



Reality Sensing, Mining and Augmentation
for Mobile CitizenGovernment Dialogue
FP7-288815

D1.2 - Mobile Sensor App with Mining Functionality

Dissemination level:	PU - Public
Contractual date of delivery:	Month 27, March 2014
Actual date of delivery:	Month 26, February 2014
Workpackage:	WP1 - Reality Sensing and Mining
Task:	T1.2, T1.3
Type:	Prototype
Approval Status:	PMB Final Draft
Version:	1.0
Number of pages:	46
Filename:	D1-2.tex
Abstract In this deliverable we present several reality mining methods and an improved version of the sensor collection service. These methods allow the extraction of context information from sensor data collected on the mobile device. In particular we include a Human Activity Recognition component, a Service Line Detection component and a Traffic Jam Detection component. These components are implemented in a generic way and integrated into the Live+Gov toolkit.	
The information in this document reflects only the author's views and the European Community is not liable for any use that may be made of the information contained therein. The information in this document is provided as is and no guarantee or warranty is given that the information is fit for any particular purpose. The user thereof uses the information at its sole risk and liability.	



This work was supported by the EU 7th Framework Programme under grant number IST-FP7-288815 in project Live+Gov (www.liveandgov.eu)

Copyright 2013 Live+Gov Consortium consisting of:

- Universitt Koblenz-Landau
- Centre for Research and Technology Hellas
- Yucat BV
- Mattersoft OY
- Fundacion BiscayTIK
- EuroSoc GmbH

This document may not be copied, reproduced, or modified in whole or in part for any purpose without written permission from the Live+Gov Consortium. In addition to such written permission to copy, reproduce, or modify this document in whole or part, an acknowledgement of the authors of the document and all applicable portions of the copyright notice must be clearly referenced.

All rights reserved.

History

Version	Date	Reason	Revised by
0.1	2013-10-14	Outline	Heinrich Hartmann
0.2	2014-01-06	Added section on topic modeling	Christoph Kling
0.3	2014-01-07	Revised Outline	Heinrich Hartmann
0.4	2014-01-07	Alpha Version	Heinrich Hartmann
0.5	2014-01-21	Added section on Service Line Detection	Christoph Schaefer
0.6	2014-01-22	Added section on Distributed Geo Matching	Daniel Janke
0.7	2014-01-24	Added description of Sensor Collector	Heinrich Hartmann
0.8	2014-01-31	Added Traffic Jam Module Description	Laura Niittylä
0.9	2014-02-04	Added SVM Classifier Description	Elisavet Chatzilari
0.10	2014-02-06	Revised section on Human Activity Recognition	Heinrich Hartmann
0.11	2014-02-20	Incorporated feedback from internal reviews.	Heinrich Hartmann
0.12	2014-02-20	Final corrections and revision of figures.	Heinrich Hartmann
1.0	2014-02-25	Release of final version to the consortium.	Heinrich Hartmann

Author list

Organization	Name	Contact Information
UKob	Heinrich Hartmann	Phone: +49 261 287 2759 Fax: +49 261 287 100 2759 E-mail: hartmann@uni-koblenz.de
UKob	Christoph Schaefer	Phone: +49 261 287 2786 Fax: +49 261 287 100 2786 E-mail: chrisschaefer@uni-koblenz.de
UKob	Christoph Kling	Phone: +49 261 287 2702 Fax: +49 261 287 100 2702 E-mail: ckling@uni-koblenz.de
UKob	Daniel Janke	Phone: +49 261 287 2747 Fax: +49 261 287 100 2747 E-mail: danijank@uni-koblenz.de
MTS	Laura Niittylä	E-mail: Laura.Niittyla@mattersoft.fi
CERTH	Elisavet Chatzilari	E-mail: ehatzi@iti.gr

Executive Summary

In this deliverable we present several data mining methods that are used to extract context information from the reality of citizens. In addition we report on the ongoing improvements to our sensor collection architecture and communication interfaces.

The revised sensor collector provides improvements on recording performance on low end devices, battery awareness and communication interfaces. We present a newly designed sensor exchange format that allows continues streaming of sensor information and implement a streaming API for real time monitoring of mobile sensors.

We provide a Human Activity Recognition component that has been trained to detect multiple activities of human ambulation (sitting, standing, walking, running, cycling, stairs, lying). The classifier is deployed on a mobile device, integrated into the toolkit architecture and tested in the mobility field trial. Furthermore, we have trained several different machine learning models for this task and evaluated the results against each other as well as against state of the art classifiers. Our best classification accuracy remains with 84 % still below the reported accuracy of 97% in the literature. We are confident that we can improve our method quickly based on the built up infrastructure.

The Service Line Detection component is capable of detecting the precise service line used by a citizen based on the GPS sensor information captured on the mobile device. The detection is performed in a two-stage process that compares the received coordinates to real-time GPS positions of public transportation vehicles and in a second step to interpolated vehicle positions from time-table and vehicle track data. Also this service is fully integrated into the Live+Gov SaaS toolkit and has been successfully tested in the first field trial.

The Traffic Jam Detection component is able to detect irregularities in public transportation systems that provide real-time data about vehicles. It has been integrated into the Live+Gov toolkit as part of the Server Side Mining Service (C9). Although a thorough evaluation of the detection quality is difficult to perform, a comparison to other commercially available services shows promising results, with only minor differences.

The sensor collector as well as other mobile components remain limited to the Android platform. Despite several efforts, that are explained in Section 2.3, our prototypes based on a hybrid cross-platform development framework were not be able to deliver the intended results. Therefore, as described in D1.1, we have to fall back on native development of the individual components on multiple mobile platforms.

In Section 3.4 we present research results on Distributed Geo Matching, that allow to link social media posts (most importantly from Twitter¹) to vehicles of public transportation that they originate from. The approach addresses the problem of handling the high volume of incoming vehicle position updates in a distributed computing architecture.

We conclude this deliverable with a report on Geographic Topic Analysis in Section 3.5. The presented method is able to extract latent topics form reports of issues in the urban environment. These topics take into account information about category of the reported issue as well as geographic proximity. This new method was shown to be superior to current state of the art geographic topic models in the publication [7].

¹<http://www.twitter.com>

Abbreviations and Acronyms

API	Application Programming Interface
GPS	Global Positioning System
GSM	Global System for Mobile Communications
HTML	HyperText Markup Language
HTTP	Hypertext Transfer Protocol
JSON	JavaScript Object Notation
REST	Representational State Transfer
SVM	Support Vector Machine
URL	Uniform Resource Locator
UUID	Universal Unique Device Identifier
WP	Work Package
WIFI	Wireless Fidelity (IEEE 802.11), WLAN
WLAN	Wireless Local Area Network
XML	Extensible Markup Language

Table of Contents

1	Introduction.....	8
2	Sensor Collection Service.....	9
2.1	Mobile Sensor Collector	9
2.2	Sensor Storage Service	15
2.3	Cross-Platform Development	17
3	Reality Mining Methods.....	18
3.1	Human Activity Recognition	18
3.2	Service Line Detection.....	29
3.3	Traffic Jam Detection	32
3.4	Distributed Geo-Matching.....	37
3.5	Geographic Topic Analysis	42
4	Conclusion	45
5	References.....	46

List of Figures

1	Live+Gov Toolkit Architecture	8
2	Sensor Collection Architecture	10
3	Sensor collection testing GUI	13
4	Example rows of the ssf format	14
5	Sensor Storage Architecture	15
6	Schema of the Sensor Storage DB	16
7	Meta information of the sensor recordings	16
8	Raw data inspection view	16
9	Inspection Web Tool	16
10	Human Activity Recognition Method Overview	18
11	HAR Training Pipeline	19
12	Integrated HAR Architecture	20
13	An example of a decision tree.	22
14	Decision tree trained on manually computed features.	23
15	Hypothesis spaces for SVM classifier.	24
16	Maximization of the margin	24
17	The resulting hyperplane after training an SVM.	25
18	Illustration of the kernel trick for sample separation.	26
19	Non-linear separation of the data using the RBF kernel	26
20	Classification accuracy of HAR classifiers	27
21	The various input data for the HSL service line detection.	29
22	Visualization of the current traffic situation in Helsinki based on Jam detection .	33
23	Traffic Jam Module Architecture	34
24	An exemplary R*-tree	39
25	Distributing a geo-matching system	40
26	Document positions of reports with above-average probabilities for Topic 1-4 on the map of the Netherlands.	43

List of Tables

1	Supported sensors of the Sensor Collector	11
2	Structure of status intents	12
3	Number of accelerometer samples by activity and dataset.	27
4	Jam Detection Module Details	33

1 Introduction

The ambition of the Live+Gov project is to develop applications for closing the emerging gap between the citizens and the policy makers by the help of mobile technologies. These applications provide enhanced communication channels between the authorities and the citizens, which facilitate increased transparency, direct participation and opportunities for collaboration (cf. D2.1).

In order to meet this goals Live+Gov makes use of the advanced possibilities that are offered by modern smart-phone devices which are becoming widely available to the public. In particular, these devices are equipped with a rich set of sensors, that allow to infer information about the environment and context the citizen is currently placed in. In this deliverable we present several data mining techniques that allow to extract valuable context information from raw sensor data. For instance, we recognize the current ambulation mode of the user (e.g. walking or standing) in Section 3.1 and detect the use of public transportation in Section 3.2. These data mining methods are implemented as generic services and integrated into the Live+Gov toolkit as part of the Mobile Sensor Mining Component (C15) and the Server Side Mining Service (C8) (cf. Figure 1). The extracted context information can be used to personalize the user experience (cf. WP3) of the use case applications or analyze the behavior and service needs of the citizens.

In this deliverable we also provide an improved and extended version of the Mobile Sensor Collection Component (C14). The improvements contain in particular enhancements of the communication interfaces. We provide a new real-time streaming API as well as a revised sensor data format (ssf) cf. Section 2.1.3.

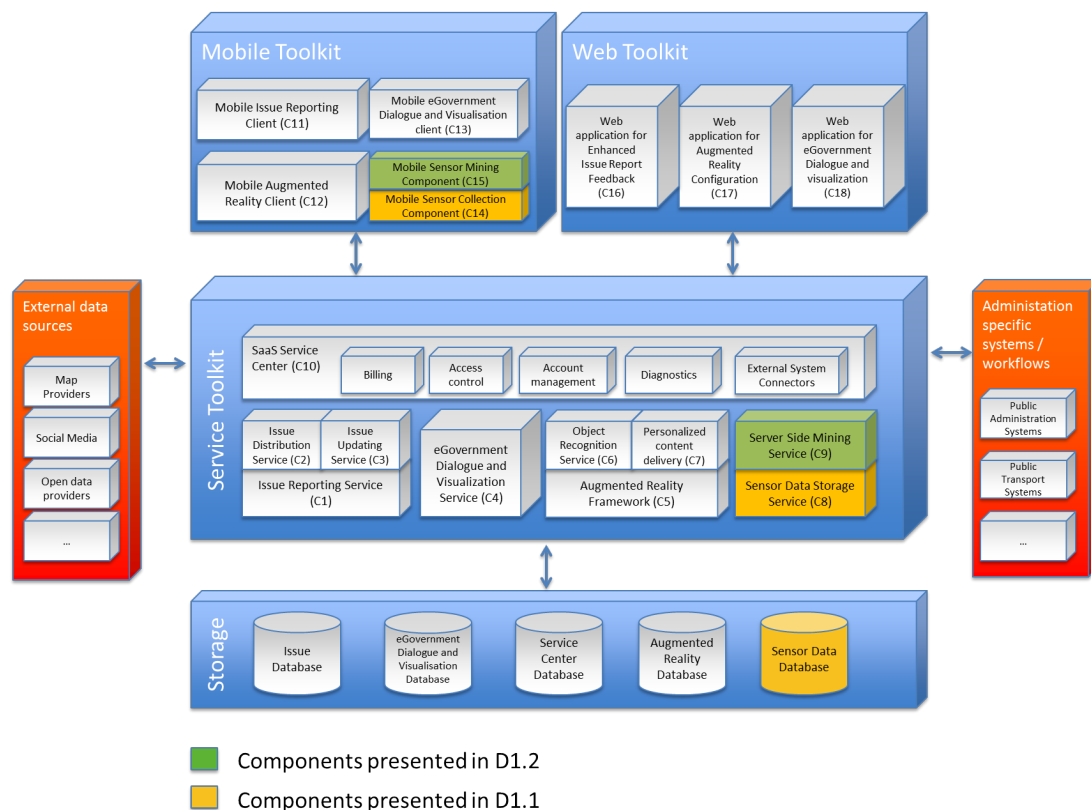


Figure 1: Live+Gov Toolkit Architecture

2 Sensor Collection Service

2.1 Mobile Sensor Collector

In this chapter we present the a revised sensor collection component. It offers the following improvements over the last version presented in D1.1:

High performance recordings on Low-End devices During the field trial some low-end devices showed difficulties to record the sensor samples with sufficient frequency. We streamlined the architecture in order to reduce the overhead caused by the processing of the samples. For example, we traded the convenience of a SQL database in favor of the performance benefits of a flat file solution.

