

A Retraining Methodology for Enhancing Agent Intelligence

Pericles A. Mitkas, Andreas L. Symeonidis and Ioannis N. Athanasiadis

Electrical and Computer Engineering Dept.,

Aristotle University of Thessaloniki, GR541 24, Thessaloniki, Greece

and

Intelligent Systems and Software Engineering Laboratory,

Informatics and Telematics Institute – CERTH, GR570 01, Thessaloniki, Greece

Tel.: +30 2310 996 349, Fax: +30 2310 996 398

Email: mitkas@eng.auth.gr, asymeon@ee.auth.gr, ionathan@ee.auth.gr

Abstract—Data mining has proven a successful gateway for discovering useful knowledge and for enhancing business intelligence in a range of application fields. Incorporating this knowledge into already deployed applications, though, is highly impractical, since it requires reconfigurable software architectures, as well as human expert consulting. In an attempt to overcome this deficiency, we have developed Agent Academy, an integrated development framework that supports both design and control of multi-agent systems (MAS), as well as agent training. We define agent training as the automated incorporation of logic structures generated through data mining into the agents of the system. The increased flexibility and cooperation primitives of MAS, augmented with the training and retraining capabilities of Agent Academy, provide a powerful means for the dynamic exploitation of data mining extracted knowledge. In this paper, we present the methodology and tools for agent retraining. Through experimental results with the Agent Academy platform, we demonstrate how the extracted knowledge can be formulated and how retraining can lead to the improvement - in the long run - of agent intelligence.

I. INTRODUCTION

In a highly complex and competitive business environment, companies must take swift, yet sound decisions that rely on corporate logic and domain knowledge. Diffusing, however, this knowledge into the software processes of the company is a difficult task, which requires reconfigurable software architectures and human expert involvement. A unified approach for discovering useful corporate knowledge and embedding it into the company's software would therefore be highly desirable.

The most dominant solution for discovering non-trivial, implicit, previously unknown and potentially useful [7] knowledge is Data Mining (DM), a technology developed to support the tremendous data outburst and the imperative need for the interpretation and exploitation of massive data volumes. DM issues concerning data normalization, algorithm

complexity and scalability, result validation and comprehension have already been successfully dealt with [1], [14], [24]. Numerous approaches have been adopted for the realization of autonomous and versatile DM tools, which feature all the appropriate pre- and post-processing steps that constitute the process of Knowledge Discovery in Databases (KDD) [6], [7], [20]. The ultimate goal of DM is the extraction of a valid knowledge model (i.e. decision rules, decision trees, association rules, clusters, etc) that best describes the data trends and underlying patterns.

On the other hand, despite the support corporate software provides on process coordination and data organization, it often - especially legacy software - lacks advanced capabilities, thus limiting company competitiveness. The increasing demand for sophisticated software comprising collaborative, yet autonomous, units, which can regulate, control and organize all distributed activities in the business process, has oriented AI researchers towards the employment of Agent Technology (AT) in a variety of disciplines [15], [25]. The versatility and generic nature of the AT paradigm has shown that inherently distributed problems, which require the synergy of a number of elements for their solution, can be efficiently implemented as a multi-agent system (MAS) [8]. The coupling of DM and AT principles is, therefore, expected to enable the development of highly reconfigurable systems that incorporate domain knowledge and provide decision making capabilities. The exploitation of useful knowledge extracted by the use of DM may considerably improve agent infrastructures, while also increasing reusability and minimizing customization costs.

Going briefly through related work, attempts to couple DM and AT already exist. Galitsky and Pampapathi [12] use both inductive (DM) and deductive (AT) approaches, in order to model and process the claims of unsatisfied customers. Deduction is used for describing the behaviors of agents (humans or companies), for which we have complete information, while induction is used to predict the behavior of agents, whose actions are uncertain to us. A more theoretical approach on the way DM extracted knowledge can contribute to AT performance has been presented by Fernandes [9],

who attempts to model the notions of data, information, and knowledge in purely logical terms, in order to integrate inductive and deductive reasoning into one inference engine. Kero et al. [17], finally, propose a DM model that utilizes both inductive and deductive components. They model the discovery of knowledge as an iteration between high-level, user-specified patterns and their elaboration to (deductive) database queries. One the other hand, they define the notion of a meta-query that performs the (inductive) analysis of these queries and their transformation to modified, ready-to-use knowledge.

Advancing on earlier research efforts to couple the two technologies, we have developed Agent Academy [2], [19], an integrated platform for developing MAS architectures and for enhancing their functionality and intelligence through the use of DM techniques.

