

Individual prognosis of diabetes long-term risks: A CBR Approach

E. Armengol¹, A. Palaudàries² and E. Plaza^{1*}

¹Artificial Intelligence Research Institute (IIIA-CSIC)
Campus UAB, 08193- Bellaterra, Catalonia, Spain

² Unitat d'Endocrinologia. Hospital de Mataró
Carretera de Cirera, s/n. 08304-Mataró, Catalonia, Spain

Abstract

In this paper we present DIRAS, an application supporting the physicians to determine the risk of complications for individual diabetic patients. The risk pattern of each diabetic patient is obtained using a Case-based Reasoning method called LID. Case-based Reasoning is an Artificial Intelligence technique based on solving new situations according to past experiences. For each patient, the LID method determines the risk of each diabetic complication according to the risk of already diagnosed patients. In addition, LID builds a description that can be viewed as an explanation of the obtained risk.

1. Introduction

Lucas and Abu-Hanna [1] define *prognosis* as the prediction of the course and outcome of disease processes. Usually, computer systems supporting the physicians to take decisions use *prognosis models* that predict the outcomes of a particular process. Techniques used to build prognosis models range from a construction by hand to statistical techniques or Artificial Intelligence (AI) techniques –see [1] for the description of specific techniques. In particular, we have been investigating *Case-based Reasoning* [2] an AI technique based in the human capability to solve new situations by learning from the past situations already solved.

In this paper we present DIRAS (Diabetes Individualized Risk Assessment System), an application whose goal is to predict the risk of complications for diabetic patients. Diabetes Mellitus

*Corresponding author. Phone: 34-935809570, Fax: 34 - 935809661, e-mail: enric@iiia.csic.es

is one of the most frequent human clinic diseases since it affects around a 3% of the European population and around one hundred million people in the world. There are two major types of diabetes: type 1 (or insulin-dependent) and type 2 (or non insulin-dependent). The diabetes type 1 usually develops in children or people less than 40 years old. This form of diabetes is characterized by an insufficient production of insulin at the pancreas. People with this type of diabetes need daily injections of insulin. If not diagnosed and treated with insulin, the person can lapse into a life-threatening coma. Diabetes type 2, the most common one, usually develops in adults over the age of 40 being more common among adults over 55. Usually, people with diabetes type 2 have overweight and sedentary lifestyle. In diabetes type 2 the pancreas produces insulin but the body does not use it effectively. The consequences are the same that those of the diabetes type 1 although its symptoms appear gradually, and they tend to be vague. Some people with diabetes type 2 must inject insulin, but most are controlled with a combination of weight loss, exercise, and prescription of oral diabetes medication.

A bad management of both forms of diabetes will produce microcomplications (such as blindness, renal failure or polyneuropathy), and macrocomplications (such as gangrene and amputation, aggravated coronary heart disease or stroke). Therefore, main concern in the management of the diabetes is reducing the risk of a patient developing a new long-term complication and the risk of progression in the complications already present. The prediction of the individual risk to develop long-term complications is based on the analysis of a large quantity of data (e.g. diabetes type, diabetes duration, age, cholesterol, and metabolic control degree) that have to be continuously evaluated. The therapeutic goals to offer a good life quality to the patient depend on this analysis. Because diabetes mellitus has a high prevalence, sometimes physicians taking care of diabetic patients have not a specialized formation in diabetes and, consequently, the management of these patients may be less accurate. For this reason, a system supporting an individualized assessment of the patients can be a useful tool to improve both the management and the treatment of diabetes.

We have developed the DIRAS application supporting the physicians to determine the risk of complications for each patient according to the clinical data of that patient. We call *risk pattern* to the set of assessments concerning diabetic complications. DIRAS obtains a risk pattern where the risk of diabetic macrocomplications (ischemic cardiopathy, low extremities vasculopathy, and stroke) and diabetic microcomplications (nephropathy, retinopathy and polyneuropathy) are explicitly assessed. The DiabCare Quality Network (<http://www.diabcare.de>) is a European consortium having as goal the improvement of the care in diabetes type 1 and type 2. Basically, the goal of this project is to implement effective measures for the prevention of complications such as blindness due to diabetes, number of people entering to an end-stage diabetic renal failure, etc. DiabCare manages groups of diabetic patients using statistical tools. The main contribution of DIRAS is focusing on individual patients instead of populations of patients. For this purpose, we

use an Artificial Intelligence technique called *Case-based Reasoning* (2) for assessing the risk of complications in patients with diabetes type 1 or type 2.

The structure of this paper is the following. Section 2 explains the goal and the knowledge structures used by the DIRAS application. Section 3 presents LID, a Case-based Reasoning method used by DIRAS to assess the complication risks of a patient. Finally, in section 4 the results of DIRAS are discussed.