Improved battery awareness Power consumption in the sensor collection process is largely driven the GPS sensor. In May 2013 Google released a new location API, that takes advantage of multiple location provider and reduces the power consumption significantly. The revised component makes use of this new API, and offers a fallback in case the service is not supported by the device.

Revised export format We have improved the data exchange formats used for publishing samples to other components and storage on the central server. In Section 2.1.3 we specify the ssf-Sensor Stream Format, which is inspired by the [Common Log Format](#) used by many web servers. It allows easy human inspection and processing by standard (UNIX) tools like “grep” and “sed” while being reasonably memory efficient.

Streaming API The streaming API is a new communication channel between the mobile sensor collector and the central server. It allows to transfer the incoming sensor data directly to the server in real time.

Integration into Live+Gov Service Center The revised sensor collector and sensor storage service is compatible with the Live+Gov service center. This feature allows central user management as well as health monitoring and centralized log aggregation of the services.

Seamless Extendability The revised architecture can be easily extended by further stream processing components. This is realized by a dispatcher thread that distributes sensor samples to registered components over a simple interface (again using ssf). Implemented extensions include the Human Activity Recognition component (c.f. section 3.1) and a GPS-sample publisher component that is used by the Service Line Detection service (cf. section 3.2).

2.1.1 Architecture Description

The new architecture is sketched in Figure 2. It allows asynchronous collection and processing of the collected sensor samples and can be easily extended by plugins which are executed by the dispatcher thread. It consists of the following individual components:

Sensor Event Thread The sensor event thread registers callback listeners for all configured sensors. When sensor events occur, the received values are serialized in the ssf format (cf. 2.1.3) and pushed onto the SensorQueue.

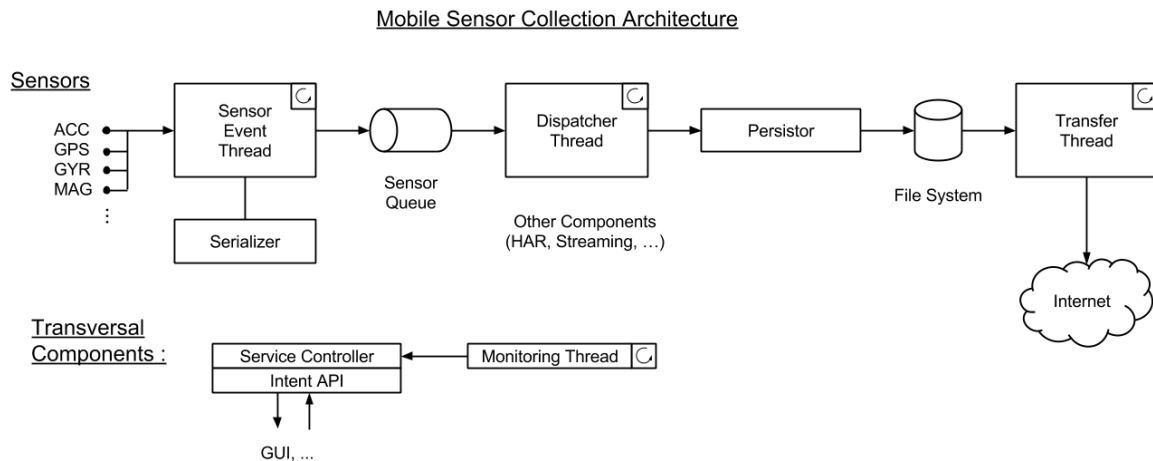


Figure 2: Sensor Collection Architecture

SensorQueue The sensor queue is a synchronized queue object that stores sensor values as strings.

DispatcherThread The dispatcher thread reads sensor values from the sensor queue and passes them to several services that can be registered. Included services are the persistor service and the activity recognition component, the streaming service and GPS publication pipeline (for use in the Service Line Detection API).

Persistor The persistor service is executed by the DispatcherThread and appends the received sensor samples into a predefined file. Also a variant with zip-file compression is included.

TransferThread The transfer thread is an independent thread that transfers the persisted samples to the Live+Gov servers.

Monitoring Thread The monitoring thread polls the different components and gathers run-time information like numbers of processed samples or the transfer state.

Service Controller The service controller starts and stops all threads, and takes care of proper configuration of all services. It listens to intents specified in the API and changes the state of the service accordingly.

The supported sensors are summarized in Table 1. They are unchanged from D1.1.

2.1.2 Intent API

The communication with other Live+Gov toolkit components is facilitated through the exchange intent messages³. The API was slightly improved since deliverable D1.1. We removed the explicit “Start/Stop Service” intents, since the service now shuts down automatically when no services are requested. The following list summarizes all provided intent controls. New intents are marked with an asterisk (*).

Start/Stop Recording Starts and stops recording of sensor samples. Samples are persisted in a local sensor file.

³<http://developer.android.com/guide/components/intents-filters.html>

Sensor	Description
GPS	Global position in latitude and longitude. If available, the GPS samples are gathered via Google's new Play Services location provider ² . If the Play Services are not available, the service falls back to accessing the GPS samples directly.
Accelerometer	Measures the acceleration force in m/s^2 that is applied to a device on all three physical axes. Also the low-pass and high-pass filtered variants 'Linear Acceleration' and 'Gravity' offered by the Google API can be captured.
Rotation Vector	Measures the orientation of a device by providing the three elements of the device's rotation vector.
Gyroscope	Measures a device's rate of rotation in rad/s around each of the three physical axes.
Magnetic field	Measures the ambient geomagnetic field for all three physical axes in μT .
WLAN	A list of all available wireless local area networks in the transmission range.
Bluetooth	A list of all available bluetooth clients in the transmission range.
GSM	Cellular network operator and the radio cell the mobile phone is connected with.
Google Activity	Google recently released an Activity Recognition Library. The sensor collector is able to record activities as sensor values as well.

Table 1: Supported sensors of the Sensor Collector

Delete Samples* Delete all stored samples on the device.

Send Annotation This intent allows users to annotate their recording by a string value that will be recorded by a “tag-sensor”.

Enable/Disable Streaming* When enabled, all collected sensor samples are streamed over to a central server over the network. See section 2.1.4 for more details.

Transfer samples Controls the sample transfer state machine. When enabled and a network connection is available sensor samples are transferred in batches to the server backend.

Enable/Disable sample broadcast When enabled all recorded sensor samples are broadcasted in the form of intents into the system. This allows other components, e.g. feature extraction and context mining, to make use of the recorded sensor data.

Request Status Report When this intent is received a status report is broadcasted to the system. A summary of returned information is shown in Table 2.

Set user ID* Set the user id to a given value. This intent is used to connect the mining results to the user data in the Live+Gov service center.

Get GPS samples* Returns a summary of the latest GPS samples gathered by the system in *ssf* format. It is intended to be used by the service line detection service described in Section 3.2.

When a control intent is received by the component, the appropriate action is taken. After execution has finished a status report intent is broadcasted to the system which provides information about whether the requested method call was successful. Further technical description of the component and the API are included in D4.1.

Key	Type	Description
running	boolean	True if service is running.
recording	boolean	True if sensor samples are recorded.
storage	boolean	True if samples are stored in the database.
transfer	boolean	True if continues transfer of samples is enabled.
broadcast	boolean	True if samples are broadcasted.

Table 2: Structure of status intents

Component Testing

The sensor collector comes accompanied with a testing GUI shown in Figure 3, that provides buttons for the implemented intent controls API. The returned status information is displayed in the form of plain text logs, button states and spinners. The GUI can be used to test the functionality of the component and for the gathering of training data. In addition we provide automated unit-test cases for the intent API along with the source code.

2.1.3 Sensor Stream Format (*ssf*)

The sensor stream format (*ssf*) which is used for transferring sensor values from a mobile device (or another sensor source) to the Server Side Storage Service (C9). The file format is inspired

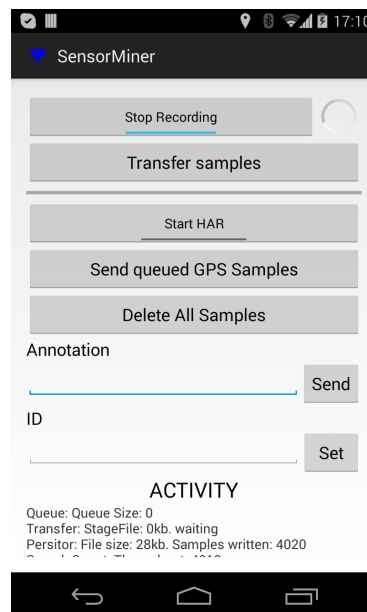


Figure 3: Sensor collection testing GUI

by the [Common Log Format](#) used by many web servers. It replaces the XML-based format used in D1.1. We view incoming sensor values as events and record them as a simple stream of CSV-rows. Each sensor has an individual prefix but writes into the same file.

The key advantages of this format as opposed to the previously used XML format are:

- Human readability. The format is easy to understand by humans.
- Flexibility. Every line can be interpreted without the context of the file. Therefore ssf-files can be arbitrarily sliced, concatenated and filtered using UNIX tools like *grep*, *cat*, *sed*, *awk*.
- Streaming support. Lines can be individually transferred over tcp sockets or messaging systems. Allowing immediate inspection of the recorded samples on a remote system.

The rows of the ssf-format have the following structure in ABNF⁴:

```
SSF = SENSOR_PREFIX "," TIME_STAMP "," USER_ID "," SENSOR_VALUES

SENSOR_PREFIX = "GPS" | "ACC" | "LAC" | "GRA" | "GYR" | "MAG" | \
               "WIFI" | "BLT" | "GSM" | "ACT" | "ERR" | "TAG"

TIME_STAMP    = *DIGIT
USER_ID       = *(ALPHA | DIGIT)
SENSOR_VALUE  = *VCHAR
```

⁴<http://tools.ietf.org/html/rfc5234>

The individual fields have the following semantics:

- **SENSOR PREFIX.** Identifies the sensor producing the sample. The following sensor prefixes are supported:
GPS (GPS sensor), ACC (Accelerometer), LAC (Linear Acceleration), GRA (Gravity), GYR (Gyroscope), MAG (Magnetometer), WIFI (Wifi networks), BLT (Bluetooth), GSM (GSM cells), ACT (Google Play Services Activity), ERR (Error Value), TAG (Annotations entered by the user).
- **TIME STAMP.** UNIX timestamp in milliseconds e.g. 1377609577214
- **USER ID.** An ID that identifies records from the same user. The default value is the device-id that is provided by Android.
- **SENSOR VALUES.** Sensor values in individual formats. In order to adhere to the CSV standard no , -symbol and new-lines may be used in this field.
 - ACC/GYR/MAG/LAC/GRA x,y,z-values separated by space characters.
 - GPS lat,lon,alt-values separated by space characters
 - ACT Activity Name, Confidence separated by space characters.
 - WIFI List of access-points, where each visible accesspoint is separated by a ; and written as
Escaped SSID String/Escaped BSSID String/Frequency in MHz/RSSI in dBm
 - BLT List of devices where each visible device is separated by a ; and written as
Escaped Address/Device Major Class/Device Class/Bond State/Optional Escaped Name/Optional RSSI
 - GSM The state of the device written as Service State/Roaming State/Manual Carrier Selection State/Escaped Carrier Name/Escaped Signal Strength followed by :, followed by a possibly empty list of cells, where each cell is separated by a ; and written as:
Escaped Cell Identity/Cell Type/RSSI in dBm

An example snippet from an ssf-file is provided in Figure 4.

```
ACC,1377605748123,5,0.9813749 0.0021324 0.0142523
GPS,1377605748156,5,50.32124 25.2453 136.5335
WIFI,1341244415,wifiUser,"WiFi AP"/"00:12:42"/2412/-45; \\
    "Another WiFi AP"/"33:13:53"/2437/-56
TAG,1378114981049,anotherUser,"test tag"
BLT,1385988380374,bluetoothUser,"C8:F7:33:B7:B5:B4" \\
    /computer/computer laptop/bonded/"LAPTOP"/-46
```

Figure 4: Example rows of the ssf format

2.1.4 Streaming Service

The revised version of the sensor collector includes a streaming service, that transfers recorded samples directly to a remote machine. The streaming service uses the ZeroMQ networking

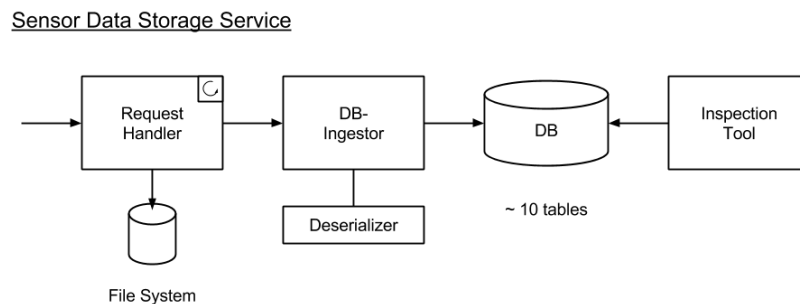


Figure 5: Sensor Storage Architecture

library⁵ to transfer single ssf lines to the Live+Gov machine running a streaming server. The streaming server appends the received samples to a file (`stream.ssf`).