Agent Academy (AA) agents are developed over the Java Agent Development Framework (JADE) [5], which conforms to the FIPA specifications [10]. The MAS ontologies are developed through the *Agent Factory* module of AA. Data to be mined are imported to AA in XML format and are handled by the *Data Miner* module, a DM suite that expands the Waikato Environment for Knowledge Analysis (WEKA) tool [24]. The extracted knowledge structures are represented in PMML (Predictive Model Markup Language), a language that efficiently describes clustering, classification and association rule knowledge models [13]. The resulting knowledge is then incorporated into the agents of the MAS by the use of the *Agent Training Module* of AA. All necessary data files (application data, agent behavior data, knowledge structures, and agent ontologies) are stored into AA's main database, the *Agent Use Repository*. Agents can be periodically recalled for retraining, since appropriate agent tracking tools have been incorporated into Agent Academy, in order to monitor agent activity after their initial deployment.

It is through retraining that we intend to show how certain DM techniques can be used to augment agent intelligence and therefore improve MAS overall performance. The rest of the paper is organized as follows: Section II defines the formal model for training and retraining agents through Agent Academy and specifies all the necessary notations. Section II outlines the already developed mechanism for training and retraining, while Section IV describes the various training and retraining options for the improvement of agent intelligence and presents some indicative experimental results. Finally, Section V summarizes and concludes the paper.

II. A FORMAL MODEL FOR AGENT (RE)TRAINING

To develop a MAS application with Agent Academy, the steps below must be followed:

- a. Creation of the application ontology.
- b. Creation of agent behaviors.
- c. Creation of agent types, realizing the created behaviors.
- d. Data mining on agent type-specific datasets.
- e. Generation of knowledge models for each agent type.
- f. Creation of the application agents (of the different agent types)

- g. Incorporation of the extracted knowledge models into the corresponding agents.
- h. MAS instantiation.
- i. Agent monitoring.
- j. Retraining of the MAS agents on a periodic basis.

Let O be the ontology of the MAS. Let $A = \{A_1, A_2, \dots, A_S\}$ be the set of attributes described in O and defined on $\mathcal{D}$, the application data domain. Let $D \subseteq \mathcal{D}$ be a set of application data, where each dataset tuple is a vector $T = \{t_1, t_2, \dots, t_S\}$, and $t_s, s = 1 \dots S$, is the value of the corresponding attribute A_s . Missing values are allowed within T .

In order to initially train a certain type $Q_i, i = 1, \dots, k$, of application agents, we use a subset of the application dataset, containing the attributes that are relevant to this specific type. We therefore define $D_{IQ_i} \subseteq D_{IT}$, where D_{IQ_i} is the initial training dataset for agent type Q_i , and D_{IT} is the initial application dataset. In most cases $D_{IT} = D$. For each Q_i we perform data mining on its corresponding dataset D_{IQ_i} , in order to extract a useful knowledge model KM_o ($o = 1, \dots, p$). This model will be incorporated into all agents of the same type $Q_i(j), j = 1, \dots, m^1$. We then instantiate the MAS and monitor its agents.

During the retraining phase, each agent can be retrained individually. The available datasets include: a) the initial dataset D_{IT} , b) a new non-agent dataset² D_{NQ_i} , and c) all the datasets $D_{Q_i(j)}$, each containing the tuples representing the actions (decisions) taken by the respective agent. It must be denoted that $D_{Q_i} = D_{Q_i(1)} \oplus D_{Q_i(2)} \oplus \dots \oplus D_{Q_i(m)}$. The symbol $\oplus$ represents the concatenation of two datasets, an operation that preserves multiple copies of tuples. There are five different options of agent retraining, with respect to the datasets used:

- A. $D_{IQ_i} \oplus D_{NQ_i}$. Retrain the agent using the initial dataset along with a new, non-agent dataset D_{NQ_i} .
- B. $D_{NQ_i} \oplus D_{Q_i}$. Retrain the agent using a non-agent dataset D_{NQ_i} along with D_{Q_i} , a dataset generated by all the Q_i -type agents of the application. AA agents are monitored and their actions are recorded, in order to construct the D_{Q_i} dataset.
- C. $D_{IQ_i} \oplus D_{NQ_i} \oplus D_{Q_i}$. Retrain the agent using all the available datasets.
- D. $D_{IQ_i} \oplus D_{Q_i}$. Use the initial dataset D_{IQ_i} along with the agent generated data.
- E. $D_{IQ_i} \oplus D_{Q_i(j)}$. Use the initial dataset D_{IQ_i} along with $D_{Q_i(j)}$, the generated data of the j -th agent.

A schematic representation of the training and retraining procedure is given in Fig. 1.

