



Adaptation Knowledge Discovery from a Case Base

Mathieu d'Aquin, Fadi Badra, Sandrine Lafrogne, Jean Lieber, Amedeo Napoli, Laszlo Szathmary

► To cite this version:

Mathieu d'Aquin, Fadi Badra, Sandrine Lafrogne, Jean Lieber, Amedeo Napoli, et al.. Adaptation Knowledge Discovery from a Case Base. 2006, pp.795–796. hal-00110287

HAL Id: hal-00110287

<https://inria.hal.science/hal-00110287>

Submitted on 27 Oct 2006

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Adaptation Knowledge Discovery from a Case Base

M. d'Aquin¹ and F. Badra¹ and S. Lafrogne¹ and J. Lieber¹ and A. Napoli¹ and L. Szathmary¹

Abstract. In case-based reasoning, the adaptation step depends in general on domain-dependent knowledge, which motivates studies on adaptation knowledge acquisition (AKA). CABAMAKA is an AKA system based on principles of knowledge discovery from databases. This system explores the variations within the case base to elicit adaptation knowledge. It has been successfully tested in an application of case-based decision support to breast cancer treatment.

1 INTRODUCTION

Case-based reasoning (CBR [4]) aims at solving a target problem thanks to a case base. A case represents a previously solved problem. A CBR system selects a case from the case base and then adapts the associated solution, requiring domain-dependent knowledge for adaptation. The goal of adaptation knowledge acquisition (AKA) is to extract this knowledge. The system CABAMAKA applies principles of knowledge discovery from databases (KDD) to AKA. The originality of CABAMAKA lies essentially in the approach to AKA that uses a powerful learning technique that is guided by a domain expert, according to the spirit of KDD. This paper proposes an original and working approach to AKA, based on KDD techniques.

CBR and adaptation. A case in a given CBR application is usually represented by a pair $(pb, Sol(pb))$ where pb represents a problem statement and $Sol(pb)$, a solution of pb . CBR relies on the *source cases* $(srce, Sol(srce))$ that constitute the *case base* CB . In a particular CBR session, the problem to be solved is called *target problem*, denoted by tgt . A case-based inference associates to tgt a solution $Sol(tgt)$, with respect to the case base CB and to additional knowledge bases, in particular $\mathcal{O}$, the *domain ontology* that usually introduces the concepts and terms used to represent the cases.

A classical decomposition of CBR consists in the steps of retrieval and adaptation. *Retrieval* selects $(srce, Sol(srce)) \in CB$ such that $srce$ is judged to be similar to tgt . The goal of adaptation is to solve tgt by modifying $Sol(srce)$ accordingly.

The work presented hereafter is based on the following model of adaptation, similar to *transformational analogy* [1]:

- ① $(srce, tgt) \mapsto \Delta pb$, where Δpb encodes the similarities and dissimilarities of the problems $srce$ and tgt .
- ② $(\Delta pb, AK) \mapsto \Delta sol$, where AK is the adaptation knowledge and where Δsol encodes the similarities and dissimilarities of $Sol(srce)$ and the forthcoming $Sol(tgt)$.
- ③ $(Sol(srce), \Delta sol) \mapsto Sol(tgt)$, $Sol(srce)$ is modified into $Sol(tgt)$ according to Δsol .

Adaptation is generally supposed to be domain-dependent in the sense that it relies on domain-specific adaptation knowledge. There-

fore, this knowledge has to be acquired. This is the purpose of *adaptation knowledge acquisition* (AKA).

A related work in AKA. The idea of the research presented in [3] is to exploit the variations between source cases to learn adaptation rules. These rules compute variations on solutions from variations on problems. More precisely, ordered pairs $(srce-case_1, srce-case_2)$ of similar source cases are formed. Then, for each of these pairs, the variations between the problems $srce_1$ and $srce_2$ and the solutions $Sol(srce_1)$ and $Sol(srce_2)$ are represented $(\Delta pb$ and $\Delta sol)$. Finally, the adaptation rules are learned, using as training set the set of the input-output pairs $(\Delta pb, \Delta sol)$. The experiments have shown that the CBR system using the adaptation knowledge acquired from the automatic system of AKA shows a better performance compared to the CBR system working without adaptation. This research has strongly influenced our work that is globally based on similar ideas.