This simple service allows direct monitoring of the recorded samples, via simple UNIX command line tools.

<code>tail -f stream.ssf</code>	prints the recoded samples to the console.
<code>grep <SENSOR PREFIX></code>	filters out samples of a specific type.
<code>logtop</code>	shows the frequencies of the individual incoming sensor values.

2.2 Sensor Storage Service

The sensor storage service runs on the Live+Gov servers and stores the sensor samples collected by the mobile sensor collection component. The storage component offers a HTTP/REST interface to upload sensor data. The uploaded samples are stored in a PostgreSQL⁶ database with PostGIS⁷ plugin for Geo queries.

The architecture is sketched in 5. It consists of the following components.

- **Request Handler.** Java/Tomcat Servlet that provides the HTTP interface for sample upload. The received samples are stored in the file system for backup in order to have a backup and passed to the DB ingester thread. When both operations where successful a response code "202 Accepted" is sent back to the client.
- **DB Ingester.** Stores sensor samples in a database. The samples are provided in ssf format and written to a PostgreSQL database with schema described in Figure 2.2.
- **Inspection tool.** This component allows to inspect the stored samples in the database in a convenient way and prepare and export them for usage in the data mining task. Figure 9 shows a screenshot of two views of the front end. We are currently expanding this tool to also include visualizations of the data mining results.

For integration with the Live+Gov service center, we provide an external service, that probes the upload URL in regular intervals and sends health checks and uploads log files to the central server as described in deliverable D4.2.

⁵<http://www.zeromq.org>

⁶<http://www.postgresql.org/>

⁷<http://postgis.net/>

```

sensor_gps (trip_id INT, ts BIGINT, lonlat GEOGRAPHY(Point));
sensor_accelerometer (trip_id INT, ts BIGINT,
    x FLOAT, y FLOAT, z FLOAT);
sensor_linear_acceleration (trip_id INT, ts BIGINT,
    x FLOAT, y FLOAT, z FLOAT);
sensor_gravity (trip_id INT, ts BIGINT,
    x FLOAT, y FLOAT, z FLOAT);
sensor_tags (trip_id INT, ts BIGINT, tag TEXT);
sensor_google_activity (trip_id INT, ts BIGINT, activity TEXT);
har_annotation (trip_id INT, ts BIGINT, tag TEXT);

```

Figure 6: Schema of the Sensor Storage DB

Trips Raw HAR						
Filter: [search this table]						
Trip ID	User ID	Start Time	End Time	Duration	Name	
515	mrar	2014/02/04 01:06:54	2014/02/04 01:06:54	01:00:00		X
514	mrar	2014/02/04 12:59:32	2014/02/04 01:00:09	01:00:36		X
513	mrar	2014/02/03 11:00:28	2014/02/03 11:01:16	01:00:47		X
512	mrar	2014/02/03 10:57:16	2014/02/03 10:58:27	01:01:11		X
511	2b1a8f3a6cb9f62	2014/02/03 07:59:47	2014/02/03 08:12:20	01:12:33		X
510	2b1a8f3a6cb9f62	2014/02/03 07:58:51	2014/02/03 07:58:56	01:00:05		X
509	43	1970/01/02 01:16:43	2014/01/31 08:33:49	08:17:06		X
508	43	1970/01/02 01:00:49	2014/01/31 08:07:51	08:07:02		X
507	43	1970/01/01 10:56:15	2014/01/31 06:46:06	08:49:50		X
506	42	1970/01/01 10:54:25	2014/01/31 04:31:11	06:36:46		X
505	48	2014/01/31 04:28:11	2014/01/31 04:37:36	01:09:24		X
504	39	1970/01/03 12:56:02	2014/01/31 07:12:28	07:16:25		X
503	62	2014/01/30 05:46:47	2014/01/30 06:24:25	01:37:38		X
502	62	2014/01/30 03:28:36	2014/01/30 03:45:24	01:16:47		X
501	43	1970/01/01 06:31:48	2014/01/30 04:13:41	10:41:53		X
500	40	1970/01/02 12:39:24	2014/01/30 07:36:53	07:57:29		X
499	39	1970/01/02 05:59:25	2014/01/30 07:10:01	02:10:35		X
498	43	1970/01/01 05:02:57	2014/01/30 07:50:05	03:47:08		X
497	43	1970/01/01 04:07:13	2014/01/29 05:42:44	02:35:31		X
496	48	2014/01/29 07:40:05	2014/01/29 07:57:50	01:17:45		X
495	42	2014/01/29 05:10:34	2014/01/29 05:39:21	01:28:47		X
494	43	1970/01/01 11:48:55	2014/01/28 10:23:39	11:34:43		X
493	43	1970/01/01 10:32:41	2014/01/28 05:18:14	07:45:32		X
492	40	1970/01/01 07:47:28	2014/01/28 05:27:05	10:39:36		X
491	HHL	2014/01/27 07:12:55	2014/01/27 07:26:09	01:14:04		X
490	46	1970/01/01 06:43:06	1970/01/01 06:43:06	01:00:00		X
489	43	1970/01/01 05:23:38	1970/01/01 05:23:38	01:00:00		X
488	43	1970/01/01 02:51:30	1970/01/01 02:51:30	01:00:00		X
487	43	1970/01/01 10:18:44	1970/01/01 10:18:44	01:00:00		X
486	60	1970/01/14 12:34:49	1970/01/14 12:34:49	01:00:00		X
485	60	1970/01/13 16:22:37	1970/01/13 16:22:37	01:00:00		X
484	60	1970/01/13 05:38:27	1970/01/13 05:38:27	01:00:00		X
483	41	1970/01/03 03:52:21	2014/01/24 06:53:35	04:01:14		X

Figure 7: Meta information of the sensor recordings

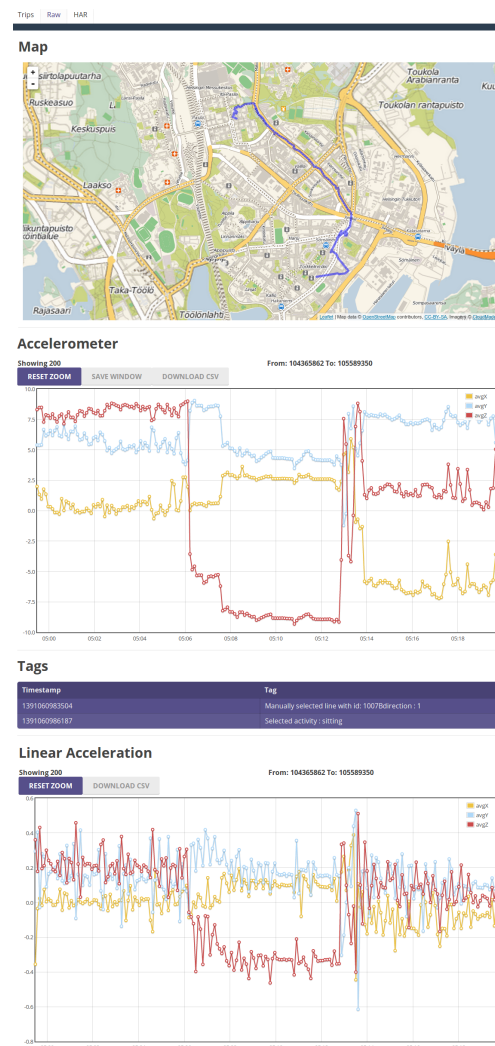


Figure 8: Raw data inspection view

Figure 9: Inspection Web Tool

2.3 Cross-Platform Development

In order to reach as many citizens as possible with our efforts, it is important to have the mobile components available on multiple platforms (e.g. Android, iOS). In deliverable D1.1 we have evaluated different methods for multiple platform sensor collection and reality mining. The basic three possible pathways are (cf. D4.1):

1. **Native development** of all components for all platforms.
2. **Web-based applications** exploit HTML5 and advanced APIs offered by mobile browsers to provide services on the mobile phone.
3. **Hybrid development** uses native wrapper frameworks like PhoneGap⁸ or Titanium⁹ that provide access to native platform and hardware features.

The evaluation criteria included necessary sensor sampling frequencies, native communication facilities and support for background services. It was found, that the Titanium framework was able to fulfill all three requirements. On the basis of this evaluation we proposed to implement the mobile reality mining services using this framework.

In due course we have started to implement prototypes of sensor processing applications, such as the *TitaniumSensorCollector*-package which is provided alongside with this deliverable. It turned out, however, that it is not possible to collect sensor data while running in the background with Titanium. Our prior evaluation covered both requirements, but it was not tested if it was cover both requirements at the same time. We were very surprised to learn that those requirements cannot be met in parallel.

For the practical use of the component, the support for sensor recordings in the background is very important. Indeed, the classification of human activities is performed while the device is worn in a pocket, in which case the display is typically switched off and the application is running in the background.

As stated in D1.1, we propose independent **native implementation** as fallback solution in this case. Our initial development is focused on the Android platform. Using a conversion toolkit we have been able to port some components (e.g. the sensor collector) to Blackberry. For support of iOS further development is necessary.

It has to be noted, that independent native development is associated with high additional effort, as the development environments and APIs are very different from each other. It has to be assessed if the resources of the project shall be spent on this task, or rather in improving the quality of the existing Android implementations.

In the view of these difficulties we have taken steps to reduce the amount of computation performed on the mobile device. In analogy to the image recognition task described in D3.1, it is possible to deploy parts of the logic on a central server. Currently our Service Line Detection component (c.f. section 3.2) is implemented in such a way, thereby paving the way for an iOS port of the Mobility Field Trial Application. A similar solution is planned for the Human Activity Recognition Component described in section 3.1.

⁸<http://phonegap.com/>

⁹<http://www.appcelerator.com/>

3 Reality Mining Methods

In this chapter we describe several reality mining methods that have been developed in the Live+Gov project.

The first three methods Human Activity Recognition (Section 3.1), Service Line Detection (Section 3.2) and Traffic Jam Detection (Section 3.3) are fully implemented and tested in field trials. The work on Distributed Geo-Matching (Section 3.4) and Geographical Topic Analysis (Section 3.5) are of more scientific nature and have already lead to a publication [7].

3.1 Human Activity Recognition

In this section we describe the implementation of the human activity recognition ('HAR') component on the mobile device. The algorithmic foundations of this data mining task have been described in detail in deliverable D2.2 in section 3.2 "Mobile Sensing and activity recognition". We include a brief summary here (3.1.1) for the sake of completeness. Subsection 3.1.2 contains the descriptions of the implementation. In subsection 3.1.5 we evaluate our method on two data sets and compare it to the state of the art approaches.

3.1.1 HAR Method Summary

The process of activity recognition uses a pipeline of signal processing and machine learning techniques. It consists of two phases: the "training phase" and the "integration phase".

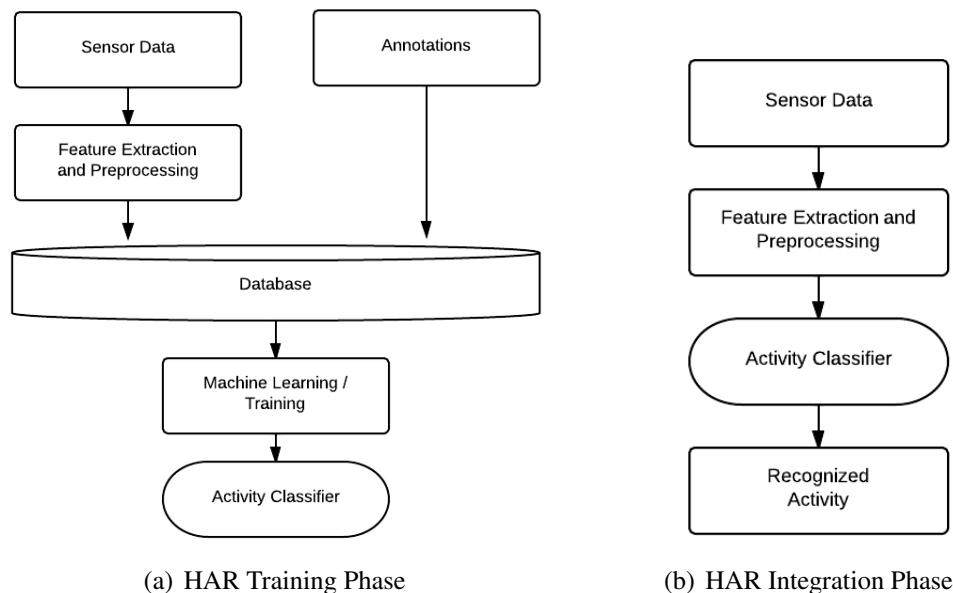


Figure 10: Human Activity Recognition Method Overview

In the training phase (cf. Figure 10(a)) a group of volunteers is asked to perform the targeted activities for a certain amount of time, while recording sensor samples with the device in their pocket. The gained training data stored in a database and used to train a classifier of the activities.

