

Optimizing the configuration of an heterogeneous architecture of sensors for activity recognition, using the extended belief rule-based inference methodology



Macarena Espinilla^{a,*}, Javier Medina^a, Alberto Calzada^b, Jun Liu^b, Luis Martínez^a, Chris Nugent^b

^a Computer Sciences Department, University of Jaén, Campus Las Lagunillas s/n, Jaén, 23071, Spain

^b School of Computing and Mathematics, University of Ulster, Jordanstown, BT37 0QB, U.K.

ARTICLE INFO

Article history:

Received 6 February 2016

Revised 4 September 2016

Accepted 30 October 2016

Available online 5 November 2016

Keywords:

Heterogeneous architecture of sensors

Optimization

Efficient power-aware

Activity recognition

Ambient intelligence

Feature selection

Data-driven

ABSTRACT

Smart environments are heterogeneous architectures with a broad range of heterogeneous electronic devices that are with high in processing capabilities for computing, considering low power consumption. They have the ability to record information about the behavior of the people by means of their interactions with the objects within an environment. This kind of environments are providing solutions to address some of the problems associated with the growing size and ageing of the population by means of the recognition of activities that can offer monitoring activities of daily living and adapting the environment. In order to deploy low-cost smart environments and reduce the computational complexity for activity recognition, it is a key issue to know the subset of sensors which are relevant for activity recognition. By using feature selection methods to optimize the subset of initial sensors in a smart environment, this paper proposes the adaption of the extended belief rule-based inference methodology (RIMER+) to handle data binary sensors and its use as the suitable classifier for activity recognition that keeps the accuracy of results even in situations where an essential sensor fails. A case study is presented in which a smart environment dataset for activity recognition with 14 sensors is set. Two optimizations with 7 and 10 sensors are obtained with two feature selection methods in which the adaptation of RIMER+ for smart environment provides an encouraged performance against the most popular classifiers in terms of robustness.

© 2016 Elsevier B.V. All rights reserved.

1. Introduction

The number of elderly will reach 2 billion by the year 2050 and a key issue for this people is to stay as long as possible in their own homes in order to have a healthy ageing and wellbeing [44,45]. One of the most common diseases in this group is related to cognitive processes such as dementia. These illnesses are currently incurable; hence efforts are focused towards delaying their progression. In the early stages of dementia, it is useful to provide support in the form of prompting through the completion of activities of daily living (ADL) in addition to offering a series of reminders for tasks such as medication management, eating or grooming [12,19,25].

Smart environments are built heterogeneous architectures with a broad range of multiples and different electronic devices that are with high in processing capabilities for computing, considering low power consumption. These environments have the ability to record information about the behavior of the person by means of his/her interaction with the objects within an environment [8]. So, smart environments are residences with a heterogeneous architecture of sensor in which sensors are connected to a range of objects or locations and networked and used to identify people in the environment and their actions [5]. This kind of environments are adapted to perceive the user contexts in order to help people in their ADL providing a smart solutions to address some of the problems associated with the growing size of the population.

The sensor-based activity recognition [1,5,31,35,48] is at the core of smart environments. So, the process of activity recognition aims to recognize the actions and goals of one or more person within the environment based on a series of observations of actions and environmental conditions. It can, therefore, be deemed as a complex process that involves the following steps: (i) to choose and deploy the appropriate sensors to objects within the smart

* Corresponding author.

E-mail addresses: mestevez@ujaen.es (M. Espinilla), jmquero@ujaen.es (J. Medina), a.calzada@ulster.ac.uk (A. Calzada), j.liu@ulster.ac.uk (J. Liu), martin@ujaen.es (L. Martínez), cd.nugent@ulster.ac.uk (C. Nugent).

environment; (ii) to collect, store and process information and, finally, (iii) to classify activities from sensor data through the use of computational activity models.

Traditionally, approaches used for sensor-based activity recognition have been divided into two main categories: Data-Driven (DDA) and Knowledge-Driven (KDA) Approaches [5].

The former, DDA, are based on machine learning techniques in which a pre existent dataset of user behaviors is required. A training process is carried out, usually, to build an activity model which is followed by testing processes to evaluate the generalization of the model in classifying unseen activities [21,41]. The advantages of the DDA are the capabilities of handling uncertainty and temporal information. However, these approaches require large datasets for training and learning, and suffer from the data scarcity or the cold start problem [5].

Regarding KDA, an activity model is built through the incorporation of rich prior knowledge gleaned from the application domain, using knowledge engineering and knowledge management techniques [5,7]. KDA has the advantages of being semantically clear, logically elegant, and easy to get started. Nonetheless, they are weak to deal with uncertainty and temporal information as well as the activity models can be considered as static and incomplete.

In order to take advantage of the main benefits provided by DDAs and KDAs and to avoid some of their common disadvantages, some hybrid approaches have been developed such as the extended belief rule-based inference methodology (RIMER+) that was proposed in [34]. Knowledge Base of RIMER+ is based on Extended Belief Rule-Bases (E-BRB), which are able to capture (i) sample data and expert knowledge in a homogeneous way; (ii) nonlinear and causal relationships; and (iii) several types of uncertainty related to expert knowledge and data.

In the literature, we can find multiple smart environments with different heterogeneous architectures that are equipped with a large set of multiple and different sensors in order to carry out an activity recognition process. For example in the smart environment presented in [9] are involved 39 sensors and in the smart environment shown in [47] are involved 52 sensors. Although there has been significant progress in sensor-based activity recognition with promising results, which offer improvements with the problems associated with the growing size of the population, it still remains expensive to deploy a full set of sensors within a smart environment [36,50].

Therefore, the selection of an appropriate set of the sensors in the heterogeneous architecture, which are placed in objects within the environment, is an important issue in order to effectively monitor and capture the user's behavior along with the state change of the environment. This selection has a direct influence on the sensor data that will be used in the activity classification process. So, the current challenge is to know what configuration of initial sensors is the essential for activity recognition, i.e., the optimization of sensors in the heterogeneous architecture on a smart environment. This optimization process should imply the selection of a subset of the original sensors without loss of accuracy for activity recognition, even if one of the essential sensors fails because there are any technical failure or any sensor are deactivated, subsequently returning wrong values.

Feature selection methods [13] provide a way to select a subset of relevant features to generate a classifier or a model from a dataset obtained from a real process. Thus, these methods can identify which of the features are relevant or if there are interdependency relations between them. The use of feature selection methods in smart environments has been conducted in some studies [17,18,20]. However, none of them has shed light on the classifier which should be used in conjunction with feature selection methods in order to keep the accuracy, providing robustness.

In this paper, we focus on the use of feature selection methods to optimize the subset of initial sensors for activity recognition in a smart environment with binary sensor in order to reduce costs from a technology perspective, maintaining accuracy for activity recognition as well as reducing the computational complexity. Furthermore, this paper suggests the adaptation of RIMER+[34] like an approach for activity recognition, which is called R+DRAH. This new adaptation is focused on handle binary sensor data due to the fact the traditional RIMER+ works with values between 0 and 1. So, this new version offers accuracy in a smart environment for activity recognition with binary sensor optimization, being also robust in situations in which a relevant sensor fails. A case study for activity recognition, in which two optimizations of sensors are obtained using feature selection, is carried out. An evaluation of R+DRAH against the most popular DDA classifiers is performed, considering the situation when a sensor failure in order to provide robustness.

