be taken into account is synchronizing the replay of the distributed timeline at the client-side. For example, in the music synthesis scenario we mention above, if a timeline starts with the piano user playing first and the drums user coming second and we know that the drums user has a better network connection than the pianist, we still want the piano to come first during replay regardless the fact that the drummer would get their events faster. This can be achieved by having the middleware exchange a series of control messages between their connected clients and each other in a fashion similar to the TCP handshaking mechanisms, thus ensuring that all sides are ready to start replaying.

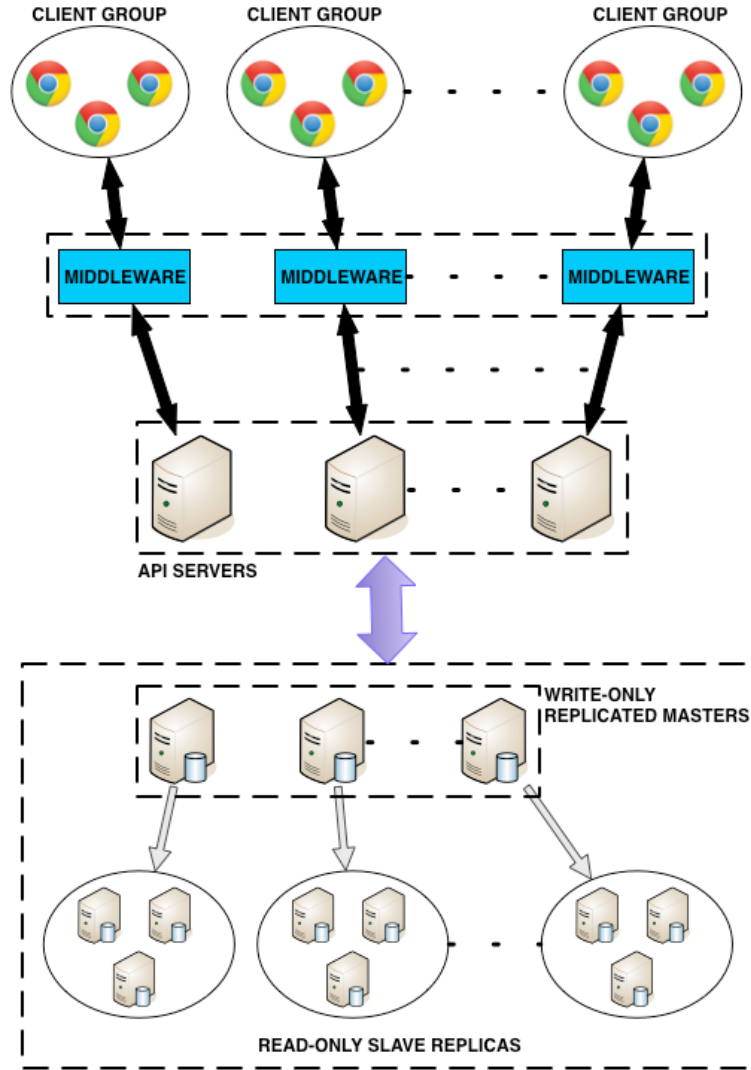


Figure 6: An abstract schematic representation of the architecture of the logger with a rich middle-ware layer, distributed back-end servers and a layer of replicated & grouped database servers.

Chapter 5: Towards A Fully Generic Event Logger

In this chapter, we want to explore concepts which would take the logging system from the prototype-like form it is now and transform it to a more generic & reliable system, while improving its scalability, efficiency & performance capabilities where possible.

In principle, we expect the logger to evolve towards the point of being fully generic by accommodating three main qualities:

- **Application agnostic:** the logger should be configured support different types of applications regardless their modality or design.
- **Event-type agnostic:** the logger must be able to log, process & replay any type of events the application usage could generate.
- **Context aware:** the logging system should be able to log events that are not necessarily created by direct application usage but occur as side-effects in the context of the usage itself.

Of course, the above points may resemble a somewhat utopian expectation set, however we believe that there are techniques & tools one can experiment with and hopefully move towards one or more of the directions we present here. We will elaborate on such techniques & tools for the remainder of this chapter.