In the integration phase (cf. 10(b)), the trained classifier is embedded into the mobile device.

Both phases rely on the preprocessing steps of "windowing" and "feature generation". The stream of incoming sensor data is divided into time windows of fixed size (typically 1-10 sec.) and for each window a set of features is computed. These features are filtering out relevant information from the raw signal. Typical features include mean values and standard deviations as well as frequency domain features like Fourier modes.

Deliverable D2.2 contains detailed lists of all sensors and features and data mining methods that are used in the literature as well as discussions of quality of the recognition results.

3.1.2 Component Description

The Human Activity Recognition Component implements the two phases described in Figure 10. The training phase is executed offline on a server and uses the following processing pipeline (cf. Figure 11).

Classification Architecture

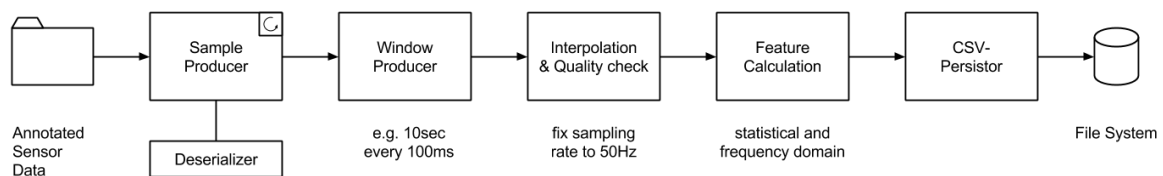


Figure 11: HAR Training Pipeline

1. **Annotated Sensor Data.** The training data is stored as ssf files in the file system. The annotations are represented as folder structure. In this way it becomes very easy to add new training data to the repository. Using the inspection web tool, one can select appropriate time slices and export the corresponding samples as “.ssf” file. The exported file is then moved to the folder corresponding to the recorded activity.
2. **Sample Producer.** The sample producer service reads the sensor samples from the file system and pushes them onto a generic sample processing pipeline. The interface is designed a way that resembles the incoming data stream on the mobile device.
3. **Window Producer.** The window producer takes as constructor parameters a window-length and a delay. Incoming sensor samples are grouped in windows of the given length and passed further down as a single object.
4. **Interpolation and Quality Check.** The recorded sensor data can be subject so several irregularities, like frequency disparities or small outtakes. To ensure that those irregularities do not pollute the training data we perform quality checks on the individual windows and interpolate the samples to a constant frequency of 50Hz.
5. **Feature Calculation.** The classification calculates different features extracted from windows of raw sampling data. A description of the used features is included in section 3.1.3.
6. **CSV-Persistor.** The calculated feature vectors are stored on the file system as CSV file.

In the next step the stored annotated feature vectors are used for training of machine learning classifiers. For this task we relied on available external tools like WEKA¹⁰ and libsvm¹¹. More details of the classification methods are given in section 3.1.4. The trained classifiers were then exported as java methods and integrated into the mobile HAR service (cf. Figure 12).

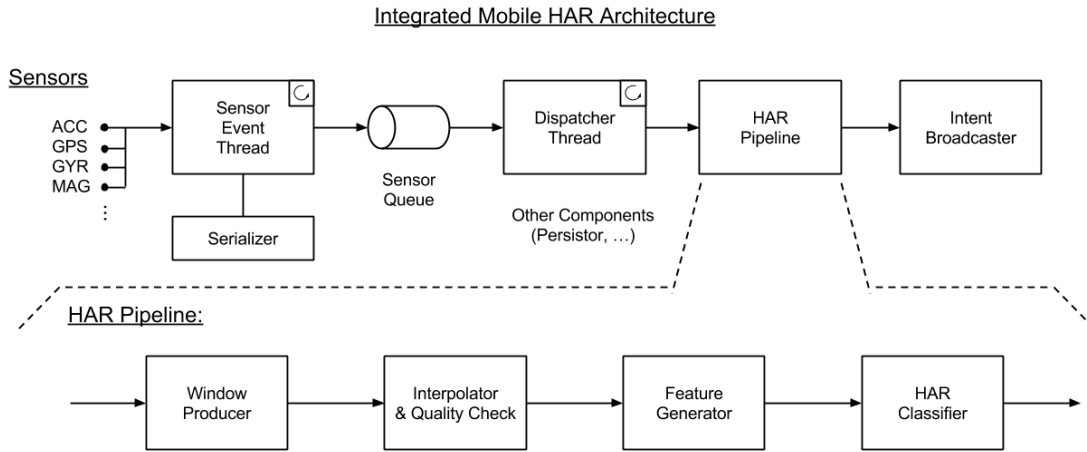


Figure 12: Integrated HAR Architecture

The integrated mobile HAR architecture consists of the following components.

1. **Dispatcher Thread.** The dispatcher thread is part of the sensor collection architecture described in Chapter 2. It emits all gathered samples in the ssf format.
2. **HAR Pipeline.** The HAR pipeline reuses most parts described in Figure 11. The Sample Producer is replaced by the Dispatcher Thread as source for sensor data.
3. **HAR Classifier.** The HAR classifier component warps exported java method and outputs the results of the classification as string objects. The current implementation uses a decision tree.
4. **Intent Broadcaster.** The classification results are broadcasted as an intent for use in other parts of the system (e.g. GUI, storage).

3.1.3 Feature Calculation

For the classification of activities, we cannot directly use the raw stream of sensor data. The reason is that machine learning classifiers cannot operate on data streams directly, but need to get explicit vectors as input. Moreover, these vectors cannot have arbitrary dimension. If we feed in the raw sensor signal of only 5 seconds recording, we get a vector of dimension

$$d = 50Hz \cdot 5sec \cdot 3 axes = 1000,$$

which is too much to handle for most classification algorithms.

¹⁰<http://www.cs.waikato.ac.nz/ml/weka/>

¹¹<http://www.csie.ntu.edu.tw/~cjlin/libsvm/>

With the calculation of features, we try to reduce the dimensionality of the input signal, while keeping the essential information contained in the signal. We use two different methods for feature calculation *manual selection* and an automated *PCA* based feature set.

Manually Selected Features

The following manual features have been selected:

Index	Name	Description
0	id	User ID of the recorded samples
1	tag	Name of the recoded activity
2	xMean	Mean values of the individual axes
3	yMean	
4	zMean	
5	xVar	Variance of the individual axes
6	yVar	
7	zVar	
8	s2Mean	Mean value of the length of the acceleration vector.
9	s2Var	Variance of the length of the acceleration vector
10	tilt	Average tilting angle of the device towards the vertical axes.
11	energy	Total energy of the recording.
12	kurtosis	Kurtosis ¹² measure of "peakedness" of the length of the acceleration vector.
13-24	S2Bins	Histogram over the length of the acceleration vector.
25-37	FFTBins	Histogram over the absolute values of the Fourier Modes of the length of the acceleration vector.

Our feature list has been guided by the literature (cf. D1.2, [9]). These consist of statistical features (mean and variance), that are computed for each of the three axes as well as for the length of the acceleration vector $s_2 = \sqrt{x^2 + y^2 + z^2}$. An explicit tilt-feature has been included to allow better separation of the "sitting" and "standing" classes.

We also include frequency domain features. The "energy" feature measures the total spectral density¹³ of the signal. Furthermore a set of histogram features measures the energy of different Fourier modes of the signal. We have divided energy spectrum into 10 equally distributed bins, covering roughly 95% of the observed energy spectrum in the training data. Two more bins added for outliers above and below the bin coverage.

PCA Based Features

There is a lot of freedom in the manual choice of feature vectors. Moreover, it is not clear that the heuristically selected features are able to separate the categories well. An automated way of selecting features is provided by the *Principal Component Analysis* (PCA) method.¹⁴ This method selects a number of linear features that maximize variance of the resulting data. We use a reduced vector size of 266 dimensions to keep 90% of the variance in our data set.

¹³http://en.wikipedia.org/wiki/Spectral_density

¹⁴http://en.wikipedia.org/wiki/Principal_component_analysis

3.1.4 Classifier Training

In this section we present more details about the employed machine learning methods for human activity recognition. Both methods rely on a training set of annotated feature vectors, and cannot be used directly on the data-streams. As was explained in D2.1 the two presented methods are the most popular approaches in the literature and have led to good classification accuracies in earlier work.

Decision Tree HAR classifier

A decision tree uses a hierarchy of simple binary decisions to classify a feature vector into a set of given categories. To be more precise, let $x_{\bullet} = (x_1, \dots, x_n)$ a feature vector with numeric feature values $x_i \in \mathbb{R}$, and $C = \{c_1, \dots, c_N\}$ a set of target categories. The individual decision rules R_k are of the form $(x_j \geq \alpha)$ or $(x_j \leq \alpha)$ for some index j and threshold value $\alpha \in \mathbb{R}$. A decision tree, organizes several of those rules R_k into a binary tree structure. The leaf nodes are linked to the corresponding target categories. Figure 3.1.4 shows an example of a tree with 9 decision nodes and 4 categories.

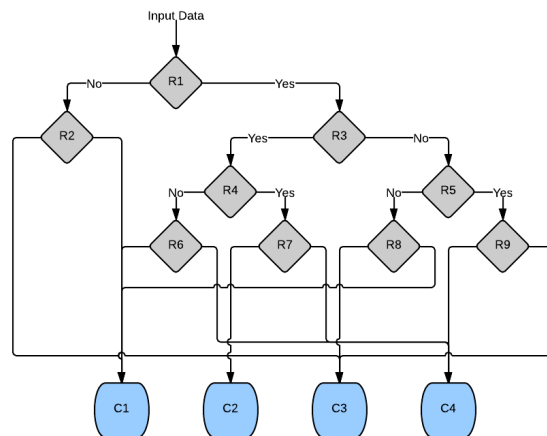


Figure 13: An example of a decision tree.

In the case of the activity recognition classifier, we have the different activities (e.g. standing, walking, running) as categories and a set of manually computed features like (standard deviations, or bin distributions of Fourier-modes) or automatically generated PCA-based features to base our decision rules on.

Of course, we do not want to build a decision tree for the HAR task by hand, but automatically learn a suitable tree on the basis of annotated training data. For this task of decision tree learning, several methods are known. Notable ones include¹⁵: ID3 (Iterative Dichotomiser 3), C4.5 (successor of ID3), CART (Classification And Regression Tree), CHAID (CHi-squared Automatic Interaction Detector).

We have chosen to use the C4.5 algorithm modified for continues attribute values [10] since it has been successfully used in the literature [9] and an open source implementation is available and integrated into the WEKA toolkit.

¹⁵http://en.wikipedia.org/wiki/Decision_tree_learning

The C4.5 algorithm uses the concept of *information gain* to build up a decision tree in a recursive, greedy fashion. For a set of annotated training vectors $T = \{(x_{\bullet}, c) | c \in C\}$ a decision rule R splits T into two subsets:

$$T = T^+ \cup T^-, \quad T^+ = \{R(x_{\bullet}, c) = \text{true}\}, \quad T^- = \{R(x_{\bullet}, c) = \text{false}\}$$

The information gain of this rule is defined as

$$IG(T, R) = H(T) - \frac{|T^+|}{|T|} \cdot H(T^+) - \frac{|T^-|}{|T|} \cdot H(T^-),$$

where $H(S)$ is the entropy of the collection $S \subset T$ cf. [11]. The C4.5 algorithm finds the decision rule with the largest information gain and uses this as the root for the decision tree. It then recurses into the subset T^+ and T^- and adds the generated nodes as childs to the root node.

The choice of the threshold values α for the decision rules is particularly crucial for the quality of the trained decision tree. A possible improvement of the threshold selection algorithm can be obtained by pre-clustering of the attribute values before the training process cf. [8].