Through AA and its training/retraining capabilities the user can formulate and augment agents' intelligence. AA supports a variety of both supervised (classification) and unsupervised

¹It should be denoted that more than one knowledge models can be incorporated into an agent type.

²We define a non-agent dataset, as the set that contains data related to the actions of agents, but has not been produced by them. For example, data may come from non-agent based applications that are still active.

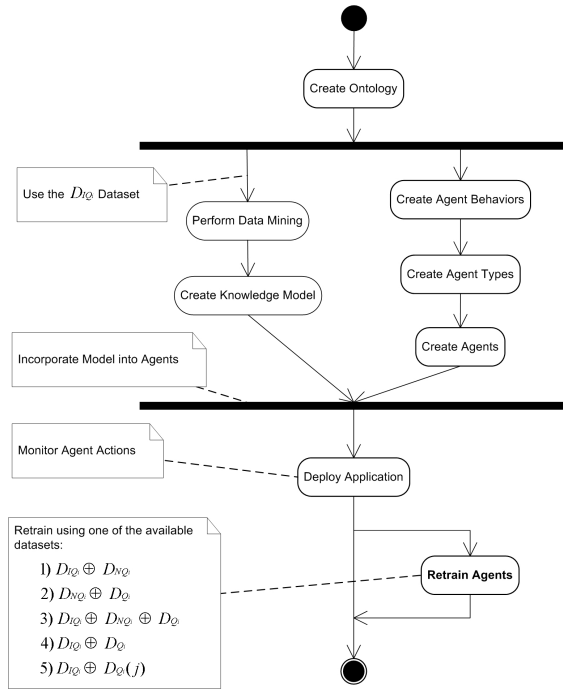


Fig. 1. Training and Retraining the agents of a MAS

learning (clustering, association rule extraction) DM techniques: ID3, C4.5, CLS, and FLR for classification, Apriori, DHP, and DIC, for association rule extraction and K-Means, PAM, EM, and κ -Profile for clustering³.

III. THE TRAINING AND RETRAINING MECHANISM

In order to enable the incorporation of knowledge into agents, we have implemented *Data Miner* as an agent-oriented tool. It is a DM suite that supports the application of a variety of classification, clustering and association rule extraction algorithms on application-specific and agent-behavior-specific data. *Data Miner* can also incorporate the extracted decision models into the AF produced agents, augmenting that way their intelligence. Apart from being a core component of the AA platform, the *Data Miner* can also function as a standalone DM tool. The mechanism for embedding rule-based reasoning capabilities into agents is illustrated in Fig. 2. Data, either application-specific or agent-behavior-specific, enter the module in XML format. Each data file contains information on the name of the agent the file belongs to and on the decision structure of the agent it will be applied on. The XML file is then inserted into the Preprocessing Unit of the *Data Miner*, where all the necessary data selection and data cleaning tasks take place. Next, data are forwarded to the Miner, where the user decides on the DM technique, as well as on the specific algorithm to employ. After DM is performed, the results are sent to the Evaluator, which is responsible for the validation and visualization of the extracted model. If

³The FLR and κ -Profile algorithms are novel algorithms, developed within the context of Agent Academy. More information on these algorithms can be found at [16], [2]

the user accepts the constructed model, a PMML document describing the knowledge model is generated. This document expresses the referencing mechanism of the agent we intend to train. The resulting decision model is then translated to a set of facts executed by a rule engine. The implementation of the rule engine is realized through the Java Expert System Shell (JESS) [11], which is a robust mechanism for executing rule-based agent reasoning. The execution of the rule engine transforms the *Data Miner* extracted knowledge into a living part of the agent's behavior.

After the MAS is instantiated, the user has the ability to monitor AA agents and their decisions. The decisions of each agent are stored separately in Agent Academy and form the $D_{Q_i(j)}$ datasets. The user can then decide, as mentioned in Section II, on the dataset s/he would like to perform retraining on.

IV. AUGMENTING AGENT INTELLIGENCE

A. Different Retraining Approaches

Retraining is performed in order to either increase or refine agent intelligence. By reapplying data mining on a new or more complete dataset, the user expects to derive a more accurate and/or more efficient knowledge models. The five retraining options defined earlier, can be classified into two main approaches: a) the type-specific, which focuses on the improvement of an agent type Q_i (options A-D) and