The remainder of the paper is structured as follows: [Section 2](#) presents related works in the activity recognition area. [Section 3](#) reviews RIMER+ that will be used in our proposal. [Section 4](#) discusses how RIMER+ is adapted in R+DRAH to focus specifically on works in smart environments. [Section 5](#) discusses the feature selection methods in order to optimize the set of sensors within a smart environment. [Section 6](#) presents an empirical study which analyzes two feature selection methods for the purpose of sensors optimization for activity recognition with R+DRAH, considering the situation when a sensor failure in order to evaluate and analyze the robustness.

2. Background

In this section, the basic notions to understand our proposal are reviewed.

2.1. Sensor-based activity recognition

Advances in sensors technology developments have encouraged research on the problem of activity recognition based on processing data obtained by sensors [5]. In order to interpret the sensor data to infer activities, it is necessary to build activity models. In this section, we review the approaches to build activity models.

Activity models can be built by means of DDA or KDA. Both approaches are reviewed in the following subsections:

2.1.1. Data-driven approaches

The DDA learn activity models from preexistent large-scale datasets of users' behaviors using data mining and machine learning techniques. These approaches imply the generation of probabilistic or statistical activity models by means of training and learning processes. So, the activity inference is based on a probabilistic or statistical classification.

Some of the most popular classifiers based on DDA approaches are briefly reviewed below:

- Naive Bayes classifier (NB) [16]. The basic idea in NB classifier is to use the joint probabilities of sensors and activities to estimate the category probabilities given a new activity. This method is based on the assumption of sensor independence, i.e. the conditional probability of a sensor given an activity is assumed to be independent of the conditional probabilities of other sensors given that activity.
- Nearest Neighbor (NN) [11]. This classifier is based on the concept of similarity [7] and the fact that patterns which are similar, usually, have the same class label. An unlabeled sample is classified with the activity label corresponding to the most frequent label among the k nearest training samples.

- Decision Table (DT) [28]. This classifier is based on a table of rules and classes. Given an unlabeled sample, this classifier searches for the exact match in the table and returns the majority class label among all matching samples, or informs no matching is found.
- Support Vector Machines (SVMs) [10]. This method focused on a nonlinear mapping to transform the original training data into a higher dimension. Within this new dimension, it searches for the linear optimal separating hyper plane. A hyperplane is a decision boundary which separates the tuples of one activity from another.

2.1.2. Knowledge-driven approaches

The KDA build activity models by using the rich domain knowledge and heuristics for activity modeling and pattern recognition. This approach, usually, involves knowledge acquisition, formal modeling, and representation. Activity models generated in this method are normally used for activity recognition or prediction through formal logical reasoning, e.g., deduction, induction, or abduction.

Depending on the way to capture, represent and use the knowledge, the knowledge driven approaches to activity modeling and recognition can be classified into two main categories:

- Logical modeling based approach [6]. Activity modeling is equivalent to knowledge modeling and representation. Knowledge representation formalisms or languages are used to represent these knowledge models and concrete knowledge instances, enabling inference and reasoning.
- Evidential theory based approach [36]. In this approach, sensor readings are used as evidence of higher level states within an activity model. These states are fused to determine more complex and higher level states until the level of belief in the activities of interest is determined.

2.2. Extended belief rule-based inference Methodology. RIMER+

In this section, the extended belief rule-based inference methodology (RIMER+) is reviewed, which can be considered as a combination of DDA and KDA, in order to take the advantages of both approaches. We will adapt this methodology for activity recognition in a smart environment.

The RIMER+ classifier, [34], utilizes the recently developed extended belief rule-based inference methodology, which combines DDA and KDA. Thus, RIMER+'s Knowledge-Base (KB) is based on Extended Belief Rule-Bases (E-BRB), which are able to capture in a homogeneous way the sample data and expert knowledge, uses casual and nonlinear relationships and, moreover, it can deal with some types of uncertainty related to expert knowledge and data.

RIMER+ uses the Evidential Reasoning (ER) algorithm [52] to infer the information included in its E-BRBs, in order to produce a prediction result based on some input of the system. Furthermore, one of the main features of RIMER+ is its E-BRB, which extends the KB used in the belief Rule-Base Inference Methodology (RIMER) [51]. Thus, the belief degree distributions are used. So, they are embedded in the consequent terms of its rules and also in each antecedent term. Take for example the following EBR:

IF Temperature is (Hot, 0), (Warm, 0.1), (Cold, 0.7)

THEN Heating is (ON, 0.2), (OFF, 0.8)

Therefore, EBR can capture the fuzziness with linguistic terms, uncertainty with beliefs, the incompleteness with ignorance or partially known belief and non linear relationships by means of IF-THEN rules. Thus, this type of extended belief rule is generic and, moreover, it provides a flexible way to incorporate hybrid input information and an efficient rule generation scheme to build E-BRBs directly from sample data [34].

As has been mentioned, E-BRBs provide the flexibility to incorporate context information in the KB. This information could

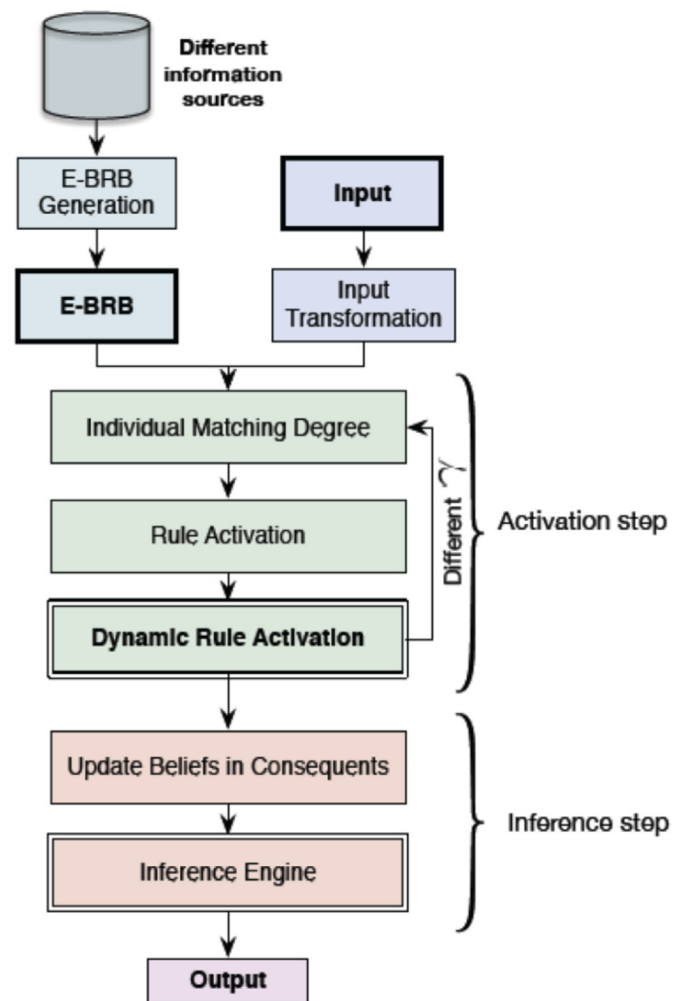


Fig. 1. R+DRA process (taken from [2]).