2 CABAMAKA

Principles. CABAMAKA deals with case base mining for AKA. Although the main ideas underlying CABAMAKA are shared with those presented in [3], the followings are original ones. The adaptation knowledge that is mined has to be validated by experts and has to be associated with explanations that make it understandable by the user. In this way, CABAMAKA may be considered as a semi-automated (or interactive) learning system. Another difference with [3] lies in the volume of the cases that are examined: given a case base CB where $|CB| = n$, the CABAMAKA system takes into account every ordered pair $(srce-case_1, srce-case_2)$ with $srce-case_1 \neq srce-case_2$ (whereas in [3], only the pairs of *similar* source cases are considered, according to a fixed criterion). Thus, the CABAMAKA system has to cope with $n(n-1)$ pairs, a rather large number of elements, since in our application $n \simeq 750$. ($n(n-1) \simeq 5 \cdot 10^5$). This is why efficient techniques of knowledge discovery from databases (KDD [2]) have been chosen for this system.

Principles of KDD. The goal of KDD is to discover knowledge from databases, with the supervision of an analyst (expert of the domain). A KDD session usually relies on three main steps: data preparation, data-mining and interpretation.

Data preparation is based on formatting and filtering operations. The formatting operations transform the data into a form allowing the application of the chosen data-mining operations. The filtering operations are used for removing noisy data and for focusing the data-mining operation on special subsets of objects and/or attributes.

Data-mining methods are applied to extract pieces of information from the data. These pieces of information have some regular

¹ LORIA (UMR 7503 CNRS-INPL-INRIA-Nancy 2-UHP), BP 239, 54506 Vandœuvre-lès-Nancy, FRANCE

properties allowing their extraction. For example, CHARM [5] is a data-mining algorithm that performs efficiently the extraction of *frequent closed itemsets (FCIs)*. CHARM inputs a database in the form of a set of transactions, each *transaction* T being a set of boolean properties or *items*. An *itemset* I is a set of items. The support of I , $\text{support}(I)$, is the proportion of transactions T of the database possessing I ($I \subseteq T$). I is frequent, with respect to a threshold $\sigma \in [0; 1]$, whenever $\text{support}(I) \geq \sigma$. I is closed if it has no proper superset J ($I \subsetneq J$) with the same support.

Interpretation aims at interpreting the output of data-mining i.e. the FCIs in the present case, with the help of an analyst. In this way, the interpretation step produces new knowledge units (e.g. rules).

Formatting. The formatting step of CABAMAKA inputs the case base CB and outputs a set of transactions obtained from the pairs (srce-case₁, srce-case₂). It is composed of two substeps. During the first substep, each srce-case = (srce, Sol(srce)) ∈ CB is formatted in two sets of boolean properties: Φ(srce) and Φ(Sol(srce)). The computation of Φ(srce) consists in translating srce from the problem representation formalism to 2^P, P being a set of boolean properties. Possibly, some information may be lost during this translation, but this loss has to be minimized. Now, this translation formats an expression srce expressed in the framework of the domain ontology O to an expression Φ(srce) that will be manipulated as data, i.e. without the use of a reasoning process. Therefore, in order to minimize the translation loss, it is assumed that if $p \in \Phi(\text{srce})$ and p entails q (given O) then $q \in \Phi(\text{srce})$. In other words, Φ(srce) is assumed to be deductively closed given O in the set P. The same assumption is made for Φ(Sol(srce)). How this first substep of formatting is computed in practice depends heavily on the representation formalism of the cases.

The second substep of formatting produces a transaction $T = \Phi((\text{srce-case}_1, \text{srce-case}_2))$ for each ordered pair of distinct source cases, based on the sets of items Φ(srce₁), Φ(srce₂), Φ(Sol(srce₁)) and Φ(Sol(srce₂)). Following the model of adaptation presented in introduction (items ①, ② and ③), T has to encode the properties of Δpb and Δsol. Δpb encodes the similarities and dissimilarities of srce₁ and srce₂, i.e.:

- The properties common to srce₁ and srce₂ (marked by “=”),
- The properties of srce₁ that srce₂ does not share (“-”) and
- The properties of srce₂ that srce₁ does not share (“+”).