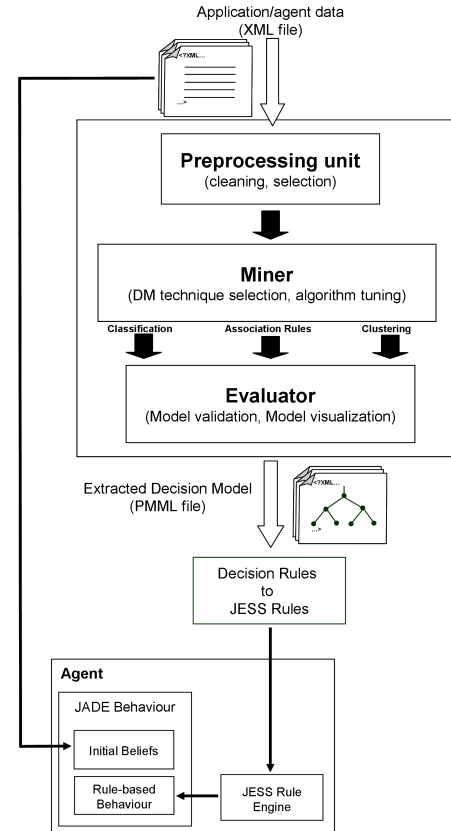


Fig. 2. The agent training/retraining mechanism in Agent Academy

b) the agent-specific, which focuses on the refinement of intelligence of an individual agent $Q_i(j)$, the j -th agent of type i (option E).

It should be denoted that we differentiate on the way we define “intelligence improvement”, since AA provides both supervised and unsupervised learning DM techniques. In the case of classification, improvement can be measured by evaluating the knowledge model extracted metrics (mean-square error, accuracy, etc.). In the case of clustering and association rule extraction, intelligence augmentation is determined by external evaluation functions. The classification algorithms provided by the AA platform are decision tree (DT) extraction algorithms. The basic prerequisites for the proper application of a DT construction algorithm are the existence of a distinct set of classes and the availability of training data. All the DT algorithms supported by the AA platform are criterion gain algorithms, i.e. algorithms that decide on the construction of the DT, according to the minimization (or maximization) of a certain criterion. In the case of ID3 and C4.5, this criterion is the information gain [21], in the case of CLS, it is record sorting [14], and in the case of FLR, the criterion is the inclusion measure [16].

The clustering algorithms provided by AA are partitioning algorithms (PAs). The objective of a PA is the grouping of the data provided into discrete clusters. Data must have high intra-cluster and low inter-cluster similarity. The splitting criterion in a PA is the Euclidean distance between data [18]. Finally, the association rule extraction algorithms provided by AA are mainly focused on transactional datasets. In order for these algorithms to decide on the strongest associations, two metrics are considered: support and confidence [3].

B. Training and Retraining in the case of Supervised Learning

Although the splitting criteria are different, all of the above-mentioned classification algorithms are applied in a similar manner. While we focus on the information gain criterion, employed by C4.5 and ID3, the approach followed can be easily adjusted to other classification algorithms of the platform. The information gain expected when splitting dataset D with respect to attribute A_s , $A_s \in A$ is given by Eq. 1:

$$Gain(D, A_s) = Info(D) - Info(D, A_s) \quad (1)$$

$Info(D)$ is the information needed to classify D with respect to C predefined distinct classes c_r (for $r = 1, \dots, C$), and is given by Eq. 2:

$$Info(D) = - \sum_{r=1}^C p(c_r) \log_2 p(c_r) \quad (2)$$

with $p(c_r)$ the ratio of D tuples that belong to class c_r . $Info(D, A_s)$ is the information needed in order to classify D , after its partitioning into V subsets D_v , $v = 1, \dots, V$, with respect to the attribute A_s . $Info(D, A_s)$, which is also denoted as the Entropy of A , is given by Eq. 3:

$$Info(D, A_s) = - \sum_{v=1}^V \frac{|D_v|}{|D|} Info(D_v) \quad (3)$$

Splitting is conducted on the attribute that yields the maximum information gain.

1) *Initial Training*: When training takes place, classification is performed on D_{IQ_i} , the initial dataset for the specific agent type. The user can decide to split the dataset into a training and a testing (and/or validation) dataset or to perform n -fold cross-validation. To evaluate the success of the applied classification scheme, a number of statistical measures are calculated, i.e classification accuracy, mean absolute error, and confusion matrix. If the extracted knowledge model is deemed satisfactory, the user may accept it and store it, for incorporation into the corresponding Q_i -type agents.

2) *Retraining Q_i* : In the case of retraining agent-type Q_i , the relevant datasets are D_{IQ_i} , D_{NQ_i} and D_{Q_i} . Retraining option C ($D_{IQ_i} \oplus D_{NQ_i} \oplus D_{Q_i}$) is the most general, containing all the available data for the specific agent type, while options A and D are subsets of option C. They are differentiated, however, since option D is particularly interesting and deserves special attention. When using datasets D_{IQ_i} and D_{NQ_i} , the user may choose among the retraining options illustrated in Table I.