Figure 14 shows an example of a trained decision tree on manually generated features.

```
J48 pruned tree
-----
Number of Leaves   :   52
Size of the tree   : 103

x18 <= 927
|  x24 <= 152
|  |  x20 <= 211
|  |  |  x4 <= -2.936108
|  |  |  |  x3 <= 1.132095
|  |  |  |  |  x18 <= 458
|  |  |  |  |  |  x35 <= 2: WALKING (6.0/1.0)
|  |  |  |  |  |  x35 > 2: CYCLING (2.0)
|  |  |  |  |  |  x18 > 458: SITTING (3.0)
|  |  |  |  |  |  x3 > 1.132095: CYCLING (527.0)
|  |  |  |  x4 > -2.936108
|  |  |  |  |  x4 <= 7.830095
|  |  |  |  |  |  x14 <= 13
|  |  |  |  |  |  |  x18 <= 87
|  |  |  |  |  |  |  |  x2 <= 3.434441
|  |  |  |  |  |  |  |  |  x3 <= -2.172633
|  |  |  |  |  |  |  |  |  |  x2 <= -3.51528: RUNNING (3.0/1.0)
|  |  |  |  |  |  |  |  |  |  x2 > -3.51528: WALKING (45.0/1.0)
|  |  |  |  |  |  |  |  |  |  x3 > -2.172633: CYCLING (5.0/1.0)
|  |  |  |  |  |  |  |  |  |  x2 > 3.434441
|  |  |  |  |  |  |  |  |  |  |  x19 <= 469: CYCLING (155.0)
|  |  |  |  |  |  |  |  |  |  |  x19 > 469: WALKING (4.0/1.0)
|  |  |  |  |  |  |  |  |  |  x18 > 87
|  |  |  |  |  |  |  |  |  |  |  |  x18 <= 415
|  |  |  |  |  |  |  |  |  |  |  |  |  x15 <= 17
|  |  |  |  |  |  |  |  |  |  |  |  |  |  x34 <= 60
|  |  |  |  |  |  |  |  |  |  |  |  |  |  |  x22 <= 138
|  |  |  |  |  |  |  |  |  |  |  |  |  |  |  |  x20 <= 98
|  |  |  |  |  |  |  |  |  |  |  |  |  |  |  |  |  x21 <= 58: CYCLING (10.0)
|  |  |  |  |  |  |  |  |  |  |  |  |  |  |  |  |  x21 > 58: WALKING (56.0)
|  |  |  |  |  |  |  |  |  |  |  |  |  |  |  |  |  |  x20 > 98
|  |  |  |  |  |  |  |  |  |  |  |  |  |  |  |  |  |  |  x17 <= 108

// TRUNCATED

x18 > 927
|  x10 <= 0.58851: SITTING (1071.0)
|  x10 > 0.58851
|  |  x3 <= -8.323328: WALKING (17.0)
|  |  x3 > -8.323328: CYCLING (15.0/2.0)
```

Figure 14: Decision tree trained on manually computed features.

Support Vector Machine HAR Classifier

A popular algorithm to train a binary classification model for mapping features to activities is the Support Vector Machines (SVMs). SVMs are known for their ability in smoothly generalizing

and coping efficiently with high-dimensionality pattern recognition problems. They define a hypothesis space that includes all the possible linear separations of the data (Fig. 15) and they choose the one that maximizes the margin between the two classes (Fig. 16).

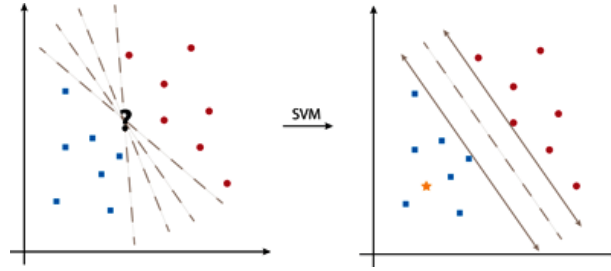


Figure 15: Left: Hypothesis space including all linear separations. Right: The selected hypothesis maximizes the margin.

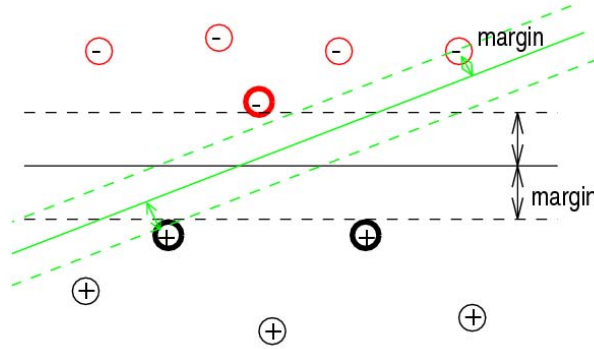


Figure 16: Maximization of the margin

The hyperplane that optimally separates the positive and the negative class (i.e. maximizes the margin) can be described by $\mathbf{w} \cdot \mathbf{x} + b = 0$, where $\mathbf{w}$ is normal to the hyperplane and $\frac{b}{\|\mathbf{w}\|}$ is the perpendicular distance from the hyperplane to the origin (Fig. 17). The optimal hyperplane can be obtained by solving the following Quadratic Programming optimization problem:

$$\min \frac{1}{2} \|\mathbf{w}\|^2 \quad \text{s.t.} \quad y_i(\mathbf{w} \cdot \mathbf{x}_i + b) - 1 \geq 0 \quad \forall i \quad (3.1)$$

In order to relax the constraints of Eq. 3.1 and allow for some misclassified points, a slack variable $\xi_i, i = 1, \dots, L$ is introduced which transforms Eq. 3.1 into:

$$\min \frac{1}{2} \|\mathbf{w}\|^2 + C \sum_{i=1}^L \xi_i \quad \text{s.t.} \quad y_i(\mathbf{w} \cdot \mathbf{x}_i + b) - 1 + \xi_i \geq 0 \quad \forall i \quad (3.2)$$

where the parameter C controls the trade-off between the slack variable penalty and the size of the margin. In the testing phase, in order to classify an unseen example x_t , its distance to the hyperplane is calculated using the formula $\mathbf{w} \cdot x_t + b$. This distance indicates the classifier's confidence that the unseen example x_t belongs to the examined class.

The previous consider linear separation of the data. However, this is rarely the case for most real world scenarios and for this reason the kernel trick has been introduced. Applying the kernel

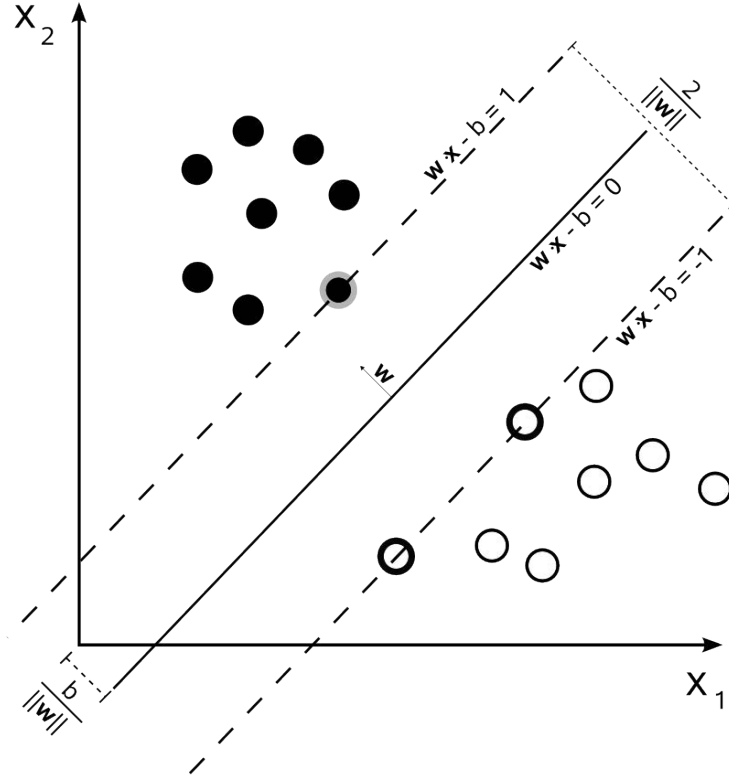


Figure 17: The resulting hyperplane after training an SVM.

trick to the cases where the classes are not linearly separable in the input feature space, we manage to map the features to a higher dimension (Fig. 18) where they can be linearly separated. For example, in the case of an RBF kernel the otherwise linear hyperplane is transformed to a hypersphere (Fig. 19). For any kernel $K(x, y)$, the classification model can be represented by a vector $\mathbf{w}$ (i.e. the model parameters), a bias scalar b and the support vectors $\mathbf{SV}_j, j = 1, \dots, N_{SV}$. For an unseen example x_t , a confidence score is extracted by computing its distance to the hyperplane of the model:

$$confidence = \mathbf{w} * \sum_{j=1}^{N_{SV}} K(\mathbf{SV}_j, \mathbf{x}_t) + b \quad (3.3)$$

In our case, an SVM model is trained on the previously extracted features in order to learn the properties that define the examined activity. The models are trained using the one versus all (OVA) technique, i.e. all positive examples of the specific activity versus all negative examples (i.e. the examples of all other activities). The distance of a vector from the hyperplane indicates our confidence that during the analysed window, the user performs the examined activity. High positive values of this score increase our confidence that this window belongs to the positive class while high negative values provide strong confidence that the performed activity is not the examined one.

3.1.5 Evaluation

We have evaluated our classifier on two different datasets.

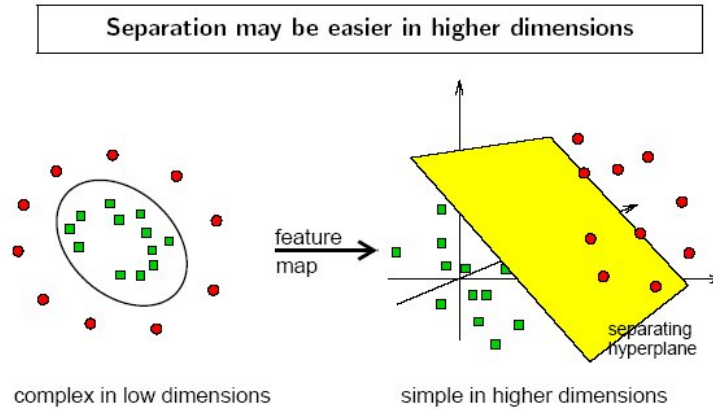


Figure 18: Illustration of the kernel trick for sample separation.

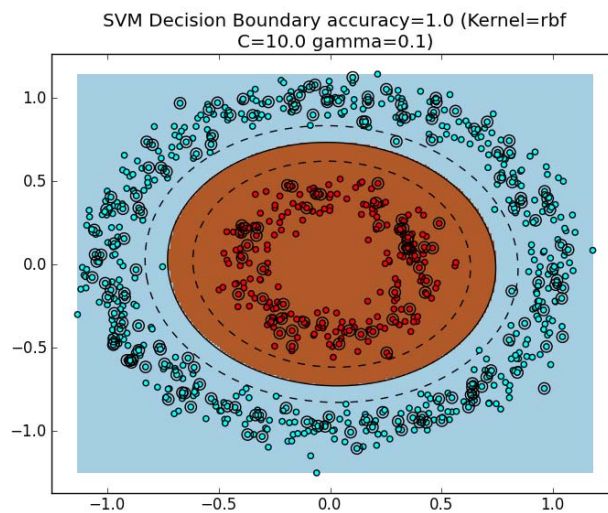


Figure 19: Non-linear separation of the data using the RBF kernel

The first dataset was gathered on the University Campus in Koblenz in December 2013. It contains a total of around 900K samples collected by 10 volunteers. The volunteers were instructed to perform the activities "walking", "running", "stairs" and "cycling" on predefined routes on the university campus. The total time effort per volunteer was about 20-25minutes and a financial reward was offered as an incentive. After the recording the samples have been inspected using our inspection tool and the beginning and ending of the activities were manually stripped in order to avoid noise from holding the device in the hand.

The other dataset was obtained from the *UCI Machine Learning Repository*¹⁶ and was gathered by Davide Anguita, et. al. [2] in 2012. It contains around 700K collected by 30 volunteers.

The number of samples per activity of both datasets are summarized in Table 3. Both datasets contain only accelerometer samples, and have been preprocessed that have been sampled at a fixed rate of 50Hz.

Both datasets have been split into a training data-set and a test data-set. The individual parts contain only full recordings of the activities. No single recording is present in both parts of the

¹⁶<http://archive.ics.uci.edu/ml/datasets/Human+Activity+Recognition+Using+Smartphones>

Activity	UCI Dataset	UKOB Dataset
sitting	113.728	80.951
standing	121.984	320.737
walking	110.208	292.024
running	0	31.916
cycling	0	436.106
stairs	188.800	30.086
lying	124.416	0
Totals	659.136	903.156

Table 3: Number of accelerometer samples by activity and dataset.

split. The volume is distributed in roughly 66/33 ratio between both parts.

We have three different parameters for the evaluation:

- **Classifier.** Decision Tree (DT) or support vector machine (SVM)
- **Features.** Manually selected (MAN) or PCA based (PCA)
- **Dataset.** Self collected (UKOB) ore reference dataset (UCI)

We get the following evaluation results:

Classifier	Features	Dataset	Accuracy in %
DT	MAN	UKOB	62
DT	PCA	UKOB	45
SVM	MAN	UKOB	87
SVM	PCA	UKOB	54
DT	MAN	UCI	83
DT	PCA	UCI	68
SVM	MAN	UCI	85
SVM	PCA	UCI	64

Figure 20 shows a plot of these accuracy values.

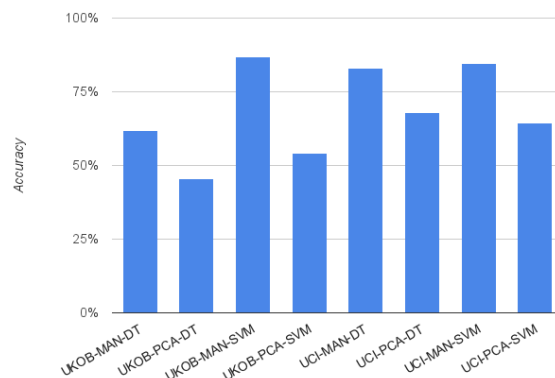


Figure 20: Classification accuracy of HAR classifiers

These figures show, that we are still around 10% away from the state of the art classifier used by [2] with 97% accuracy. There the authors use a multi-modal SVM on a similar list manually selected feature vectors and slightly parameters for window length and overlap. Our integrated application relies on a decision tree and manual features, which show significantly weaker performance (64%) than state of the art in our experiments.