be vague, uncertain and/or incomplete and heterogeneous (qualitative or quantitative). Furthermore, it provides the means to set the importance of different rules and antecedents, using rule weights (noted as θ_k for the k th rule) and antecedent relative weights (noted as δ_{ik} for the i th antecedent of the k th rule), respectively. Once the E-BRB is generated from sample data and the expert knowledge about the environment is added, it can be used to recognize future activities, setting the values of sensors as inputs for the E-BRB. This recognition method of the RIMER+ approach is based on two main processes:

1. Rule Activation: evaluates which rules need to be activated, computing their activation weights (w_k) using the similarity of their antecedents against the given inputs.
2. Rule Inference: the ER approach [52] is applied to combine the activated rules and generate the final output.

Given that an E-BRB may be generated from sample data, the quality of data might be an important issue to be concerned when generating a reliable E-BRB. In this regard, a new Dynamic Rule Activation (DRA) algorithm was proposed in [3,4] as a method to select the most relevant information to be aggregated. This is undertaken by considering that data incompleteness and inconsistency may be viewed as paired situations, given that the former appears due to the lack of information while the latter can be considered as an excess of heterogeneous information [4]. This upgraded method is denoted as R+DRA and is detailed in Fig. 1.

So, this approach enhances the performance of RIMER+, especially in multi-class datasets as is the case of the process for activity recognition.

Finally, it is noteworthy that the processes listed in Fig. 1 can be modified depending on each particular scenario. So, in the next Subsection, details of how the individual matching degree and rule activation are provided to work with binary sensor data in smart environments.

3. R+DRAH for smart environments with binary sensors

In this section, we describe the adaption of the RIMER+, which is called R+DRAH, for activity recognition in smart environments with binary sensor data

The binary representation indicates which of the sensors have been active within a time interval [40]. In order to handle this kind of data, RIMER+ provides the flexibility to modify the similarity and aggregation functions that calculate the individual matching degree and rule activation weights, as was illustrated in Fig. 1.

Therefore, in R+DRAH to handle binary data, the similarity and aggregation functions are replaced with the following similarity measure, which is very popular, so-called hamming distance [24]:

$$H(\alpha, A_k) = \sum_{i=0}^{T_k} \delta_{ik} \times d_H(\alpha_i, A_{ik})$$

where

$$d_H(\alpha_i, A_{ik}) = \begin{cases} 0, & \text{if } |\alpha_i, A_{ik}| = 0 \\ 1, & \text{otherside} \end{cases} \quad (1)$$

where α is the input vector (α_i is the input for the i th antecedent attribute) and A_k is the antecedent vector for the k th rule (A_{ik} is the i th antecedent of the k th rule).

Note that in the case study presented in this paper, each α_i and A_{ik} may only take the values 0 or 1, because the activity recognition environment is only based on binary sensors. If the similarity function H (Eq. (1)) returns zero this means that the input vector perfectly matches the antecedents of the k th rule, i.e. the current sensor values totally match the description of one activity. The higher value H returns, the more dissimilar the input vector is to the k th rule.

In order to accommodate noise in the input data of the dataset, the rule activation weights (w_k) are calculated as follows:

$$w_k = \begin{cases} 1 \times \theta_k & \text{if } H(\alpha, A_k) = 0 \\ 0.2 \times \theta_k & \text{if } H(\alpha, A_k) \leq 1 \\ 0.1 \times \theta_k & \text{if } H(\alpha, A_k) \leq 2 \\ 0 & \text{if } H(\alpha, A_k) > 2 \end{cases} \quad (2)$$

Eq. (2) shows how we can be handled data with some noise in the input vector (e.g., sensors with technical failures or deactivated, which produce wrong values) using a hybrid function of hamming distance. Although rules containing a certain amount of noise (in 1 or 2 sensors) are not completely discarded, but their activation weight is substantially lower than if no noise is found ($H(\alpha, A_k) = 0$). For the binary sensor data, the proposed H function suits better than the Euclidean distance used in [34].

4. Selection features for sensors optimization

In this paper, the proposal is to obtain a set of relevant sensors and a successful classifier for activity recognition, in order to deploy low-cost smart environments. In smart environments, it is not easy to know which the most relevant sensors are. This section discusses the use of feature selection methods to face this problem, providing a way to obtain a sensor optimization for a smart environment.

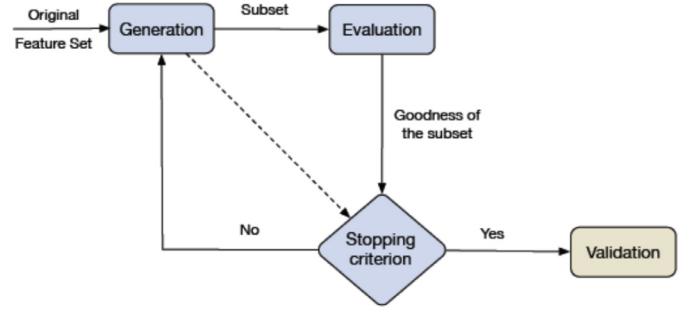


Fig. 2. Feature selection process with validation (taken from [14]).

It is really difficult to know, a priori, the relevant sensors which should be considered in a classification problem for activity recognition. Sometimes, the datasets generated by smart environments gather information from multiple sensors in order to have a complete representation of the environment domain. Nevertheless, this fact produces large datasets with redundant and irrelevant information, which could deteriorate the performance of the classification for activity recognition.

In addition, the increase of the dataset's size produces a larger computational complexity and potentially hinders the learning processes and generalization capabilities of the classifier.

The problem of identifying the most relevant sensors for activity recognition in DDA is closely related to the selection of features in a classification problem [26,43].

Feature selection methods are carried out in order to reduce the dataset, keeping as much information as possible about the domain, without a negative impact on the classification accuracy [15]. So, irrelevant and redundant features are eliminated. Reducing the number of these features clearly improves the time taken to deploy a learning algorithm and assists in obtaining a better insight into the concept of the underlying classification problem [28,29].

In order to offer a formal description of a feature selection method, the definition proposed by Dash and Liu [13] is provided:

Definition 1. Feature selection attempts to select the minimal sized subset of features according to the following criteria: i) the classification accuracy does not significantly decrease; and ii) the resulting class distribution, given only the values for the selected features, is as close as possible to the original class distribution, given all features.