Extending The Event Format For Complex Event Structures

It is a widely acceptable fact that interaction with computer platforms (including the Web) has evolved dramatically during the past decade, adding multiple and much more exciting ways that allow a user to interact with a computer, apart from the classic keyboard & mouse paradigm. The implication of this statement is that with new controller devices, new event types are defined in a programming level (e.g. the skeleton movement events generated by a Microsoft Kinect device). Furthermore, with computer software becoming more complex and programming methods such as concurrency & parallel execution becoming more widely used, one can expect that parallel or non-deterministic subsequent event occurrences become more visible within applications' execution cycles. Therefore, a logging system that aims to be generic needs to catch up with such concepts and become able to incorporate and adapt with more complex event structures and streams.

The logging system we built for this project, considering the Web platform on which it operates, indeed has such a potential. This potential mainly resides in the fact that it is built

with JavaScript which provides sufficient reflection capabilities to log client-side non-determinism in standard JavaScript running on unmodified browsers [17]. With this insight in mind, the logger could be extended to capture “behind-the-scenes” activity such as the firing of timer callbacks and random number generation [18], or even network events such as asynchronous AJAX calls or WebSocket communication. Such capabilities of the language give additional potential to the logger, that of being aware of the application context which may react to an internal or external event occurrence. We look into depth on context awareness in the following sections of this chapter.

Introducing Reactive Programming Principles

Functional Reactive Programming, or FRP, is a general framework for programming hybrid systems in a high-level, declarative manner. The key ideas in FRP are its notions of *behaviors* & *events*. Behaviors are time-varying, reactive values, while events are time-ordered sequences of discrete-time event occurrences [19].

One can already notice the particularly intuitive matching of terms & keywords between FRP and the logging system: events & behaviors is what interests us in the context of this thesis, and in FRP the same notions exist and, although they are defined in a lower & more conceptual level, it seems to be a technique which would integrate in a natural way with what we tried to do in this project.

From a programming standpoint, there are two main matters FRP can appear useful with. First is the JavaScript callback issue on the client-side event logging mechanisms. It is well known that the asynchronous & event-driven way that JavaScript operates is very much based on the primitive construct of callback functions. For example, when a browser completely loads a document, a user clicks a button, or an Ajax request is fulfilled, an event triggers a callback. JavaScript’s functional nature makes it extremely easy to create anonymous function objects that you can register as event handlers [20]. Naturally, this way of handling events is the one we used for this project as well. However, if one strives for making the logger capable of generically logging any type of events from the applications, one must make a considerable amount of effort composing callback functions for at least any type of event group the logger might encounter, if not (in the worst case) just any type of event object that could be produced by a Web-application. This approach, although it could possibly work, makes for long, repetitive code which is hard to comprehend & maintain and ultimately it may even harm the generic nature of the event logging mechanism by essentially turning it to a tailored-down solution for the application. FRP

tackles this issue by focussing more on the behavior of the target of the event occurrence, and while monitoring this behavior it internally triggers handlers -which are regular, programmer-defined JavaScript functions which are internally treated as callbacks- to process the stimuli that changed the behavior of the target element, namely the event itself. And since the event is a core entity of FRP and it is therefore defined with excellent detail and flexibility, such technique would allow for logging more complex (possibly indirect) events in different rates and configurations.

The second, and maybe more important, matter where FRP would be proven helpful is the context awareness of the logging system. In a serious gaming environment [21], we are very much interested to record, replay & reflect upon not only the direct event which caused some behavior in the application, but any side effects of that event to the rest of the application, its contextual landscape. In fact, it is very common that the context which becomes affected by a certain event occurrence is much more useful and insightful for many types of analyses. With this unique FRP feature of monitoring the behavior of an element that is a possible event target, one could enforce context awareness at configuration time, setting the logging mechanisms to watch for changes in the behavior of certain elements that are known to react in a side-effect manner when some event occurs somewhere else in the application environment. This approach would allow for very interesting post-hoc analyses with the purpose of identifying possibly complex causality patterns between event occurrences.

Smart Loggers

For the purposes of this project, the logging system was configured and built to act “stupid”, in the sense that there was no intelligence whatsoever during the record or the replay phase, with regards to conceptual manipulation of events. Essentially, the logger would just record the complete volume of events for a defined timeline, and it would eventually replay the same amount of events it recorded. Although this approach is effective for an initial stage, the logging system would definitely improve in many aspects (including the logger’s capability of being even more generic) if some degree of sophistication could be applied on the configuration phase of the logger.