For both datasets the SVM trained on manual features performs best. Although there are several parameters in the PCA feature extraction, that remain to be optimized, the evaluation results suggest that manually selected features yield better results. For the UCI dataset the decision tree is only slightly worse than the SVM based classifier whereas the UKob dataset shows a more significant drop of accuracy. This might be due to the fact that UKob dataset contains slightly different activities than UCI, e.g. a large number samples is recorded for the cycling activity, which is not present at UCI, and conversely no samples for the lying task are present.

In conclusion, we have presented here a working HAR classification architecture, that has been fully integrated into a mobile device and has been evaluated against state of the art classifiers from the literature on different datasets. Although the classification accuracies are still behind state of the art, we are confident that we can improve our methods exploiting the built up infrastructure.

3.2 Service Line Detection

One part of the user contextualization is the service line detection. The aim of this component is to recognize if a citizen uses public transport within the HSL area in Helsinki. If such a usage is detected, the right service line id with its direction and the further itinerary will be determined. Based on this information it is possible to approach higher level problems like network utilization analysis, personalized reports about traffic jams 3.3 and automated connecting train recommendations.

The service line detection is implemented as a server side component which provides a REST API for answering user queries in real time. For long term evaluations all API calls are recorded and send to the Live+Gov Service Center on a daily bases. Due to the seamless integration into the Live+Gov ecosystem, where each user gets an universal unique id, all received tracks are personalized and can be combined with additional information coming from other components. In a later analysis, the service line detection results can be augmented with the low level human activity recognition data of the same user to gain a better insight into the users behavior.

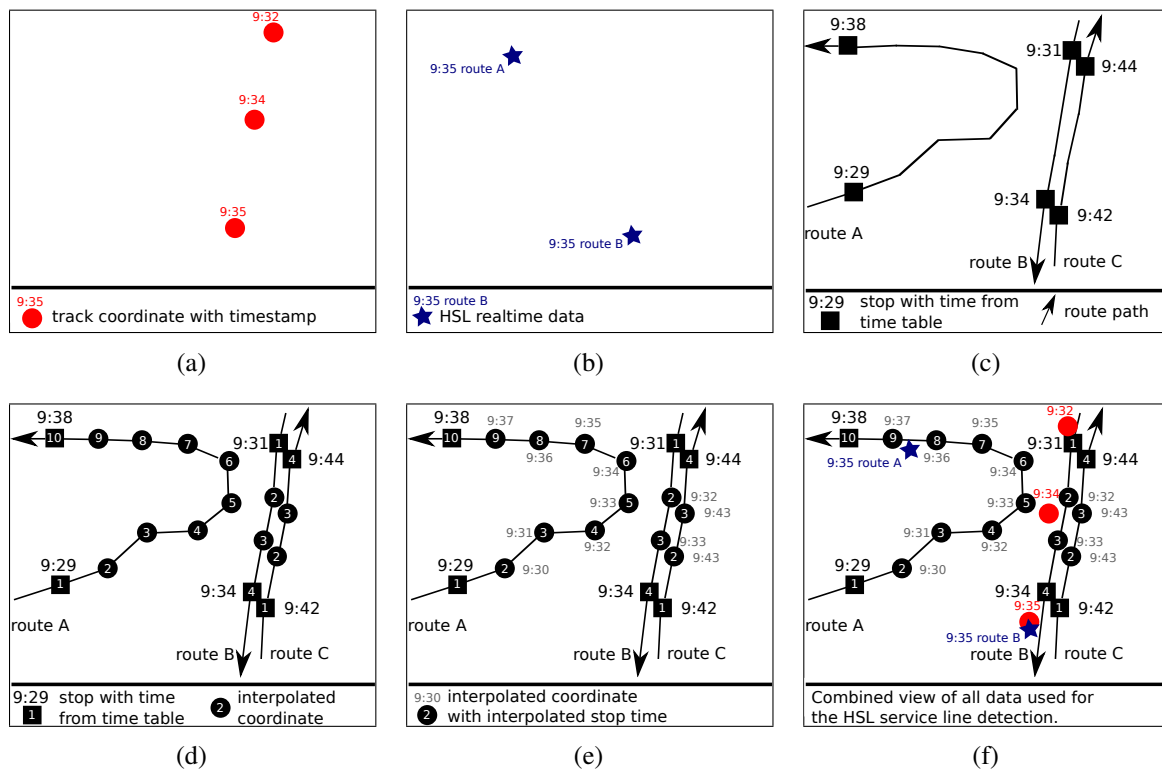


Figure 21: The various input data for the HSL service line detection.

Assigning a trajectory to a specific service line requires the access to suitable background knowledge. In the Helsinki case it is two fold: First of all, the system is fed with all HSL public transportation schedules and associated geographic information of stops and travel paths. By means of this dataset the theoretical position of every vehicle at any point in time can be calculated. A drawback of only including static time tables is that actual schedule deviations stay disregarded. This gap is closed by the HSL Live API, which is the second data source the service line detection relies on. With this API, the actual positions of some vehicles can be determined. Especially trams are tracked in realtime and can be accessed via this interface. For the majority of vehicles, such as buses, live data is not yet available. Lacking a complete

realtime coverage is why still some uncertainty remains in the overall system. Therefore, the main challenge is to combine both kinds of data to archive the best possible performance.

Just as the background data, the classification algorithm is also two-stage. Given an user trajectory containing a small set of GPS coordinates with timestamps (Figure 21(a)), in a first step the algorithm checks if the latest given timestamp is near the actual time. If so, the HSL Live API is called in order to find all vehicles which are currently around the user (Figure 21(b)). Service lines found below a certain distance to the user are scored with a high probability to be the true line belonging to the given track.

In a second step, each coordinate of the given track is compared to the static route data (Figure 21(c)) obtained from the time tables and paths. It should be noted that the public transportation schedules only contain exact arrival times for each stop, but don't provide any time data for the pathes between them. Since this data is too sparse to perform reliable service line detection on it, all paths are rasterized evenly (Figure 21(d)) and the corresponding time for each intermediate part is interpolated (Figure 21(e)). This ensures that no route section falls through the cracks when executing a simple distance query in time and space. Figure 21(f) shows the combined view of all data used for the HSL service line detection.

The result of the algorithm is the one service line which has the highest probability due to the Live API and/or the smallest distance to the given set of coordinates. During a testing phase, the optimal values of all involved parameters like distance thresholds, time and space raster sizes, and the individual scoring weights have to be learned. As one outcome of the Helsinki field trial, we expect to find a good heuristic how to setup the various ingredients of the algorithm.

3.2.1 Implementation details

As a starting point, HSL provided us with all static service line data like arrival times, stop positions, route meta data, and path files in the well known General Transit Feed Specification (GTFS)¹⁷ format. The original data set had a size of 400 MB and was imported into a PostgreSQL database. All paths and stop positions were stored as native geospatial values using PostGIS, a spatial database extender for PostgreSQL. After the import, we found 587 distinct routes, 7534 stops, and a total number of 206,153 trips in the database. In general, a service line like the bus 934 from Myyrmäki to Luhtaanmäki, has two directions, travels along its route path and stops at its stops. If a service line runs every ten minutes, each trip has its own trip id and is uniquely identifiable.

While these data are too sparse for distance queries, we interpolate the possible positions of a vehicle along its path to ensure that the distance between two consecutive trip positions is always smaller than ten meter. On the one hand, this simplifies a query like "find all trips in a distance of 20 meters around the user which arrive in plus minus one minute from now" a lot. On the other hand, the interpolation and denormalization inflates the data. Resulting in 324,440,457 trip positions annotated with arrival time and trip id, we enlarged the original data to a size of 38 GB.

This size leads the PostGIS index to its limit and prevents response times below 30 seconds for a service line detection query. Therefore, all data was partitionated horizontally in time dimension. The reason is, that all time tables repeat itself every 7 days and an usual query covers only a very

¹⁷<https://developers.google.com/transit/gtfs/>

small period of time at a specific weekday. Against this background, we divided the data into 24 parts for every day, which results in 7 times 24 subtables in the database. Finally, this setting archives average response times of less than one second for a whole service line detection query, containing of up to 200 single distance queries.

3.3 Traffic Jam Detection

In this section we are presenting a traffic jam detection module, that is able to detect irregularities in public transportation systems that provide real-time data about vehicles. It has been integrated into the Live+Gov toolkit as part of the Server Side Mining Service (C9), and was successfully used in the first Mobility Field Trial (cf. D5.1, D5.3) for the Helsinki metropolitan area.

3.3.1 Related work

Traffic fluency monitoring and jam detection is part of traffic control in all major cities and the real time information about road network fluency is of interest to all network users and traffic related authorities. Over the last years, several different methods for detecting jams have been created. In many cities traffic is monitored via road side cameras detecting fluency in the most important channels of the network and sends the image to authorities and web services. Another increasingly popular method for detecting jams is by traffic message channels (TMC) that collect information about road network status from roadside sensors and vehicles and transmits the information via radio signals to the end user. TMC information is most commonly used in navigators, both on inbuilt vehicle navigating systems and separate navigator devices.¹⁸ Over the past few years also different smartphone applications, such as Waze or Inrix have emerged on the market. In these applications information gathering is commonly done either with the vehicles connected to the system or by the users who report their journey conditions around them. Different applications provide different type of information, usually including one or several of the following information types: jams, weather, road work and accidents. What is common to all known methods for detecting jams, no service is provided specifically to the public transport users as most services are targeted for private vehicle drivers and the authorities. Based on our surveys no service is being used specifically to public transport nor uses the delay information of public transport vehicles, which makes the traffic jam detection module an unique tool.

3.3.2 Component Description

The jam detection module frequently polls vehicle location compared to the scheduled location from the Helsinki public transport service. The module analyses the vehicle situation and detects where there are traffic jams in the certain transportation authority's transportation area. The interface returns always the current jam situation, so client does not need to keep the jams in memory.

The module detects the jams based the tram location, taking into notice the delay and how fast it increases. Also, the jam is only detected if there are several vehicles filling the same criteria within the same area.

The module provides information about the nearest stop to the detected jam and also the previous stop from where the affected vehicles have come from. Based on this information current traffic situation is visualized on the application. All of Helsinki city tram network can be drawn on the map, but only lines that are currently affected by trams are shown on the map at each time. If no jams are detected in between two stops, this interval is shown with green colored line. If a

¹⁸<http://www.tisa.org/technologies/tmc/>

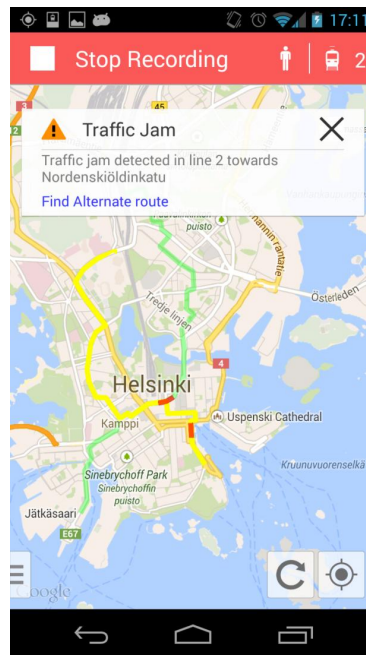


Figure 22: Visualization of the current traffic situation in Helsinki based on Jam detection

Address	213.157.92.101:80
Path	/jamdetector/JamService.svc/jams/authority
Protocol	HTTP GET
Data format	JSON

Table 4: Jam Detection Module Details

jam is detected between stops, the interval is shown with red colored line. Image 22 shows a screenshot of the application on how the jams are presented in Helsinki city center during the trial.

Architecture Description

The JamDetector module simplified architecture is visualized in the Figure 23. The public transportation systems generally differ in each city or area in multiple ways in the available data and data formats. The different public transportation systems are isolated in the module and configurable using the Spring Framework.

Live+Gov SaaS solution compatibility

The JamDetector follows Live+Gov SaaS architecture. This means that all connecting clients are authorized. Also the JamDetector provides the heart beat (health check) information to the SaaS service and sends the application log to the service.