TABLE I
RETRAINING OPTIONS FOR $D_{IQ_i} \oplus D_{NQ_i}$

Option	Dataset		Causality
	D_{IQ_i}	D_{NQ_i}	
A-1	Training	Testing	Initial model validation
A-2	Testing	Training	Model investigation on Data Independency
A-3	Concatenation and Cross-validation		New Knowledge Model

The user decides on which knowledge model to accept, based on its performance. Nevertheless, in the $D_{IQ_i} \oplus D_{NQ_i}$ case, best model performance is usually observed when option A-3 is selected. The inductive nature of classification dictates that larger training datasets lead to more efficient knowledge models.

The retraining options when the $D_{NQ_i} \oplus D_{Q_i}$ dataset is selected are illustrated in Table II:

TABLE II
RETRAINING OPTIONS FOR $D_{NQ_i} \oplus D_{Q_i}$

Option	Dataset		Causality
	D_{NQ_i}	D_{Q_i}	
B-1	Training	Testing	Indirect initial model validation
B-2	Concatenation and Cross-validation		New Knowledge Model discovery

When retraining an agent with the $D_{NQ_i} \oplus D_{Q_i}$ dataset, it is important to notice that the only information we have on the training dataset D_{IQ_i} is indirect, since D_{Q_i} is formatted based on the knowledge model the agents follow, a model inducted by the D_{IQ_i} dataset. This is why the validation of the initial model is indirect. If the D_{NQ_i} -extracted model is

similar to the D_{IQ_i} -extracted model testing accuracy is very high.

The fact that D_{Q_i} is indirectly induced by D_{IQ_i} , does not allow testing D_{Q_i} on D_{IQ_i} . Nevertheless, concatenation of the datasets can lead to more efficient and smaller classification models. Since class assignment within D_{Q_i} (the agent decisions) is dependent on the D_{IQ_i} -extracted knowledge model, a “bias” is inserted in the concatenated $D_{IQ_i} \oplus D_{Q_i}$ dataset. Let attribute A_i be the “biased” attribute and C_i the supported class. While recalculating the information gain for the $D_{IQ_i} \oplus D_{Q_i}$ dataset, we observe that the increase of $Info(D)$ is cumulative (Eq. 2), while the increase of $Info(D, A_s)$ is proportional (Eq. 3) and therefore $Gain(D, A_s)$ is increased. Clearer decisions on the splitting attributes according to the frequency of occurrence of A_i in conjunction to C_i are derived, thus leading to more efficient knowledge models. Table III illustrates the available retraining options for the corresponding dataset.

TABLE III
RETRAINING OPTIONS FOR $D_{IQ_i} \oplus D_{Q_i}$

Option	Dataset		Causality
	D_{IQ_i}	D_{Q_i}	
D-1	Concatenation and Cross-validation		More application-efficient Knowledge Model

In the most general case, where all datasets (D_{IQ_i} , D_{NQ_i} and D_{Q_i}) are available, the retraining options are similar to the ones proposed for the already described subsets and similar restrictions apply. Table IV illustrates these options.

TABLE IV
RETRAINING OPTIONS FOR $D_{IQ_i} \oplus D_{NQ_i} \oplus D_{Q_i}$

Option	Dataset			Causality
	D_{IQ_i}	D_{Q_i}	D_{NQ_i}	
C-1	Training	Testing	Testing	Initial model validation
C-2	Testing	Testing	Training	Model investigation on Data Dependency
C-3	Concatenation and Training		Testing	New Knowledge Model (more efficient) validation
C-4	Concatenation and Cross-validation			New Knowledge Model

3) *Retraining $Q_i(j)$* : When retraining a specific agent, the user is interested in the refinement of its intelligence in relation to the working environment. Let us assume that we have trained a number of agents that decide on whether a game of tennis should be conducted, according to weather outlook, temperature, humidity and wind conditions (Weather dataset, [14], [24]), and have established these agents in different cities in Greece (Athens, Thessaloniki, Patra, Chania, etc). Although all these agents rely initially on a common knowledge model, weather conditions in Thessaloniki differ from those in Chania enough to justify refined knowledge models.

In this case, we have the option to perform agent-type retraining. By the use of the $D_{IQ_i} \oplus D_{Q_i(j)}$ dataset, it is possible to refine the intelligence of the j -th agent of type i . High frequency occurrence of a certain value t_s of attribute A_s (i.e. “High” humidity in Thessaloniki, “Sunny” outlook in Chania) may produce a more “case-specific” knowledge model. In a similar to the $D_{IQ_i} \oplus D_{Q_i}$ manner, it can be seen that an increase of $Info(D, A_s)$ can lead to a different knowledge model, which incorporates instance-specific information.