Therefore, feature selection methods have a three-fold motivation. First, to simplify the classification problem due to the fact that only the relevant features are used. Second, to maintain approximately the accuracy of the classifier, after removing some of the initial features. Finally, to reduce the dimension of the dataset which, proportionally, reduces the computational complexity.

Fig. 2 illustrates the process of feature selection methods which consists of the following three steps:

- 1 Generation Procedure.** The generation procedure attempts to find out optimal feature subsets which summarize the whole set, reducing the computational complexity. In the case of a dataset that contains N features, the total number of candidate subsets to be generated is 2^N . There are different approaches to addressing this process:
 - **Complete.** This approach carries out a complete search for the optimal sub set according to an evaluation function that is utilized. The most common complete approach is the Exhaustive Search, which go over all the searching space. Thus, some algorithms use different functions that are based on various backtracking techniques, such as branch and bound, best first search

or beam search, in order to reduce the searching space and optimize the search time [39,42].

- **Heuristic.** With this approach, in each iteration, all the features that remain are considered to be selected (rejected), generating the subsets in an incremental way. In this case, the searching space is quadratic in terms of the number of features. Thus, this method has a higher-speed than the complete search, to produce the results [33].
- **Random.** This generation procedure differs with the other methods, because it is set a maximum number of possible iterations and, usually, searches a fewer number of subsets than 2^N . In order to select an optimal subset, it is necessary to assign suitable values to different parameters which take part in the random process [46].

In this paper, we are focusing on the method which guarantees the best accuracy performance, flouting the searching time. Thus, a complete generation procedure is used, concretely the Exhaustive Search. Therefore, 2^N candidate subsets will be evaluated.

2 Evaluation Functions. An evaluation function is used to evaluate the subset under examination. An evaluation function measures the goodness of a subset produced by some generation procedure, and this value is compared with the previous best. If it is found to be better, it replaces the previous best subset. Generally, an evaluation function tries to measure the discriminating ability of a feature or a subset in order to distinguish the different class labels. There are several types of evaluation functions:

- **Distance Measures** that have the idea that in a two-class problem, the most preferred features are those which induce a higher difference between the conditional probabilities of two classes [27,30].
- **Information Measures** which are based on the information gain. A feature is preferred to another if the information gain from the first feature is higher than the second one. An example of this type is entropy. These measures are employed in [29,44].
- **Dependence Measures** that quantify the ability to predict the value of one variable from the value of another one. Therefore, a feature is preferred to another if the correlation between the features with a class is higher than the correlation between another feature and the same class [37,38].
- **Consistency Measures** which deal with to find out the minimally sized sub set that satisfies the tolerable inconsistency rate, which is normally set by the user. These measures are able to remove redundant and/or irrelevant features and capable of handling some noise [14,33].
- **Classifier Error Rate Measures** that depend on the classifier itself in order to perform the feature selection [26,32].

In the presented case study, the consistency measures and the dependence measures will be used as evaluation functions due to the fact that these measures have provided a very good overall performance in classification problems [14].

2 Stopping criterion. The feature selection process is an iterative process. For this reason, it is necessary to establish a criterion which indicates when the process is finished. In this final moment, the subset of features has the highest score discovered up to that point and it is selected as the satisfactory feature subset. The choice of the stopping criterion may depend on the generation procedure and the evaluation function; possible criteria include: a subset score exceeds a threshold, a program's maximum allowed run time has been surpassed, a maximum iteration number, etc.

Finally, usually, a validation process is carried out in order to check the validity of the problem generated subset of features, by

testing and comparing the obtained results with results previously computed with the original dataset.

5. Case study

In this Section, the dataset of a smart environment, which will be used in our proposal to be optimized and used for activity recognition, is shown. Then, two optimizations are obtained by means of two feature selection methods. Finally, to evaluate the suitability in terms of robustness of the R+DRAH against four popular classifiers, some tests are carried out with the two computed optimizations.

5.1. Activity recognition dataset

The case study used in this paper is based on a popular activity recognition dataset [49]. The dataset was collected in the house of a 26-year-old male who lived alone in a three-room apartment. In the environment, 14 state-change sensors were installed.

This dataset contains 245 observations which have been generated by 14 binary sensors. The number of instances of each activity class in the dataset and the frequency of each sensor in each activity class are presented in table.

Each sensor was located in one of 14 different places within a home setting: microwave, hall-toilet door, hall-bathroom door, cups' cupboard, fridge, plates' cupboard, front door, dishwasher, toilet flush, freezer, pans' cupboard, washing machine, groceries' cupboard and hall-bedroom door. Sensors were left unattended, collecting data for 28 days in the apartment. Activities were annotated by the subject himself using a Bluetooth headset.

Seven different activities were annotated, namely: going to bed, using toilet, preparing breakfast, preparing dinner, getting a drink, taking a shower and leaving the house.

5.2. Optimizing the configuration of a heterogeneous architecture of sensors with feature selection methods

In this paper, the application of two feature selection methods is proposed to obtain two sensor optimizations of a smart environment, offering a low-cost smart environment for activity recognition.

To do so, two sensor optimizations for the smart environment are obtained. The first optimization is based on a *consistency measure* and the second one is based on a *dependence measure* due to the fact that these measures have provided a very good overall performance [14]. The feature selection methods have been implemented using the Weka software [22] that is a Java software tool with machine learning algorithms for solving real-world data mining problems, which has GNU General Public License, GPLv3.

In the first optimization, a feature selection method based on consistency measure *consistencySubsetEval* [33] is applied. This function evaluates the worth of a subset of sensors by the level of consistency in the class values when the training instances are projected on the subset of sensors. So, subsets of sensors are highly correlated with the class, considering a low intercorrelation.

In the evaluation functions based on consistency measure case, the consistency of any subset can never be lower than that of the full set of sensors. Therefore, usually, this function is used in conjunction with a generation procedure of type Random or Complete, which looks for the smallest subset with consistency equal to that of the full set of sensors.

Although the complete approach takes more time than random approach, the first approach is used in our case due to the fact that the search is conducted in the full search space, providing better results than the random method. In our context, the time to search

Table 1

Distribution of activity class and sensors.

		Going to bed	Using toilet	Preparing breakfast	Preparing dinner	Getting a drink	Taking a shower	Leaving the house
ID	Sensors\Num. Act.	24	114	20	10	20	23	34
1	Microwave	0	0	5	3	0	0	0
2	Hall-Toilet-door	6	26	1	0	0	23	1
3	Hall-Bathroom-door	0	98	0	0	0	1	0
4	Cups-cupboard	0	0	2	4	16	0	0
5	Fridge	0	0	20	6	19	0	0
6	Plates-cupboard	0	0	19	10	0	0	0
7	Front-door	0	0	0	0	0	0	34
8	Dishwasher	0	0	1	1	1	0	0
9	Toilet-Flush	0	97	0	0	0	0	0
10	Freezer	0	0	6	9	1	0	0
11	Pans-Cupboard	0	0	2	6	0	0	0
12	Washing-machine	0	0	0	0	0	0	0
13	Groceries-Cupboard	0	0	16	8	0	0	0
14	Hall-Bedroom-door	0	4	0	0	0	0	0

Table 2

Set of sensors for each optimization with feature selection.