An initial step towards making the logger a bit “smarter” would be the support of definition of rulesets during configuration, allowing for creation of implicit or explicit rules that could modify certain aspects of event logging & event replay, based on user-defined criteria. For example, if we consider the voting application we built as a testing environment for the

logger, one would want to record discreet alternating votes by all the users, meaning that the logger should ignore every event of a user after their first one, until another user generates their event. Such a requirement is absolutely valid if one wants to identify causal patterns between users with a bit smaller granularity or less noise in the events. Also, this example can be generalized as a configuration parameter which could define the resolution of events that are being logged, and it would potentially benefit the performance of the logger by decreasing the work load.

Another configuration feature which would be very useful is the support for declaring how critical time is when recording & replaying events. Coming back to the voting system example, it could be that one is interested only on the sequential patterns of event occurrences and not their exact timing information & temporal relationships. This feature would significantly improve the processing of the event objects in the middleware and the back-end servers in the logging phase, as well as the efficiency of the replay mechanisms.

Finally, and in combination with the FRP principles we described in the previous section, it would dramatically improve the depth of event logging if the logger would support defining a context within the application during the configuration phase. Allowing the user define the a set of elements or stimuli they are interested in makes for a purely generic method of creating a context-aware tool.

Chapter 6: Future Work & Conclusions

Coming to the final chapter of this thesis, before presenting the concluding thoughts for this project we want to briefly discuss about the potential of future work that would improve the logger and make it a useful tool.

Improving The Logger

So far, we discussed extensively the limitations & possible points of failure in the current design & implementation of the logging system, and we presented with ideas and insights regarding how to tackle each threat. To summarize the discussion, what follows is a short list of these limitations, together with the most immediate & direct improvement ideas:

- High event density: to prevent increasing work-load becoming a bottleneck, the first step is to separate the storage mechanisms from the API & event processing into two different servers, followed by some sort of event caching implementation in the middleware.
- (Semi) critical time ordering: in order to ensure proper event ordering and, therefore, correct behavior during replay for applications with time-critical executions, the most efficient solution would be to apply a snapshot-based algorithm [22] which would run between the clients and the middleware, and calculate the relative timestamps based on such algorithm.
- Global synchronization of timestamps: synchronizing between middleware and server-side would require a NTP-based synchronization scheme [23] which would initiate from the back-end server and update the time offsets of the middleware nodes, to ensure that the universal time signature of event occurrences is the same in all the layers.

The solutions proposed above are the first ones to be taken in order to tackle the current limitations of the logging system and make it perform & scale better.

Making Use Of Event Streams

An event logging system such as the one we built for this project is designed to produce output of big potential. Events triggered by user interaction or contextual reactions within an application make for an excellent dataset to be exploited during post-hoc analysis or data manipulation, with a wide variety of purposes. In this section, we explore this potential.

Analysis/visualization tools

When considering such a rich & large dataset, the most prominent action that comes to mind is that of *analyzing* the subject set, a method which can naturally lead towards multiple dimensions. In fact, analysis methods on event-oriented datasets has been a respectable field of study for a long time, initially engaged by the science of Mathematics [24].

What makes the output of the logging system such an excellent dataset for post-hoc analysis is:

- the high (and potentially configurable) granularity of events, which makes for a data sample of great accuracy, and
- the potential depth of information which is being carried along the recorded events.

From statistical analysis & visually extracted insights [25], to parametrized modeling [26] and behavioral pattern recognition in space & time [27], an event-oriented output can be used in many ways to extract insights with a technical, mathematical or even potentially sociological flavor.

Manipulation of time

Proceeding with a possibly more artistic approach, we will discuss the potential of manipulating the timelines produced by the logger in a more creative way. The current API of the logging system supports a set of basic operations (timeline selection, logging & replay), but it can be relatively easily extended with features which would make the recorded event streams more flexible, allowing to merge complete timelines or combine event selections into a new timeline, enabling post-hoc intervention in the temporal relationships of events etc. If we consider the collaborative music synthesis application we mention in Chapter 4, such flexibility in timeline manipulation would open up new horizons in the end result of that system, allowing the combination of both structure and randomness to create a richer artistic expression.