The analysis of different retraining options in the case of Classification indicates that there exist concrete success metrics that can be used to evaluate the extracted knowledge models and, thus, may ensure the improvement of agent intelligence.

C. Training and Retraining in the case of Unsupervised Learning

In the case of unsupervised learning, training and retraining success cannot be determined quantitatively. A more qualitative approach must be followed, to determine the efficiency of the extracted knowledge model, with respect to the overall goals of the deployed MAS.

1) *Initial Training*: To perform clustering, the user can either split the D_{IQ_i} dataset into a training and a testing subset or perform a classes-to-clusters evaluation, by testing the extracted clusters with respect to a class attribute defined in D_{IQ_i} . In order to evaluate the success of the clustering scheme, the mean square error and standard deviation of each cluster center are calculated. On the other hand, if the user decides to perform association rule extraction (ARE) on D_{IQ_i} , no training options are provided. Only the algorithm-specific metrics are specified and ARE is performed. In a similar to classification manner, if the extracted knowledge model (clusters, association rules) is favorably evaluated, it is stored and incorporated into the corresponding Q_i -type agents.

2) *Retraining by Clustering*: Clustering results are in most cases indirectly applied to the deployed MAS. In practice, some kind of an external exploitation function is developed, which somehow fires different agent actions in the case of different clusters. All the available datasets can therefore be used for both training and testing for Initial model validation, Model Data dependency investigation and New Knowledge Model discovery. A larger training dataset and more thorough testing can lead to more accurate clustering. Often retraining can result in the dynamic updating and encapsulation of dataset trends (i.e. in the case of customer segmentation). Retraining $A_i(j)$ can therefore be defined as a “case-specific” instance of retraining, where data provided by agent j , $D_{Q_i(j)}$, are used for own improvement.

3) *Retraining by Association Rule Extraction*: The ARE technique does not provide training and testing options. The whole input dataset is used for the extraction of the strongest association rules. Consequently, all available datasets (D_{IQ_i} , D_{NQ_i} , D_{Q_i} and $D_{Q_i(j)}$) are concatenated before DM is performed. This unified approach for retraining has a sole goal: to discover the strongest association rules between the

items t of D . In a similar to the clustering case manner, retraining $A_i(j)$ can be viewed as a “case-specific” instance of retraining.

V. EXPERIMENTAL RESULTS

In order to prove the added value of agent retraining, a number of experiments on Classification, Clustering and ARE were conducted. In this section, three representatives cases are discussed. These experiments are focused mainly on retraining by the use of the D_{Q_i} and $D_{Q_i(j)}$ datasets and illustrate the enhancement of agent intelligence.

A. Intelligent Environmental Monitoring System

The first experiment was performed for the O₃RTAA System, an agent-based intelligent environmental monitoring system developed for assessing ambient air-quality [4]. A community of software agents is assigned to monitor and validate multi-sensor data, to assess air-quality, and, finally, to fire alarms to appropriate recipients, when needed. Data mining techniques have been used for adding data-driven, customized intelligence into agents with successful results [16].

In this work, we focused on the Diagnosis Agent Type. Agents of this type are responsible for monitoring various air quality attributes including pollutants’ emissions and meteorological attributes. Each one of the Diagnosis Agent instances is assigned to monitor one attribute through the corresponding field sensor. In the case of sensor breakdown, Diagnosis Agents take control and perform an estimation of the missing sensor values using a data-driven Reasoning Engine, which exploits DM techniques.

One of the Diagnosis Agents is responsible for estimating missing ozone measurement values. This task is accomplished using a predictive model comprised of the predictors and the response. For the estimation of missing ozone values the predictors are the current values measured by the rest of the sensors, while the response is the level of the missing value (Low, Medium, or High). In this way, the problem has been formed as a classification task.

For training and retraining the Ozone Diagnosis Agent we used a dataset, labeled C2ONDA01 and supplied by CEAM, which contained data from a meteorological station in the district of Valencia, Spain. Several meteorological attributes and air-pollutant values were recorded on a quarter-hourly basis during the year 2001. There are approximately 35,000 records, with ten attributes per record plus the class attribute. The dataset was split into three subsets: one subset for initial training (D_{IQ_i}), a second subset for agent testing (D_{Q_i}) and another subset for validation (D_{Val}) containing around 40%, 35% and 25% of the data, respectively.