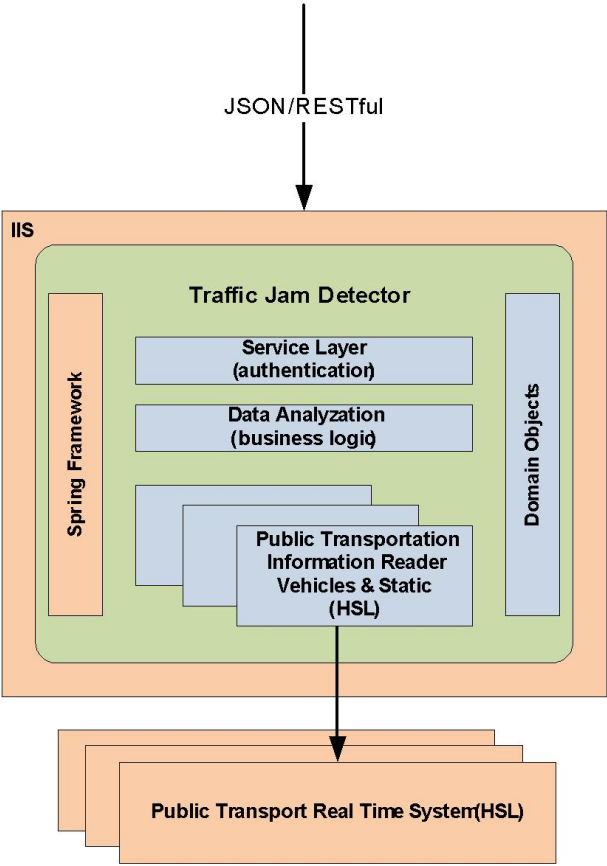


Figure 23: Traffic Jam Module Architecture

API description

The API lists the jams detected in the certain public transportation area in JSON format. The jam object contains information about the vehicles participating the detected jam and vehicles nearest stop/station. This information can be used to visualize the jam location for the users.

3.3.3 Evaluation

The module development begun with evaluating different methods for creating the detection algorithms and creating a way to detect jams reliably. The tram location data from Helsinki provides reliable information about the delay of a vehicle in real time and was considered the most efficient and reliable source of data for the jam detection.

A comprehensive and reliable detection algorithm was created through experimenting in the early stages of the module development. The parameters to be used were searched the most suitable values through thorough consideration in order to generate reasonable and accurate alerts that would correspond with the logical situation in traffic and would fit in the definition of a jam. After some weeks of monitoring the detections, the defined values were also re-evaluated and re-defined before the first pilot based on the previous observations.

Based on the first trial results, the amount of detected jams is seen relatively high, on average 3 jams were received on each time the API was called. This would indicate the parameters needing to be set tighter in the future if only greater jams are wanted to be detected. When evaluating the accuracy it was quickly discovered that validating the detection is rather challenging when not having the possibility to verify the situation on site. User questionnaires showed that no users were travelling in the areas where alerts were given at the time of the alert and therefore they could not reliably state their opinion about the accuracy of alerts. Also, no known major jams took place during the trial in the pilot area and therefore manual validation could not be done during the trial. Manual comparison to other known services providing jam alerts was done. In all cases the detected jams were of a short duration, mainly less than a minute and therefore no reliable statement about accuracy could be done at this stage. In the comparison we used two popular services: Waze¹⁹, which is a mobile application providing traffic data in over 30 countries and V-traffic²⁰, a national service in Finland for providing traffic data in major cities. Based on the comparison to other services the results seem promising as alerts were given in the same areas at the same time, only with minor differences.

When comparing the module to existing services it is obvious that there are existing systems that provide information with more geographical coverage and collect the data from greater number of sources. Thus, these services also provide more alerts as they also cover roads that are not included in the current module coverage, the Helsinki tram network. What needs to be addressed when comparing the module to existing services is that the other services do not focus on public transport and using only these services would not help to meet the requirements of the mobility use case in full. However, the possibility to support the jam detection module with existing services needs to be further studied.

We believe that when expanding the data to also cover other public transport vehicles, the coverage in the module will increase and the quality and the usefulness of the module improves.

¹⁹<http://www.inrixtraffic.com/>

²⁰<http://www.v-traffic.fi/>

For example in the Helsinki region, all public transport busses, trams, trains and the ferry will be covered in the new public transport information system to be implemented in 2016. Then the module will provide a much better service with greater geographical coverage and expanded fleet. Unfortunately at the time of development, location data is not yet available for other vehicles than trams but the possibility to include these vehicles in the jam detection has been taken into notice from the start.

As no reliable evaluation is possible at this stage we will continue validating the module, analyzing the results and developing the algorithms. This process will continue until the final trial where improvements and reliable accuracy of the module are hoped to have been achieved.

3.4 Distributed Geo-Matching

Nowadays, the means of public transportation become equipped with physical sensors like GPS sensors or light barriers. Those sensors are used to keep track of the current vehicle position or to estimate the number of passengers. But those sensors do not detect, for instance, emergencies, accidents or the passengers' opinions on bus lines. To overcome these limitations, another kind of sensors is required.

In the last decade, many social networks or micro-blogging services like TwitterTM arose. They are frequently used by humans to publish their opinions, observations, etc. which then can be requested by a web API. Therefore, the users of those services can be seen as social sensors. The increasing number of mobile phones leads to a situation in which more and more passengers become social sensors by publishing statements about crowded buses or rioting passengers.

The evaluation of the social sensor data requires that the bus or train can be identified to which a message is related. This task is simplified by the fact that most mobile phones have integrated GPS sensors. They measure the current geographical position, which is added to the published message together with a timestamp. This position information can be used to identify the nearest train or bus. A more detailed description of this matching problem will be given in the next section.

The amount of data produced by the different types of sensors can exceed the processing capabilities of a single computer. Therefore, a distributed approach is required, which will be explained in section 3.4.2.

3.4.1 Problem of Matching Physical and Social Sensor Data

In order to find the bus or train a message is related to, the data measured by the physical sensors have to be processed first. In this context, the relevant data consists of the longitude and latitude ($\mathbb{R} \times \mathbb{R}$) measured by the GPS sensors of the vehicles. Before transmitting these data, they have to be extended by the unique id (V_{id}) of the current vehicle and a timestamp (T) to keep track of the chronological order. The set of all possible position is defined as Pos in equation (3.4).

$$Pos := V_{id} \times T \times \mathbb{R} \times \mathbb{R} \quad (3.4)$$

The relevant data of the social sensors are the messages published by the passengers. Each message contains a timestamp, i.e., the publishing time and the GPS position of the mobile phone from which the message was sent. The set of all possible messages Mes is defined in equation (3.5) where C is the set of all possible message contents.

$$Mes := T \times \mathbb{R} \times \mathbb{R} \times C \quad (3.5)$$

The data of the different sensors are transmitted as a not necessarily finite stream of data as defined in the equations (3.6) and (3.7).

$$posStream : \mathbb{N} \rightarrow Pos \quad (3.6)$$

$$mesStream : \mathbb{N} \rightarrow Mes \quad (3.7)$$

The range of $posStream$, i.e., the set of vehicle positions received by the data stream $posStream$ are the trajectories of the different vehicles. A trajectory is defined in equation (3.8) as a discrete function mapping some point in time to a geographical position.

$$traj : T \mapsto \mathbb{R} \times \mathbb{R} \quad (3.8)$$

But this definition is not adequate because the timestamps of messages may vary from the timestamps contained in the position data of the vehicles. Therefore, a continuous trajectory function $\widehat{traj}$ is required. This is reached by, for instance, a linear interpolation of $traj$.

Equation (3.9) shows the function $allTraj$ which maps the unique identifier of a vehicle on its corresponding continuous trajectory.

$$allTraj : V_{id} \mapsto (T \mapsto \mathbb{R} \times \mathbb{R}) \quad (3.9)$$

In order to relate a message m to a bus or train, the vehicle must be determined which has the smallest euclidean distance $dist$ to the position of the sender at the point in time when m was sent. Furthermore, a vehicle has a maximum length $maxDist$ and all vehicles with a greater distance from m can be ignored. Thus, each message can be seen as a range nearest-neighbour query rnn as defined in equation (3.10).

$$\begin{aligned} rnn : Mes \rightarrow Pos \cup \{\text{null}\} \text{ with} \quad (3.10) \\ rnn((t, lon_m, lat_m, cont)) &:= (v_{Id}, t, lon_v, lat_v) \\ &\text{if } allTraj(v_{Id})(t) = (lon_v, lat_v) \wedge dist((lon_m, lat_m), (lon_v, lat_v)) \leq maxDist \wedge \\ &\quad \nexists v'_{Id} \in V_{id} : dist((lon_m, lat_m), allTraj(v'_{Id})(t)) < dist((lon_m, lat_m), allTraj(v_{Id})(t)) \\ rnn((t, lon_m, lat_m, cont)) &:= \text{null} \\ &\text{otherwise} \end{aligned}$$

Putting it all together, the problem of finding the nearest vehicle for each message is defined as followed:

Problem Finding range nearest vehicle for each message

Input: $posStream, mesStream, maxDist$

Output: $rnnStream : \mathbb{N} \rightarrow Mes \times Pos \bullet rnnStream(i) := (mesStream(i), rnn(mesStream(i)))$

An implementation which solves this problem has to deal with the fact, that the data received by different sensors are not equally delayed. For instance, the data measured from physical sensors may be faster transmitted then messages received from social networks.

Another difficulty is, that the amount of data received by the different input streams may exceed the processing capabilities of a single computer. One solution for this problem would be to discard incoming data if the load of the single machine would be too high. This could lead to the loss of valuable data. In order to avoid this, a distributed approach is required which scales horizontally.

Some approaches like PLACE* [13] distribute the incoming data according to a static mapping of computers on geographical regions. The disadvantage of this static mapping can be seen, for instance, if there exists a sport stadium in some region. If a sport event takes place, then there are many vehicles and passengers producing a huge amount of data. Therefore, the corresponding computer can only deal with a relatively small region. But if no event takes place, this computer has almost nothing to do. In order to improve this, the mapping has to be dynamic, i.e., a dynamic load balancing is required.

3.4.2 Distributed Geo-Matching Approach

The schema of an intuitive stream-processing system which solves the problem described above can be seen in Figure 25(a). It consists of a single component (C) which receives the incoming

data streams of vehicle positions and messages. It caches the current positions and processes the range nearest-neighbour queries initiated by the received messages. The answers are transmitted via the outgoing data stream *rnnStream*.

In order to answer range nearest-neighbour queries, *C* has to cache the current position of all vehicles. The problem of the delayed sensor data requires the caching of more vehicle positions than just the latest. Thus, a sliding window approach is implemented with a window size exceeding all regular delay times. This is realized by deleting all vehicle positions outside the window whenever a new position with a timestamp later than the timestamp of the latest cached position is received.

The data structure used as cache should answer range nearest-neighbour queries efficiently. In spatio-temporal databases R*-trees [5] are used for this purpose. Similar to a B-tree, the inserted elements are stored only in the leafs. Each node except the root has a minimal and a maximal number of children or elements. If an insertion is performed on a node already containing the maximal number of elements, then it is split and the resulting inner node is inserted in the parent node. In contrast to B-trees, each node of an R*-tree represents the minimal bounding rectangle (MBR) of all elements contained in its subtree as illustrated in Figure 24. The presented tree consists of a root node (dark gray), seven leafs (light gray) and several vehicle positions (black squares).

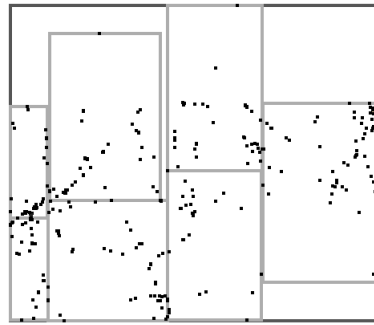


Figure 24: An exemplary R*-tree

During the processing of a range nearest-neighbour query, only the children of a node are visited whose MBR intersects with the queried range. Thus, the performance of the query processing depends on the number of subtrees to be traversed. In order to reduce this number, the MBRs of the child nodes should be disjoint. As shown in [5] the algorithms used for insertion and deletion are designed to ensure a minimal overlap of MBRs. Another advantage is, that the MBRs are dynamically adjusted according to the currently contained elements.

The problem of the system shown in Figure 25(a) is that its processing capabilities are quite limited. Therefore, a distributed approach is required which spreads the load on several computers. Figure 25(b) shows such a system. It consists of several caches C_0 to C_n each of them owning an own R*-tree for caching current vehicle positions and responding to queries.