ID	Sensors	Initial sensors	Optimization 1	Optimization 2
1	Microwave	Yes	Yes	No
2	Hall-Toilet-door	Yes	No	Yes
3	Hall-Bathroom-door	Yes	Yes	Yes
4	Cups-cupboard	Yes	Yes	No
5	Fridge	Yes	Yes	Yes
6	Plates-cupboard	Yes	Yes	Yes
7	Front-door	Yes	Yes	Yes
8	Dishwasher	Yes	No	No
9	Toilet-Flush	Yes	Yes	Yes
10	Freezer	Yes	Yes	No
11	Pans-Cupboard	Yes	No	No
12	Washing-machine	Yes	No	No
13	Groceries-Cupboard	Yes	Yes	No
14	Hall-Bedroom-door	Yes	Yes	Yes
Number of Sensors		14	10	7

the optimization is not relevant due to the fact that the optimization only be calculated once.

Regarding the dataset described in the previous section, the feature selection method selects a subset of 10 sensors, which are indicated in the second column of the [Table 2](#).

In the second optimization, a feature selection method based on dependence measure, *consistencySubsetEval* [23] is applied. This function evaluates the worth of a subset of sensors by considering the individual predictive ability of each sensor with the redundancy degree between them. So, the subsets of sensors which are highly correlated with the class are preferred, taking into account a low intercorrelation. In this case, the generation procedure of type Complete is also used in the optimization for the same reason.

The feature selection method based on dependence measure generates an optimization with 7 sensors, which are indicated in the third column of [Table 2](#).

So, we have obtained two optimizations. On the one hand, a smart environment is obtained with a cost reduction of 29% with the first sensor optimization due to the fact that this optimization has 10 sensors, removing 4 sensors of the initial set. On the other hand, a smart environment is obtained with a cost reduction of 50% with the second sensor optimization because this optimization has 7 sensors, eliminating 7 sensors of the initial set.

In the following section, the two obtained optimizations, with 10 and 7 sensors, are evaluated for activity recognition, comparing with four popular classifiers.

5.3. Recognition activity with sensor optimizations. Results and discussion

Once the two optimizations of the smart environment have been generated, this paper proposes the use of the adaptation of RIMER+, R+DRAH, as the most suitable classifier in terms of robustness for activity recognition in a smart environment with a sensor optimization.

To prove the robustness of R+DRAH in situations in which a sensor has some type of failure, a series of tests for each optimization is carried out. So, R+DRAH is evaluated comparing with other four popular classifiers used as DDA approaches for activity recognition. These four classifiers are NB, NN, DT and SVMs that were reviewed in the [Section 2.1](#).

Three commonly-used types of tests were executed to evaluate the accuracy rate based on a cross-validation: 10-fold Cross-Validation (CV10), 4-fold Cross Validation (CV4) and 2-fold Cross-Validation (CV2). The main advantage of the cross-validation is that all the samples in the dataset are eventually used for training and testing. So, this validation avoids the problem of considering how the data is divided. The accuracy mean rate for the three cross-validations (C10, CV5 and CV2) is also computed to be considered.

For the optimization 1, [Table 3](#) shows the assessment of the classifiers with 10 sensors as well as when one sensor has some type of failure with the optimization 1. It is noteworthy that the sensors with IDs 2, 8, 11 and 12 are not included because the optimization 1 removes these sensors in the selection feature process.

Regarding the situation in which the set of 10 sensors is available in the optimization 1, all classifiers have a higher performance than the 91%. In this ideal situation, SVM obtains the best results

Table 3
Assessment of accuracy and robustness of optimization 1.

		Opt. 1		CV10	CV4	CV2	Mean			
		NB		96,33	95,92	94,69	95,65			
		NN (k = 10)		94,69	91,84	91,84	92,79			
		DT		95,51	94,69	93,47	94,56			
		SVM		96,73	97,14	95,10	96,32			
		R+DRAH		93,47	95,51	92,24	93,74			
F. sensor 1	CV10	CV4	CV2	Mean	F. sensor 7	CV10	CV4	CV2	Mean	
NB	96,32	96,73	95,51	96,19	NB	87,34	86,94	85,71	86,66	
NN (k = 10)	91,83	91,84	91,84	91,83	NN (k = 10)	83,67	82,86	82,86	83,13	
DT	93,46	94,69	93,88	94,01	DT	85,71	85,71	84,49	85,30	
SVM	97,56	96,33	95,10	96,33	SVM	88,16	88,16	86,12	87,48	
R+DRAH	97,14	97,14	94,70	96,33	R+DRAH	95,51	95,51	92,24	94,42	
F. sensor 3	CV10	CV4	CV2	Mean	F. sensor 9	CV10	CV4	CV2	Mean	
NB	90,61	89,80	88,57	89,66	NB	90,61	89,80	88,57	89,66	
NN (k = 10)	86,93	86,12	86,12	86,39	NN (k = 10)	86,93	86,12	86,12	86,39	
DT	88,57	88,57	87,35	88,16	DT	88,57	89,39	85,71	87,89	
SVM	91,02	90,20	88,57	89,93	SVM	91,02	91,02	89,39	90,48	
R+DRAH	95,92	95,92	92,65	94,83	R+DRAH	95,92	95,92	93,88	95,24	
F. sensor 4	CV10	CV4	CV2	Mean	F. sensor 10	CV10	CV4	CV2	Mean	
NB	93,46	94,29	94,29	94,01	NB	95,10	96,33	96,73	96,05	
NN (k = 10)	93,06	93,06	87,76	91,29	NN (k = 10)	91,84	91,84	91,84	91,84	
DT	93,87	94,69	93,47	94,01	DT	95,91	95,10	93,88	94,96	
SVM	97,14	97,55	95,10	96,60	SVM	95,51	95,51	95,10	95,37	
R+DRAH	97,55	97,55	94,69	96,60	R+DRAH	96,73	96,73	94,69	96,05	
F. sensor 5	CV10	CV4	CV2	Mean	F. sensor 13	CV10	CV4	CV2	Mean	
NB	92,24	92,65	92,65	92,51	NB	96,32	95,51	94,29	95,37	
NN (k = 10)	92,24	91,84	82,45	88,84	NN (k = 10)	92,24	92,65	92,24	92,38	
DT	92,65	92,65	93,47	92,92	DT	95,91	95,10	93,47	94,83	
SVM	94,28	93,88	93,06	93,74	SVM	94,28	93,88	93,88	94,01	
R+DRAH	96,73	96,73	92,65	95,37	R+DRAH	96,73	95,92	93,47	95,37	
F. sensor 6	CV10	CV4	CV2	Mean	F. sensor 14	CV10	CV4	CV2	Mean	
NB	95,91	95,10	94,29	95,10	NB	87,75	86,94	85,71	86,80	
NN (k = 10)	91,83	91,02	88,98	90,61	NN (k = 10)	84,08	83,27	83,27	83,54	
DT	94,69	95,10	94,10	94,96	DT	84,08	86,53	80,41	83,67	
SVM	95,10	95,92	94,29	95,10	SVM	88,16	88,16	86,53	87,62	
R+DRAH	95,92	95,92	93,47	95,10	R+DRAH	95,92	95,92	93,88	95,24	

with a rating mean of 96,32%, followed by NB, DT, R+DRAH and NN.