Expanding the users' conscience

Perhaps the most intricate aspect of working with the logger's output and replay functionality is the stage of user reflection upon their usage. Going beyond the technical dimension and moving towards more sociological & psychological concepts, the idea of allowing the user to see how they used a certain application and get stimulated to make thoughts such as "*Why did I act this way here?*" or "*How did I do this the wrong way?*", or even "*What would happen if at this point in time I did something different?*" can take the overall ap-

proach of application usage to a potentially very interesting direction. Looking at the matter from an optimistic standpoint, one can argue that such direction could cultivate a more educated generation of users, armed with awareness & responsibility for their actions and the effect they have in the context where they occur, very valuable qualities for people both as computer users and as members of the society.

General Conclusions

Revisiting the initial research questions of this thesis, we conclude that indeed there are ways to record & replay events in various environments & application schemes, from static monolithic single user applications, to distributed collaborative ones. Furthermore, the current technological tools that scientists & programmers have at their disposal allow to capture sufficient amounts of information depth regarding the event itself and the context reaction caused by its occurrence.

We chose to work with the ever-growing Web platform, basing this decision on the fact that the Web is perhaps the most widely used platform nowadays and there has not been very extensive research regarding event logging & replay mechanisms for this platform, with [28] being one -and perhaps the only adequate- exception.

To test and evaluate our prototype event logger, we built a concept Web-application representing a real-life scenario, and we built it in a way that it operates in a distributed, ad-hoc environment and is used in a real-time & collaborative manner.

We came across & tackled numerous technical & conceptual challenges, identified flaws and limitations upon which we reflected and considered a number of approaches to further improve the logging system in many aspects: performance, efficiency, scalability, post-hoc analysis & user education.

As a final thought, we argue that such a system can introduce a fresh look into data collection for computational analysis and behavioral research, and at the same time bring maturity to computer users.

Appendix A: Key Operations Code Snippets

Client-Side Event Registration & Triggering

```
function registerEvent(e) {
  postRequest(
    'event/new', //web service method
    {
      'data' : {
        'event' :
        {
          'type'      : e.type,
          'ts'        : parseInt(e.timeStamp),
          'target' : {
            'id'       : e.target.id,
            'class'    : e.target.className,
            'tagName'  : e.target.tagName.toLowerCase(),
            'textContent' : e.target.textContent,
            'type'     : e.target.type,
            'outerHTML' : e.target.outerHTML
          }
        }
      },
      'source' : cid //client id
    }
  ),
  null
};
}
```

```
function fireEvent(e) {
  var targetSelector = e.target.tagName;

  if (e.target.id != '') {
    targetSelector += '#' + e.target.id;
  }
  else if (e.target.class != '') {
    targetSelector += '.' + e.target.class;
  }
  //schedule event trigger
  setTimeout(function() {
    $(targetSelector).trigger(e.type);
  }, e.rts);
}
```

```
$("#button")
.not(".event-control") //exclude logger controls
.on("click", function(event) {
  if (event.isTrigger == undefined) {
    //native event -- send to server
    registerEvent(event);
  }
});
```

```
    }  
  }  
);
```

Enabling CORS In The Middleware

```
//enable CORS in middleware  
var allowCrossDomain = function(req, res, next) {  
  res.header('Access-Control-Allow-Origin', 'http://localhost:3000');  
  res.header('Access-Control-Allow-Methods', 'GET,POST');  
  res.header('Access-Control-Allow-Headers', 'Content-Type');  
  
  next();  
}
```

Server-Side Client-To-Socket Binding & Event Distribution

```
socket.on('client.bind', function(data) {  
  //bind client id with socket id  
  globalState.clients[data.client_id] = socket.id;  
});  
  
for (var i = 0; i < events.length; i++) {  
  //find socket id for current client id  
  var sid = globalState.clients[events[i]['_id']];  
  //emit timeline chunk of current client to their socket  
  io.sockets.socket(sid)  
    .emit('events.distributed', { data: events[i]['events'] });  
}
```

Appendix B: Journal

This section is sort-of a surreal (and incomplete) diary where I log some thoughts about the phases of developing my Master's Thesis project in a romanticized, anecdotal and literary manner. Enjoy and do not take the following text too seriously.