The initial training of the Diagnosis Agent was conducted using Quinlan’s C4.5 [21] algorithm for decision tree induction, using the D_{IQ_i} subset. This decision tree was embedded in the Diagnosis Agent and the agent used it for deciding on the records of the D_{Q_i} subset. Agent decisions along with the initial application data were used for retraining the Diagnosis Agent (Option D: $D_{IQ_i} \oplus D_{Q_i}$). Finally, the Diagnosis Agent with the updated decision tree was used

for deciding on the cases of the last subset (D_{Val}). The retrained Diagnosis Agent performed much better compared to the initial training model, as shown in Table V. The use of agent decisions included in D_{Q_i} has enhanced the Diagnosis Agent performance on the D_{Val} subset by 3.65%.

TABLE V
CLASSIFICATION ACCURACIES FOR THE DIAGNOSIS AGENT

Option	Dataset		
	D_{IQ_i}	D_{Q_i}	D_{Val}
Number of instances	11,641	10,000	7,414
Training	Used	73.58%	71.89%
Retraining		Used	74.66%

B. Speech Recognition Agents

This experiment was based on the “vowel” dataset of the UCI repository [23]. The problem in this case is to recognize a vowel spoken by an arbitrary speaker. This dataset is comprised of ten continuous primary features (derived from spectral data) and two discrete contextual features (the speaker’s identity and sex) and contains records for 15 speakers. The observations fall into eleven classes (eleven different vowels). The vowel problem was assigned to an agent community to solve. Two agents $Q_i(1)$ and $Q_i(2)$ were deployed to recognize vowels. Although of the same type, the two agents operate in different environments. This is why the dataset was split in the following way: The data of the first nine speakers (D_{IQ_i}) were used as a common training set for both $Q_i(1)$ and $Q_i(2)$. The records for the next two speakers were assigned to $Q_i(1)$ and those of the last two speakers were assigned to $Q_i(2)$.

The procedure followed was to evaluate the retraining performance of each one of the agents (Option E: $D_{IQ_i} \oplus D_{Q_i(j)}$). After initial training with D_{IQ_i} , each of the $Q_i(1)$ and $Q_i(2)$ was tested on one of the two assigned speakers, while the second speaker was used for the evaluation of the retraining phase. Quinlan’s C4.5 algorithm was applied. The classification accuracy, which is similar to that reported by P.D. Turney [22], is illustrated in Table VI.

TABLE VI
SPEECH RECOGNITION AGENTS CLASSIFICATION ACCURACY

	$Q_i(1)$		
	D_{IQ_i}	$D_{Q_i(1)}$	$D_{Val(1)}$
Number of speakers	9	1	1
Initial Training	Used	53.03%	46.97%
Retraining		Used	56.06%

	$Q_i(2)$		
	D_{IQ_i}	$D_{Q_i(2)}$	$D_{Val(2)}$
Number of speakers	9	1	1
Initial Training	Used	33.33%	28.78%
Retraining		Used	43.93%

It is obvious in this case that retraining using $D_{Q_i(j)}$ leads to considerable enhancement of the agents’ ability to decide

correctly. The decision models that are induced after the retraining procedure outperformed the validation speakers. The improvement by the mean of classification accuracy was improved by 36% in average.

C. The Iris Recommendation Agent

In order to investigate retraining in the case of clustering, we used the Iris UCI Dataset [23], a dataset widely used in pattern recognition literature. It has four numeric attributes describing the iris plant and one nominal attribute describing its class. The 150 records of the set were split into two subsets: one subset (75%) for initial training (D_{IQ_i}) and a second subset (25%) for agent testing (D_{Q_i}). Classes-to-clusters evaluation was performed on D_{IQ_i} and $D_{IQ_i} \oplus D_{Q_i}$ (Option D) and the performance of the resulted clusters was compared on the number of correctly classified instances of the dataset (Table VII).

TABLE VII
THE IRIS RECOMMENDATION AGENT SUCCESS

	Dataset		Correctly classified
	D_{IQ_i}	D_{Q_i}	
Number of instances	113	37	
Initial Training	Used	–	83.19%
Retraining		Used	88.67%

Again, retraining with the $D_{IQ_i} \oplus D_{Q_i}$ dataset leads to the improvement of clustering results. The new knowledge models obtained with the above retraining options can be easily incorporated into agents following the already implemented training/retraining mechanism, which is described next.

VI. CONCLUSIONS

Work presented in this paper explains how DM techniques can be successfully coupled with AT, leading to dynamically created agent intelligence. The concepts of training and retraining are formulated and special focus is given on retraining. Through this procedure, where DM is performed on new datasets (D_{NQ_i} , D_{Q_i} and $D_{Q_i(j)}$), refined knowledge is extracted and dynamically embedded into the agents. The different retraining options in the cases of Supervised and Unsupervised Learning are outlined in this paper and experimental results on different types of retraining are provided. Finally, the training and retraining mechanism is presented. Based on our research work we strongly believe that data mining extracted knowledge could and should be coupled with agent technology, and that training and retraining can indeed lead to more intelligent agents.