In situations where a sensor can fail with the optimization 1, R+DRAH obtains the best results when there is a lack of information. Furthermore, the results obtained by R+DRAH are overcome by several points to the results obtained by other classifiers. This fact is shown in cases when sensors with IDs: 3, 7, 9 and 14 fail. In other cases, when sensors with IDs: 1, 4, 5, 6, 10 and 13 are not operative, the good results of R+DRAH are equivalent to the classifier SVM or NB classifiers.

The best results with 10 sensors are obtained with the SVM classifier, but when a sensor is not operative, R+DRAH provides similar results or the best results.

In the second series of tests, Table 4 shows the assessment of classifiers for the optimization 2 with 7 sensors and when one sensor has some type of failure. It is noteworthy that sensors with IDs: 1, 4, 8, 10, 11, 12 and 13 are not included because this optimization removes these sensors in the selection feature process.

When the set of 7 sensors is available in the optimization 1, all classifiers have a higher performance than the 90%. However, compared with the first optimization, the best classifier is R+DRAH with an accuracy mean rate of 96,33% followed by NB with a 95,78% and DT, SVM and NN.

In situations where a sensor can fail with the optimization 2, R+DRAH also obtains the best results. Moreover, the results ob-

tained by R+DRAH are overcome by several points to the results obtained by other classifiers in all cases with loss of information.

We can conclude that the conjunction of features selection methods with the R+DRAH as classifier for activity recognition provides an encouraged performance to obtain in a low-cost smart environment. Furthermore, it is remarkable that with optimizations with fewer sensors, the classifier R+DRAH obtains more successful results.

Moreover, on the first hand, the datasets generated by smart environments gather information from multiple sensors, which can have redundant and irrelevant information. The performance of the classification for activity recognition overcomes this problem by means of features selection methods, generating low-cost environment with suitable sensors. On the other hand, R+DRAH provides a suitable performance due to its ability to handle data with some noise, for example, sensors with technical failures, false positives or temporary disconnections. This is mainly based on the dynamic rule activation method to select the most relevant information to be aggregated in order to active the rules to infer the activities.

Finally, regarding to other classifiers, such as the SVM (the second best classifier), which are a black box classifier; the R+DRAH is a white box generating interpretable fuzzy rules, which can be analyzed, adjusted and tuned by human experts.

Table 4

Assessment of accuracy and robustness of optimization 2.

Opt. 2	CV10	CV4	CV2	Mean	F. sensor 6	CV10	CV4	CV2	Mean
NB	95,51	95,92	95,92	95,78	NB	87,35	87,76	87,76	87,62
NN (k = 10)	93,88	93,88	90,20	92,65	NN (k = 10)	85,71	85,71	85,71	85,71
DT	95,92	95,10	93,88	94,97	DT	84,90	85,31	86,12	85,44
SVM	95,10	94,29	93,47	94,29	SVM	87,35	86,12	85,71	86,39
R+DRAH	98,37	95,92	94,69	96,33	R+DRAH	98,78	96,33	95,10	96,74
F. sensor 2	CV10	CV4	CV2	Mean	F. sensor 7	CV10	CV4	CV2	Mean
NB	95,51	95,92	95,92	95,78	NB	95,10	95,51	95,51	95,37
NN (k = 10)	93,88	93,88	89,80	92,52	NN (k = 10)	93,47	93,47	82,04	89,66
DT	95,92	95,10	93,88	94,97	DT	94,69	93,88	92,24	93,61
SVM	95,92	95,92	94,29	95,37	SVM	95,10	94,29	93,47	94,29
R+DRAH	99,59	99,59	97,96	99,05	R+DRAH	99,18	96,73	95,51	97,14
F. sensor 3	CV10	CV4	CV2	Mean	F. sensor 9	CV10	CV4	CV2	Mean
NB	95,10	95,10	95,10	95,10	NB	93,47	93,47	93,47	93,47
NN (k = 10)	93,47	93,47	89,80	92,24	NN (k = 10)	91,84	91,84	88,16	90,61
DT	94,29	93,47	92,24	93,33	DT	92,24	92,65	91,43	92,11
SVM	93,88	93,47	93,47	93,61	SVM	92,65	92,65	91,84	92,38
R+DRAH	99,18	99,18	97,96	98,77	R+DRAH	99,18	99,19	99,18	99,18
F. sensor 5	CV10	CV4	CV2	Mean	F. sensor 14	CV10	CV4	CV2	Mean
NB	93,88	94,29	94,29	94,15	NB	93,47	93,47	93,47	93,47
NN (k = 10)	93,88	93,88	93,88	93,88	NN (k = 10)	91,84	91,84	84,49	89,39
DT	93,47	92,65	92,24	92,79	DT	92,24	92,24	90,61	91,70
SVM	93,88	93,06	93,47	93,47	SVM	93,06	92,24	91,02	92,11
R+DRAH	98,78	96,33	95,10	96,74	R+DRAH	99,18	96,73	96,73	97,55

6. Conclusions

The sensors' optimization can produce the benefit to reduce costs from a technological perspective, maintaining accuracy for activity recognition, whilst, moreover, has the additional benefit of reducing sensor data and the computational complexity. In this paper, the use of feature selection methods has been proposed in order to optimize the set of sensors of a smart environment in conjunction with R+DRAH, which has been adapted of the RIMER+ for manage binary sensor data in a smart environment. A case study applying two feature selection methods, based on consistency and dependence measure for a smart environment with 14 sensors, has been carried out. Thus, two optimizations with 7 and 10 sensors have been obtained in which R+DRAH have been shown as the suitable classifier against four of the most popular data-driven approaches for activity recognition in situations where an essential sensor fails. Our future works are focused on extending the experimentation in order to offer performance statistics with different dataset according to the dataset size (small and large), the balance among the number of classes (balanced and unbalanced) and, finally, the presence of noise and inconsistencies.

Acknowledgments

This contribution has been supported by research projects: TIN2015-66524-P and UJA2014/06/14. Invest Northern Ireland is acknowledged for partially supporting this project under the Competence Centre Program Grant [RD0513853](#) Connected Health Innovation Centre.

Supplementary materials

Supplementary material associated with this article can be found, in the online version, at [doi:10.1016/j.micpro.2016.10.007](#).

References

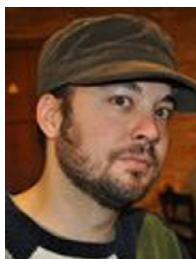
- [1] M. Alam, E. Hamida, Surveying wearable human assistive technology for life and safety critical applications: standards, challenges and opportunities, *Sensors (Switzerland)* 14 (5) (2014) 9153–9209.