Furthermore, this system consists of balancer component *B*. It receives the incoming streams of vehicle positions and messages and decides to which of the different caches each single datum is sent. Similar to the root of an R*-tree, it keeps track of the MBRs of the different caches. A new vehicle position is basically sent to the cache with the nearest MBR. If *B* receives a new message, it is only forwarded to the caches whose MBR intersects with the queried range. In the best case this is only one cache. But it could be several caches as well. Thus a special combiner

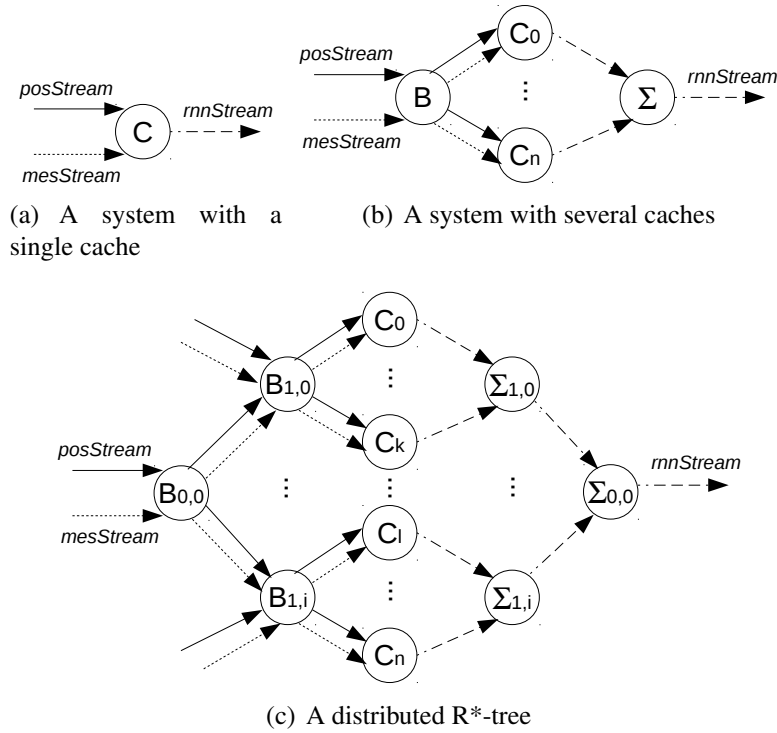


Figure 25: Distributing a geo-matching system

component (Σ) is required which receives the nearest vehicle of each cache involved in the query processing and returns the overall nearest neighbour.

A further responsibility of B is the load balancing, i.e., the number of cached vehicle positions and queries in progress should be similar for all caches. Therefore, B counts the number of vehicle positions and messages forwarded to the different caches. The load balancing is done by distributing the received vehicle positions according to special target ranges for the different caches. In the case of a balanced system, these regions are equal to the MBRs. But if the load of a cache C_i exceeds the load of a cache C_j with a neighboured MBR, then the target range of C_i is reduced and the range of C_j increased. During the progress of the sliding window the MBRs of the caches will approach their target ranges.

In regular intervals, B sends special requests to all caches to send the information about their current loads and MBRs back (not shown in Figure 25(b)). If B receives the responses of all caches, it updates its calculated values. These update requests are performed asynchronously, i.e., the distribution of the incoming data is not blocked. This leads to a situation where the information received by the caches are not up to date any more. Therefore, B caches all forwarding decisions from the time the updated request is sent until all caches have responded. Then these decisions are applied on the received load information and MBRs.

The disadvantage of this system is, that the single balancer component B and combiner component Σ can become bottlenecks. A future solution could be found when thinking of B as the root of a distributed R*-tree whose children are the trees stored in the different caches. If it is close to its processing limits, it creates child balancer components which are responsible for subregions of the overall MBR. Each of the children get an input stream for all the data of its MBR. Data for regions which are not covered by the children are still received by the root. Whenever a balancer component is split, the combiner components are split analogously. Such

a system is illustrated in Figure 25(c).

3.5 Geographic Topic Analysis

The BuitenBeter dataset consists of 13,811 geo-tagged issues reported to the office of public order in the Netherlands between July and September 2012. Each issue is assigned to one of 16 categories such as "poor road conditions", "Dirt on the street" or "Graffiti".

Geographical co-occurrences of arbitrary categorical observations - such as public order issues of the BuitenBeter dataset - can be used to discover latent underlying factors that explain observations which co-occur more frequently than expected by chance. One method for detecting these latent factors is probabilistic topic modelling, where grouped observations are modelled as being sampled from mixtures of latent topics, which are probability distributions over the set of distinct observations.

For the BuitenBeter dataset, we e.g. might expect that areas where graffiti are reported will see many reports of the category "Dirt on the street" since both categories are related to urban environments. We model those environments as hidden "topics".

Existing approaches for geographical topic modelling adopt topic models such as latent Dirichlet allocation [6] and extend the models by assigning distributions over locations to topics, or by introducing latent geographical regions. In models which extend topics for spatial distributions (such as two-dimensional normal distributions) [12], topics with a complex (i.e. non-Gaussian) spatial distribution cannot be detected. In models with latent, Gaussian distributed regions [14], documents within a complex shaped topic area do not influence the topic distribution of distant documents within the same area. Therefore, topics with a complex spatial distribution such as topics distributed along coastlines, rivers or country borders are harder to detect by such methods. More elaborate models introduce artificial assumptions about the structure of geographical distributions, e.g. by introducing hierarchical structures [1] in advance.

We introduce a novel geographical topic model which captures dependencies between geographical regions to support the detection of topics with complex, e.g. non-Gaussian distributed spatial structures [7].

Our model differs from existing approaches in several aspects:

We model locations and words separately, as the separation of spatial clusters and document semantics allows us to define meaningful neighbour relations between spatially adjacent clusters. Our topic model takes a set of geographical clusters as input.

We also expect that geographical clusters adjacent in space exhibit similar topic distributions: Most geographical topics cannot be approximated by a simple spatial probability distribution such as a Gaussian distribution and for these complex topic areas, coherent sets of multiple spatial distributions are a reasonable approximation. Therefore we detect and smooth the topic distribution of these adjacent regions to increase the probability of detecting coherent topic areas.

The detection of geographical clusters is straight-forward: We use a mixture of a fixed number of Fisher distributions – distributions on a three dimensional unit sphere similar to isotropic Gaussian distributions on a plane. The parameters are fit using the approximation given by Banerjee et al. [4] in an expectation-maximisation algorithm. In our model, each cluster is associated with a distribution over the set of topics, sampled from a Dirichlet process (DP) with a common base measure which is itself drawn from a DP with Dirichlet distributed multinomial distributions over the set of categories as base measure. For the smoothing of topic distributions of adjacent clusters, we first define the adjacency relation using the Delaunay triangulation [3]

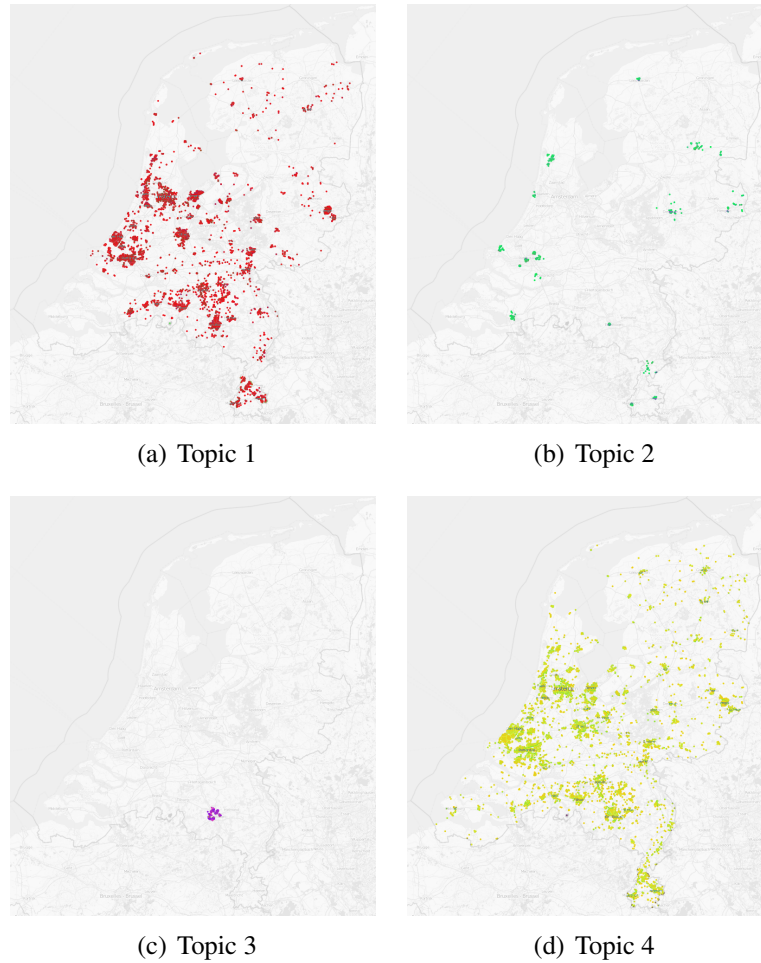


Figure 26: Document positions of reports with above-average probabilities for Topic 1-4 on the map of the Netherlands.

of the cluster centroids. Each public order issue report draws an own topic distribution from a DP with a weighted mixture of the topic distribution of its own geographical cluster and its neighbour clusters as base measure. Finally, each observed report category is drawn from its document-specific topic distribution. For a detailed description of the model and its parameters, we refer to our paper [7].

To train our model on the BuitenBeter dataset, we use 300 clusters and the initial settings for the parameters $\beta = 0.5$, $\gamma = 1.0$, $\alpha_0 = 1.0$, $A = 99999$, $\delta = 1.0$. By setting A to such a high value, we give up document-specific topic distributions and sample the categories directly from a mixture of region-specific topic distributions. All other parameters are set to the values used in the paper [7], except for the hyperparameters for A , which is fixed.

For creating a comprehensible report, the detected topics can be characterised by their most-likely categories:

Topic 1: *Dirt on the street, Other, Graffiti*

Topic 2: *Damaged street light, Bad road*

Topic 3: *Obstacle by trees, Weed, Loose paving stones, Bad road, Idea/wish*

Topic 4: *Other, Weed, Loose paving stones, Bad road, Damaged street light, Idea/wish*

The position of documents with an above-average probability for each topic are shown in Figure 26

We see that, as expected, Topic 1 is only observed in the area of larger cities, whilst Topic 4 is present across the whole country. Topic 2 occurs mostly within city centres. Topic 3 is observed in the area of Eindhoven for which different categories were available in the BuitenBeter application.

4 Conclusion

In this deliverable we presented five different data mining methods for extracting context information from the reality of citizens as well as an improved sensor collection architecture.

Three of these methods were integrated into the Live+Gov toolkit and tested in the field trial. Two methods are ongoing research projects, and have already led to a publication.

We have presented a Traffic Jam Detection component is able to detect irregularities in public transportation systems, a Service Line Detection component that is capable of detecting the precise service line used by a citizen, and a Human Activity Recognition component that classifies modes of human ambulation. These components have already shown acceptable performance and will be improved further for the upcoming second field trials.

5 References

- [1] Amr Ahmed, Liangjie Hong, and Alex Smola. Hierarchical geographical modeling of user locations from social media posts. In *WWW*, 2013.
- [2] Davide Anguita, Alessandro Ghio, Luca Oneto, Xavier Parra, and JorgeL. Reyes-Ortiz. Human activity recognition on smartphones using a multiclass hardware-friendly support vector machine. In Jos Bravo, Ramn Hervs, and Marcela Rodrguez, editors, *Ambient Assisted Living and Home Care*, volume 7657 of *Lecture Notes in Computer Science*, pages 216–223. Springer Berlin Heidelberg, 2012.
- [3] Franz Aurenhammer. Voronoi diagrams – a survey of a fundamental geometric data structure. *ACM Comput. Surv.*, 23(3):345–405, 1991.
- [4] Arindam Banerjee, Inderjit S. Dhillon, Joydeep Ghosh, and Suvrit Sra. Clustering on the unit hypersphere using von Mises–Fisher distributions. *JMLR*, 6:1345–1382, 2005.
- [5] Norbert Beckmann, Hans-Peter Kriegel, Ralf Schneider, and Bernhard Seeger. The r^* -tree: An efficient and robust access method for points and rectangles. In *Proceedings of the 1990 ACM SIGMOD International Conference on Management of Data*, SIGMOD '90, pages 322–331, New York, NY, USA, 1990. ACM.
- [6] David M. Blei, Andrew Y. Ng, and Michael I. Jordan. Latent Dirichlet allocation. *JMLR*, 3:993–1022, March 2003.
- [7] Christoph C. Kling, Jerome Kunegis, Sergej Sizov, and Steffen Staab. Detecting non-gaussian geographical topics in tagged photo collections. In *WSDM*, 2014.
- [8] Sotiris Kotsiantis and Dimitris Kanellopoulos. Discretization techniques: A recent survey. *GESTS International Transactions on Computer Science and Engineering*, 32(1):47–58, 2006.
- [9] O. Lara and M. Labrador. A survey on human activity recognition using wearable sensors. *Communications Surveys Tutorials, IEEE*, PP(99):1–18, 2012.
- [10] J. Ross Quinlan. Improved use of continuous attributes in c4. 5. *Journal of Artificial Intelligence Research*, 1996.
- [11] Claude E. Shannon. A mathematical theory of communication. 27(3):379–423, July 1948.
- [12] Sergej Sizov. GeoFolk: latent spatial semantics in Web 2.0 social media. In *WSDM*, pages 281–290, 2010.
- [13] Xiaopeng Xiong, H.G. Elmongui, Xiaoyong Chai, and W.G. Aref. Place: A distributed spatio-temporal data stream management system for moving objects. In *Mobile Data Management, 2007 International Conference on*, pages 44–51, 2007.
- [14] Zhijun Yin, Liangliang Cao, Jiawei Han, Chengxiang Zhai, and Thomas S. Huang. Geographical topic discovery and comparison. In *WWW*, pages 247–256, 2011.