The Beginning

"Alright, let's do this.", that's what I thought. A system which will "keep track of time", it will know what buttons you pressed and it will do the same for you afterwards. It will replay your actions in front of you, and you will be able to see what went wrong with whatever it is you were doing. Sounds exciting, right? Maybe also a bit creepy for the more "Orwellian" ones (including myself..?). But there is no need to bring in the politics so soon, after all it's just an academic project, just my Master's Thesis. I finally got to this, after what seems to be an eternity of studying, working, and trying to find balance for everything else. And all I can think is: "let's get this done and get over with it".

But, where to start? I mean, I do want to get over with this but I still want to work for it, deliver something solid out of my efforts. And I certainly do not want to let my supervisor down, to whom I look up so much that, sometimes, I think whether I am just unworthy of collaborating with him on such a project. And with such things in mind, before I even began to answer that very first question, much more complicated (and, at the time, fairly pointless) questions started arising. How would it work? How would it be used? How would it be installed? Is it acceptable to have a user do this or that to use this system? Will the time I have be enough? How am I going to manage this when I have an almost full-time job at the same time? Chaos. The mind becomes blocked, incapable of comprehending anything related to the project anymore. Just, blank.

I was lucky enough to arrange that the second reader for my Thesis is a pretty technical person, and a very nice guy as well. So I went to see him, and discuss with him about some specifics of my project, hoping that he would help me solve all this technical madness I had inside my head. I explained him what my picture is with regards to how this system should work and what I figured that the scope was, and he would then point out risks, questions I should consider answering and several approaches for defining the scope of my project. Ultimately, he had a few valuable pieces of advice for me:

- try to think of the simplest case possible and work towards that first
- try to make it as generic as you can

- regardless of the implementation success, it is important that you explain in your report why you made the design decisions that you will make

So I decided to break it down and look at the building blocks of the system, the absolutely basic notions, the ingredients that would just make it work, for one simple case. And from there, we can see how it goes. After all, one needs to walk one step at a time in order to climb the mountain, and that's the way it is. Maybe this sounds obvious, but I had to repeatedly remind it to myself in order to realize it and make this chaos go away from my head.

Pencil & Paper

After I managed to let go of this first wave of crippling stress, I started asking basic questions. "What is an event?", that was the first question I had to answer, since my system will be an event logger. I started sketching out the event object, thinking about what properties it would have and what meta-properties would be useful to create for tracking it, saving it and bringing it back to be (re)played. And the most vital of event properties is of course its type, which indicates what sort of action triggered a specific event, and for my system it would mean the opposite (what sort of action will this event trigger). I knew that the system should to work with Web-based applications, so I had to define a starting pool of event types that would exist in this context.

The notion which popped up right after defining the event was the one where I have two events, one happening after the other. I decided to loosely name that one "timeline", basically a string of subsequent events. Immediately, the need of determining the temporal relationships between events became obvious. Thus, a new property was added to the event object: a relative timestamp which indicates how much time has passed between the occurrence of the previous event and the current one.

Alright, so I had the "soul" of my project in place, its basic conceptual building blocks. It was time to start building the "skeleton", and this meant that I had to pick up my pencil and draw some lines & boxes. Data flows, communication schemes, some naive architecture which demonstrates a hypothetical trip of events & timelines from one end to the other. A host application, one back-end server, maybe two back-end servers... I had to put my thoughts into practice. I needed a proof-of-concept to see if what I was sketching out was moving towards a good direction.