- [2] A. Calzada, J. Liu, C. Nugent, H. Wang, L. Martinez, Sensor-based activity recognition using extended belief rule-based inference methodology, in: *Engineering in Medicine and Biology Society (EMBC), 2014 36th Annual International Conference of the IEEE, IEEE, 2014*, pp. 2694–2697.
- [3] A. Calzada, J. Liu, H. Wang, A. Kashyap, Dynamic rule activation for extended belief rule bases, in: *International Conference on Machine Learning and Cybernetics (ICMLC), 4, IEEE, 2013*, pp. 1836–1841.
- [4] A. Calzada, J. Liu, H. Wang, A. Kashyap, A new dynamic rule activation method for extended belief rule-based systems, *Knowledge and Data Engineering, IEEE Transactions on*, 2014 PP(99), 1–1, [doi:10.1109/TKDE.2014.2356460](#).
- [5] L. Chen, J. Hoey, C. Nugent, D. Cook, Z. Yu, Sensor-based activity recognition, *Syst. Man Cybern. Part C* 42 (6) (2012) 790–808.
- [6] L. Chen, C.D. Nugent, H. Wang, A knowledge-driven approach to activity recognition in smart homes, *IEEE Trans. Knowl. Data Eng.* 24 (6) (2012) 961–974.
- [7] Y. Chen, E. Garcia, M. Gupta, A. Rahimi, L. Cazzanti, Similarity-based classification: concepts and algorithms, *J. Mach. Learn. Res.* 10 (2009) 747–776.
- [8] D. Cook, J. Augusto, V. Jakkula, Ambient intelligence: technologies, applications, and opportunities, *Pervasive Mob. Comput.* 5 (4) (2009) 277–298, [doi:10.1016/j.pmcj.2009.04.001](#).
- [9] D. Cook, M. Schmitter-Edgecombe, Assessing the quality of activities in a smart environment, *Methods Inf. Med.* 48 (5) (2009) 480–485, [doi:10.3414/ME0592](#).
- [10] C. Cortes, V. Vapnik, Support vector networks, *Mach. Learn.* 20 (1995) 273–297.
- [11] T. Cover, P. Hart, Nearest neighbor pattern classification, *IEEE Trans. Inf. Theory* 13 (1) (1967) 21–27, [doi:10.1109/TIT.1967.1053964](#).
- [12] B. Das, D. Cook, M. Schmitter-Edgecombe, A. Seelye, Puck: an auto mated prompting system for smart environments: toward achieving auto mated prompting-challenges involved, *Pers. Ubiquitous Comput.* 16 (7) (2012) 859–873, [doi:10.1007/s00779-011-0445-6](#).
- [13] M. Dash, H. Liu, Feature selection for classification, *Intell. Data Anal.* 1 (3) (1997) 131–156.
- [14] M. Dash, H. Liu, Consistency-based search in feature selection, *Artif. Intell.* 151 (1) (2003) 155–176.
- [15] P.A. Devijver, J. Kittler, *Pattern Recognition: A Statistical Approach*, 761, Prentice-Hall, London, 1982.
- [16] P. Domingos, M. Pazzani, On the optimality of the simple bayesian classifier under zero-one loss, *Mach. Learn.* 29 (1997) 103–137.
- [17] H. Fang, L. He, H. Si, P. Liu, X. Xie, Human activity recognition based on feature selection in smart home using back-propagation algorithm, *ISA Trans.* 53 (5) (2014) 1629–1638, [doi:10.1016/j.isatra.2014.06.008](#).
- [18] H. Fang, R. Srinivasan, D. Cook, Feature selections for human activity recognition in smart home environments, *Int. J. Innovative Comput. Inf. Control* 8 (5 B) (2012) 3525–3535.
- [19] K.D. Feuz, D.J. Cook, C. Rosasco, K. Robertson, M. Schmitter-Edgecombe, Automated detection of activity transitions for prompting, *IEEE Trans. Hum.-Mach. Syst.* (2014), [doi:10.1109/THMS.2014.2362529](#).
- [20] H. Ghasemzadeh, N. Amini, R. Saeedi, M. Sarrafzadeh, Power-aware computing in wearable sensor networks: an optimal feature selection, *IEEE Trans. Mobile Comput.* 14 (4) (2015) 800–812, [doi:10.1109/TMC.2014.2331969](#).
- [21] T. Gu, L. Wang, Z. Wu, X. Tao, J. Lu, A pattern mining approach to sensor-based human activity recognition, *IEEE Trans. Knowl. Data Eng.* 23 (9) (2011) 1359–1372.