ACKNOWLEDGEMENTS

Work presented here has been partially supported by the European Commission through the IST initiative (IST project No 2000-31050).

REFERENCES

- [1] P. Adriaans and D. Zantige. *Data Mining*. Addison-Wesley, 1996.
- [2] The Agent Academy Consortium. Requirements and Specifications Document, 2000. Available at <http://AgentAcademy.itigr/>.
- [3] A. Amir, R. Feldman, and R. Kashi. A new and versatile method for association generation. *Journal of Information Systems*, 22(6-7):333–347, 1997.
- [4] I.N. Athanasiadis and P.A. Mitkas. An agent-based intelligent environmental monitoring system. *Management of Environmental Quality*, 15(3):238–249, 2004.
- [5] F. Bellifemine, A. Poggi, and G. Rimassa. Developing multi-agent systems with jade. In *Seventh International Workshop on Agent Theories, Architectures, and Languages*, 2000. Available at: <http://jade.cse.it>.
- [6] M.S. Chen, J. Han, and P.S. Yu. Data Mining: An overview from a database perspective. *IEEE Transactions on Knowledge and Data Engineering*, 8(6):866–883, 1996.
- [7] U. Fayyad, G. Piatetsky-Shapiro, and P. Smyth. Knowledge Discovery and Data Mining: Towards a unifying framework. In *Second International Conference on Knowledge Discovery and Data Mining*, pages 82–88, 1996.
- [8] J. Ferber. *Multi-Agent Systems - An introduction to Distributed Artificial Intelligence*. Addison-Wesley, 1999.
- [9] A.A.A. Fernandes. Combining inductive and deductive inference in knowledge management tasks. In *Proc. 11th International Workshop on Database and Expert Systems Applications, First International Workshop on Theory and Applications of Knowledge Management*, pages 1109–1114, 2000.
- [10] Foundation for Intelligent Physical Agents. FIPA specifications, 2000. Available at: <http://www.fipa.org/specs/>.
- [11] E. J. Friedman-Hill. *Jess in Action*. Manning Publications, 2003.
- [12] B. Galitsky and R. Pampapathi. Deductive and inductive reasoning for processing the claims of unsatisfied customers. In *Proc. 16th Int. Conf. on Industrial and Engineering Applications of Artificial Intelligence and Expert Systems*, pages 21–30, 2003.
- [13] Data Mining Group. Predictive Model Markup Language Specifications (PMML), ver. 2.0, 2001. Available at: <http://www.dmg.org>.
- [14] J. Han and M. Kamber. *Data Mining: Concepts and Techniques*. Morgan Kaufmann, 2001.
- [15] N.R. Jennings and M.J. Wooldridge, editors. *Agent technology: Foundations, applications and markets*. Springer Verlag, 1998.
- [16] V.G. Kaburlasos, I.N. Athanasiadis, and P.A. Mitkas. Fuzzy Lattice Reasoning (FLR) classifier and its application on improved estimation of ambient ozone concentration. *Submitted for publication*, 2004.
- [17] C. Kero, L. Russell, S. Tsur, and W.M. Shen. An overview of data mining technologies. In *The KDD Workshop in the 4th International Conference on Deductive and Object-Oriented Databases*, 1995.
- [18] J.B. MacQueen. Some methods for classification and analysis of multivariate observations. In *Proc. 5th Berkeley Symposium on Mathematical Statistics and Probability*, pages 281–297, 1967.
- [19] P.A. Mitkas, D. Kehagias, A.L. Symeonidis, and I.N. Athanasiadis. A Framework for Constructing Multi-Agent Applications and Training Intelligent Agents. *Agent-Oriented Software Engineering IV*, LCNS 2935. pages 96–109. Springer-Verlag, 2004.
- [20] G. Piatetsky-Shapiro and W.J. Frawley. *Knowledge Discovery in Databases*. MIT Press, 1992.
- [21] J.R. Quinlan. *C4.5: Programs for Machine Learning*. Morgan Kaufmann, 1993.
- [22] P.D. Turney. Robust classification with context-sensitive features. In *Proc. 6th International Conference on Industrial and Engineering Applications of Artificial Intelligence and Expert Systems*, pages 268–276, 1993.
- [23] UCI. Machine learning repository. Available at: www.ics.uci.edu/~mllearn/.
- [24] F. Witten and E. Frank. *Data Mining: Practical Machine Learning Tools and Techniques with Java Implementations*. Morgan Kaufmann, 1999.
- [25] M. Wooldridge. Intelligent agents. In G. Weiss, editor, *Multiagent Systems*. MIT Press, 1989.