Bibliography

Tanenbaum, Andrew, and Maarten Van Steen. *Distributed Systems, Principles & Paradigms*. Pearson Prentice Hall, 2007

Chodorow, Kristina. *Scaling MongoDB*. "O'Reilly Media, Inc.", 2011

Teixeira, Pedro. *Professional Node.js: Building Javascript based scalable software*. John Wiley & Sons, 2012

Phan, Thi Anh Mai. "Cloud Databases for Internet-of-Things Data." 2013

References

[1] Joshi, Shrinivas, and Alessandro Orso. "Scarpe: A technique and tool for selective capture and replay of program executions." *Software Maintenance, 2007. ICSM 2007. IEEE International Conference on*. IEEE, 2007

[2] Crockford, Douglas. "The application/json media type for javascript object notation (json)." 2006

[3] Leavitt, Neal. "Will NoSQL databases live up to their promise?." *Computer* 43.2 (2010): 12-14

[4] *Ibid*

[5] "node.js." *node.js*. N.p., n.d. Web. 16 June 2014 <<http://nodejs.org>>

[6] Tilkov, Stefan, and Steve Vinoski. "Node. js: Using JavaScript to build high-performance network programs." *IEEE Internet Computing* 14.6 (2010): 0080-83

[7] "socket.io." *socket.io*. N.p., n.d. Web. 18 June 2014 <<http://socket.io>>

[8] "Firebase - Build Realtime Apps." *Firebase*. N.p., n.d. Web. 18 June 2014 <<http://www.firebase.com>>

[9] "AngularJS — Superheroic JavaScript MVW Framework." *Angular.js*. N.p., n.d. Web. 18 June 2014 <<http://angularjs.org>>

[10] "Firebase." *AngularFire* -. N.p., n.d. Web. 28 July 2014 <<https://www.firebase.com/docs/web/binding/s/angular/quickstart.html>>

[11] Bonnet, Laurent, et al. "Reduce, you say: What nosql can do for data aggregation and bi in large repositories." *Database and Expert Systems Applications (DEXA), 2011 22nd International Workshop on*. IEEE, 2011

[12] Barth, Adam. "The web origin concept." (2011)

[13] Van Kesteren, Anne. "Cross-origin resource sharing." *W3C Working Draft WD-cors-20100727* (2010)

[14] Chandy, K. Mani, and Leslie Lamport. "Distributed snapshots: determining global states of distributed systems." *ACM Transactions on Computer Systems (TOCS)* 3.1 (1985): 63-75

[15] Khoshelham, Kourosh, and Sander Oude Elberink. "Accuracy and resolution of kinect depth data for indoor mapping applications." *Sensors* 12.2 (2012): 1437-1454

[16] "Scalability." *Wikipedia*. Wikimedia Foundation, 22 July 2014. Web. 3 July 2014. <<http://en.wikipedia.org/wiki/Scalability>>

[17] Mickens, James W., Jeremy Elson, and Jon Howell. "Mugshot: Deterministic Capture and Replay for JavaScript Applications." *NSDI*. Vol. 10. 2010

[18] *Ibid*

[19] Wan, Zhanyong, and Paul Hudak. "Functional reactive programming from first principles." *ACM SIGPLAN Notices*. Vol. 35. No. 5. ACM, 2000

[20] Tilkov, Stefan, and Steve Vinoski. "Node. js: Using JavaScript to build high-performance network programs." *IEEE Internet Computing* 14.6 (2010): 0080-83

[21] Eliëns, Anton, and Zsófia Ruttkay. "Record, Replay & Reflect—a framework for under-

standing (serious) game play." *Proc. Euromedia*. Vol. 9. 2008

[22] Mattern, Friedemann. "Efficient algorithms for distributed snapshots and global virtual time approximation." *Journal of Parallel and Distributed Computing* 18.4 (1993): 423-434

[23] Mills, David L. "Internet time synchronization: the network time protocol." *Communications, IEEE Transactions on* 39.10 (1991): 1482-1493

[24] Lewis, Peter A. W. "A computer program for the statistical analysis of series of events." *IBM Systems Journal* 5.4 (1966): 202-225

[25] Marcus, Adam, et al. "Twitinfo: aggregating and visualizing microblogs for event exploration." *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*. ACM, 2011

[26] Yacoob, Yaser, and Michael J. Black. "Parameterized modeling and recognition of activities." *Computer Vision, 1998. Sixth International Conference on*. IEEE, 1998

[27] Liu, Fang, and Rosalind W. Picard. "Finding periodicity in space and time." *Computer Vision, 1998. Sixth International Conference on*. IEEE, 1998

[28] Mickens, James W., Jeremy Elson, and Jon Howell. "Mugshot: Deterministic Capture and Replay for JavaScript Applications." *NSDI*. Vol. 10. 2010