All these properties are related to problems and thus are marked by pb. Δsol is computed in a similar way and $\Phi(T) = \Delta\text{pb} \cup \Delta\text{sol}$. For example,

$$\begin{aligned} \text{if } & \begin{cases} \Phi(\text{srce}_1) = \{a, b, c\} & \Phi(\text{Sol}(\text{srce}_1)) = \{A, B\} \\ \Phi(\text{srce}_2) = \{b, c, d\} & \Phi(\text{Sol}(\text{srce}_2)) = \{B, C\} \end{cases} \\ \text{then } & T = \{a_{\text{pb}}^-, b_{\text{pb}}^+, c_{\text{pb}}^-, d_{\text{pb}}^+, A_{\text{sol}}^-, B_{\text{sol}}^+, C_{\text{sol}}^+\} \end{aligned} \quad (1)$$

Mining. The extraction of FCIs is computed thanks to CHARM (in fact, thanks to a tool based on a CHARM-like algorithm) from the set of transactions. A transaction $T = \Phi((\text{srce-case}_1, \text{srce-case}_2))$ encodes a specific adaptation ((srce₁, Sol(srce₁)), srce₂) ↦ Sol(srce₂). An FCI extracted may be considered as a generalization of a set of transactions. For example, if $I_{\text{ex}} = \{a_{\text{pb}}^-, c_{\text{pb}}^+, d_{\text{pb}}^+, A_{\text{sol}}^-, B_{\text{sol}}^+, C_{\text{sol}}^+\}$ is an FCI, I_{ex} is a generalization of a subset of the transactions including the transaction T of equation (1): $I_{\text{ex}} \subseteq T$. The interpretation of this FCI as an adaptation rule is explained below.

Interpretation. The interpretation step is supervised by the analyst. The CABAMAKA system provides the analyst with the extracted FCIs and facilities for navigating among them. The analyst may select an FCI, say I , and interpret I as an adaptation rule. For example, the FCI I_{ex} may be interpreted in the following terms:

if a is a property of srce but is not a property of tgt,
 c is a property of both srce and tgt,
 d is not a property of srce but is a property of tgt,
 A and B are properties of Sol(srce) and
 C is not a property of Sol(srce)
then the properties of Sol(tgt) are
 $\Phi(\text{Sol}(\text{tgt})) = (\Phi(\text{Sol}(\text{srce})) \setminus \{A\}) \cup \{C\}$.

This has to be translated as an adaptation rule r of the CBR system. Then the analyst corrects r and associates an explanation with it.

Implementation. The application domain of the CBR system we are developing is breast cancer treatment: in this application, a problem pb describes a class of patients with a set of attributes and associated constraints (holding on the age of the patient, the size and the localization of the tumor, etc.). A solution Sol(pb) of pb is a set of therapeutic decisions (in surgery, chemotherapy, etc.). The requested behavior of the CBR system is to provide a treatment and explanations on this treatment proposal. This is why the analyst is required to associate an explanation to a discovered adaptation rule.

The problems, solutions and the domain ontology of the application are represented in OWL DL (recommendation of the W3C).

3 CONCLUSION

The CABAMAKA system presented in this paper is inspired by the research presented in [3] and by the principles of KDD for the purpose of semi-automatic adaptation knowledge discovery. It has enabled to discover several useful adaptation rules for a medical CBR application. It has been designed to be reusable for other CBR applications: only a few modules of CABAMAKA are dependent on the formalism of the cases and of the domain ontology, and this formalism, OWL DL, is a well-known standard. One element of future work consists in searching for ways of simplifying the presentation of the numerous extracted FCIs to the analyst. This involves an organization of these FCIs for the purpose of navigation among them. Such an organization can be a hierarchy of FCIs according to their specificities or a clustering of the FCIs in themes.

REFERENCES

- [1] J. G. Carbonell, ‘Learning by analogy: Formulating and generalizing plans from past experience’, in *Machine Learning, An Artificial Intelligence Approach*, ed., R. S. Michalski and J. G. Carbonell and T. M. Mitchell, chapter 5, 137–161, Morgan Kaufmann, Inc., (1983).
- [2] M.H. Dunham, *Data Mining – Introductory and Advanced Topics*, Prentice Hall, Upper Saddle River, NJ, 2003.
- [3] K. Hanney and M. T. Keane, ‘Learning Adaptation Rules From a Case-Base’, in *Advances in Case-Based Reasoning – Third European Workshop, EWCBR’96*, eds., I. Smith and B. Faltings, LNAI 1168, pp. 179–192. Springer Verlag, Berlin, (1996).
- [4] C. K. Riesbeck and R. C. Schank, *Inside Case-Based Reasoning*, Lawrence Erlbaum Associates, Inc., Hillsdale, New Jersey, 1989.
- [5] M. J. Zaki and C.-J. Hsiao, ‘CHARM: An Efficient Algorithm for Closed Itemset Mining’, in *SIAM International Conference on Data Mining SDM’02*, pp. 33–43, (Apr 2002).