- [22] M. Hall, E. Frank, G. Holmes, B. Pfahringer, P. Reutemann, I.H. Witten, The weka data mining software: an update, *ACM SIGKDD Explorations Newsl.* 11 (1) (2009) 10–18.
- [23] M.A. Hall, Correlation-based feature subset selection for Ph.D. thesis, University of Waikato, Hamilton, New Zealand, 1998.
- [24] R.W. Hamming, Error detecting and error correcting codes, *Bell Syst. Tech. J.* 29 (2) (1950) 147–160.
- [25] L. Holder, D. Cook, Automated activity-aware prompting for activity initiation, *Gerontechnology* 11 (4) (2013) 534–544, doi:10.4017/gt.2013.11.4.005.00.
- [26] G.H. John, R. Kohavi, K. Pfleger, Irrelevant features and the subset selection problem, in: *ICML*, vol. 94, 1994, pp. 121–129.
- [27] K. Kira, L.A. Rendell, The feature selection problem: traditional methods and a new algorithm, in: *AAAI*, vol. 2, 1992, pp. 129–134.
- [28] R. Kohavi, D. Sommerfield, Feature subset selection using the wrapper method: overfitting and dynamic search space topology, in: *KDD*, 1995, pp. 192–197.
- [29] D. Koller, M. Sahami, Toward Optimal Feature Selection, *Stanford InfoLab*, 1996.
- [30] I. Kononenko, Estimating attributes: analysis and extensions of relief, in: *Machine Learning: ECML-94*, Springer, 1994, pp. 171–182.
- [31] M. Lee, A. Dey, Sensor-based observations of daily living for aging in place, *Pers. Ubiquitous Comput.* 19 (1) (2015) 27–43, doi:10.1007/s00779-014-0810-3.
- [32] H. Liu, R. Setiono, Feature selection and classification—a probabilistic wrap per approach, in: *Proceedings of 9th International Conference on Industrial and Engineering Applications of AI and ES*, 1997, pp. 419–424.
- [33] H. Liu, R. Setiono, et al., A probabilistic approach to feature selection—a filter solution, in: *ICML*, 96, Citeseer, 1996, pp. 319–327.
- [34] J. Liu, L. Martinez, A. Calzada, H. Wang, A novel belief rule base representation, generation and its inference methodology, *Knowl.-Based Syst.* 53 (2013) 129–141, doi:10.1016/j.knosys.2013.08.019.
- [35] S. McKeever, J. Ye, L. Coyle, C. Bleakley, S. Dobson, Activity recognition using temporal evidence theory, *J. Ambient Intell. Smart Environ.* 2 (3) (2010) 253–269, doi:10.3233/AIS-2010-0071.
- [36] F. Miao, Y. He, J. Liu, Y. Li, I. Ayoola, Identifying typical physical activity on smartphone with varying positions and orientations, *BioMed. Eng. Online* 14 (1) (2015) Art. no. 32.
- [37] M. Modrzejewski, Feature selection using rough sets theory, in: *Machine Learning: ECML-93*, Springer, 1993, pp. 213–226.
- [38] A.N. Mucciardi, E.E. Gose, A comparison of seven techniques for choosing subsets of pattern recognition properties, *IEEE Trans. Comput.* 20 (9) (1971) 1023–1031.
- [39] P.M. Narendra, K. Fukunaga, A branch and bound algorithm for feature subset selection, *IEEE Trans. Comput.* 100 (9) (1977) 917–922.
- [40] F. Ordóñez, P. Toledo, A. Sanchis, Activity recognition using hybrid generative/discriminative models on home environments using binary sensors, *Sensors* 13 (5) (2013) 5460–5477.
- [41] L. San Martín, V. Peláez, R. González, A. Campos, V. Lobato, Environmental user-preference learning for smart homes: an autonomous approach, *J. Ambient Intell. Smart Environ.* 2 (3) (2010) 327–342.
- [42] J.C. Schlimmer, et al., Efficiently inducing determinations: a complete and systematic search algorithm that uses optimal pruning, in: *ICML*, Citeseer, 1993, pp. 284–290.
- [43] J.W. Shavlik, T.G. Dietterich, *Readings in Machine Learning*, Morgan Kaufmann, 1990.
- [44] J. Sheinvald, B. Dom, W. Niblack, A modeling approach to feature selection, in: *Pattern Recognition, 1990. Proceedings 10th International Conference on*, 1, IEEE, 1990, pp. 535–539.
- [45] G. Smith, S.D. Sala, R. Logie, E. Maylor, Prospective and retrospective memory in normal aging and dementia: a questionnaire study, *Memory* 8 (2000) 311–321.
- [46] H. Stoppiglia, G. Dreyfus, R. Dubois, Y. Oussar, Ranking a random feature for variable and feature selection, *J. Mach. Learn. Res.* 3 (2003) 1399–1414.
- [47] S. Szewczyk, K. Dwan, B. Minor, B. Swedlove, D. Cook, Annotating smart environment sensor data for activity learning, *Technol. Health Care* 17 (3) (2009) 161–169, doi:10.3233/THC-2009-0546.
- [48] J. Van Hoof, E. Wouters, H. Marston, B. Vanrumste, R. Overdiep, Ambient assisted living and care in the Netherlands: the voice of the user, *Int. J. Ambient Comput. Intell.* 3 (4) (2011) 25–40.
- [49] T. Van Kasteren, A. Noulas, G. Englebienne, B. Krose, Accurate activity recognition in a home setting, in: *Proceedings of the 10th International Conference on Ubiquitous Computing*, ACM, 2008, pp. 1–9.
- [50] Wang, L., Gu, T., Tao, X., Lu, J. A hierarchical approach to real-time activity recognition in body sensor networks (2012) *Pervasive Mobile Comput.*, 8 (1), pp. 115–130.
- [51] J. Yang, J. Liu, J. Wang, H. Sii, H. Wang, Belief rule-base inference methodology using the evidential reasoning approach RIMER, *IEEE Trans. Syst. Man Cybern. Part A* 36 (2) (2006) 266–285, doi:10.1109/TSMCA.2005.851270.
- [52] J.B. Yang, D.L. Xu, On the evidential reasoning algorithm for multiple at tribute decision analysis under uncertainty, *IEEE Trans. Syst. Man Cybern. Part A* 32 (3) (2002) 289–304, doi:10.1109/TSMCA.2002.802746.



M. Espinilla was born in 1983. She received the M.Sc. and Ph.D. degrees, both in Computer Science, from the University of Jaén (Jaén, Spain), in 2006 and 2009, respectively. She is currently Associate Professor in the Department of Computer Systems at University of Jaén. She has authored or co-authored over 60 publications in peer-reviewed international journals and she is also in the editor board of an international journals. She has been invited to give talks at universities and international conferences, and being chair or co-chair, program committee for a number of international conferences. Her current research interests include recommender systems and smart environments, linguistic preference modeling, decision making, fuzzy logic-based systems and sensory evaluation.



J. Medina was born in 1983. HE received the M.Sc. and Ph.D. degrees, both in Computer Science, from the University of Granada (Granada, Spain), in 2006 and 2010, respectively. He has authored or co-authored over 20 publications in peer-reviewed international journals. His current research interests include heterogeneous architectures, smart environments, cyber-physical systems, fuzzy logic.



A. Calzada received the B.Sc. degree in Computer Sciences and the M.Sc. degree in 3D computer graphics from the University of Jaen, Jaen, Spain, in 2008 and 2010, respectively. He received the Ph.D. degree in spatial decision support systems from the University of Ulster, UK, in 2014. His current research interests include spatial analysis, DSS, GIS and health-care applications. He has 17 peer-reviewed publications in international journals and conference proceedings.



J. Liu received the B.Sc. and M.Sc. degrees in applied mathematics, and the Ph.D. degree in information engineering from Southwest Jiaotong University, Chengdu, China, in 1993, 1996, and 1999, respectively. He is currently a senior lecturer in Computer Science at University of Ulster, Newtownabbey, United Kingdom. He has been working in the field of AI for many years. His current research interests include nonclassical logic and reasoning methods for intelligent systems; intelligent DSSs and information management, with applications in management, engineering, and industry field, etc. (e.g., safety and risk analysis; situation awareness and emergency systems; video scenario recognition); information fusion and data combinations; data mining and KBS; applied computational intelligence for uncertainty analysis and optimization. He is a member of the IEEE.



L. Martínez was born in 1970. He received the M.Sc. and Ph.D. degrees in Computer Sciences, both from the University of Granada, Spain, in 1993 and 1999, respectively. Currently, he is Full Professor of Computer Science Department and Head of ICT Research Centre at the University of Jaén. His current research interests are linguistic preference modeling, decision making, fuzzy logic based systems, computer aided learning, sensory evaluation, recommender systems and electronic commerce. He co-edited nine journal special issues on fuzzy preference modeling, soft computing, linguistic decision making and fuzzy sets theory and published more than 70 papers in journals indexed by the SCI as well as 30 book chapters and more than 100 contributions in International Conferences related to his areas. It is remarkable that he has been main researcher in 12 R&D projects.



C. Nugent is Professor of Biomedical Engineering at the University of Ulster. He received a Bachelor of Engineering in Electronic Systems and DPhil in Biomedical Engineering both from the University of Ulster. His research within biomedical engineering addresses the themes of Technologies to Support Independent Living, Medical Decision Support Systems specifically in the domain of computerized electrocardiology and the development of Internet based healthcare models. He has published extensively in these areas with the work spanning theoretical, clinical and biomedical engineering domains.