2. The DIRAS Application

The goal of DIRAS is to determine the risk of complications for individual diabetic patients, what we will call the *risk pattern*. For each patient, DIRAS works with five kinds of data (Figure 1): Personal-Data, Basic-Diabetes-Data, Info-Patient-Consultation, Assessment and Risk-Pattern. Personal-data contains information such as the name, address, birth date, etc. Basic-Diabetes-Data contains basic information of diabetes (such as diabetes type, duration, and whether diabetes is treated with oral drugs or insulin). Info-Patient-Consultation has data on relevant measures (e.g. glycated hemoglobin, cholesterol, blood pressure, etc), eye and foot examination, current treatments, etc.

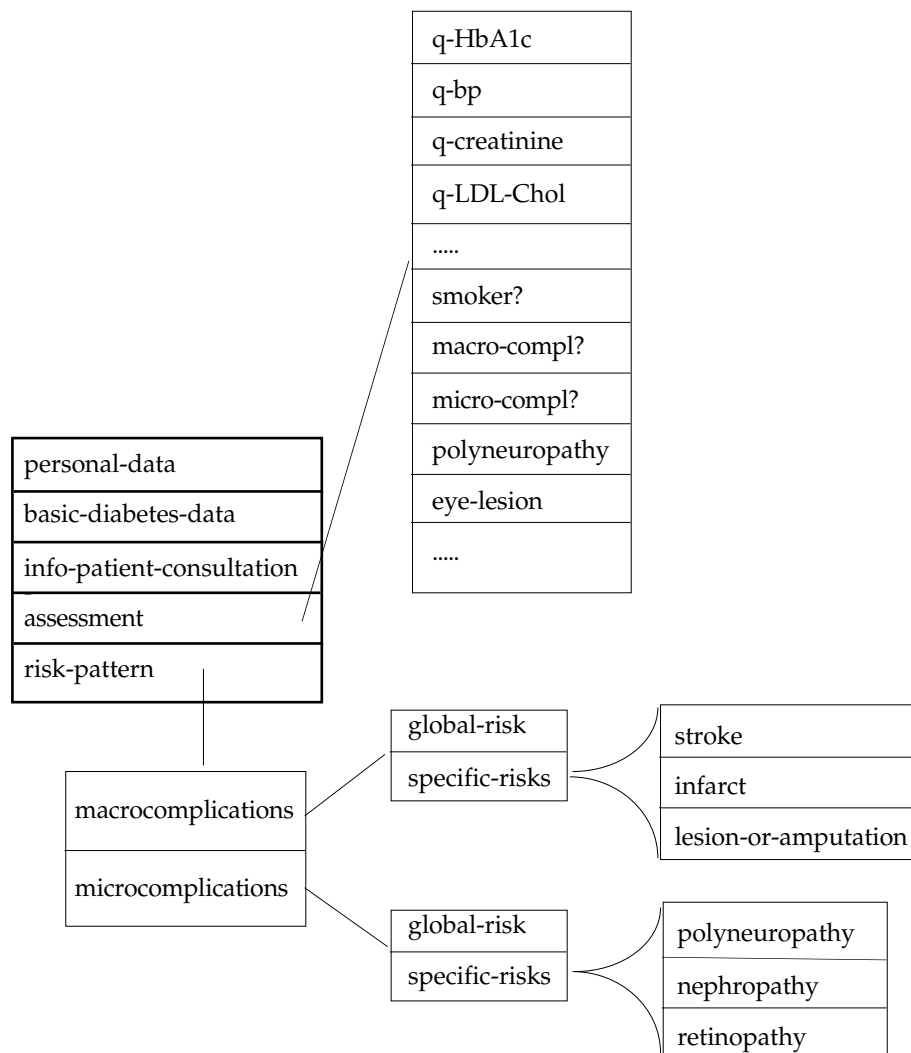


Figure 1. Description of a patient

Assessment contains knowledge obtained by applying domain knowledge provided by an expert diabetologist. This domain knowledge allows the analysis of the patient's data obtaining a high level perspective of the patient's state. For instance, Assessment holds a qualitative measure of the LDL-cholesterol (q-LDL-chol of Assessment in Fig. 1). This qualitative measure may have a value *low*, *moderate* or *high* depending on the following conditions:

```
if the patient has macrocomplications then
    if LDL-chol > 130 then q-LDL-chol = high
    else if LDL-chol > 100 then q-LDL-chol = moderate
    else low
if the patient has no macrocomplications then
    if LDL-chol > 150 then q-LDL-chol = high
    else if LDL-chol > 130 then q-LDL-chol = moderate
    else low
```

Notice that the qualitative measure depends on the presence or absence of macrocomplications. DIRAS has similar rules for qualifying other measures in an appropriate way.

In addition, from some measures and symptoms, DIRAS can infer for Assessment new facts like whether or not the patient has a specific complication. For instance the feature *micro-comp?* of Assessment contains information about the presence or absence of microcomplications in the current patient. This feature has value *true* if the patient has eye lesions, nephropathy or polyneuropathy; otherwise the feature *micro-comp?* has value *false*. In turn, the presence or absence of eye lesions, nephropathy or polyneuropathy is inferred by DIRAS using similar rules that take into account relevant data held in Info-Patient-Consultation.

Finally, the last kind of patient data is Risk-Pattern i.e. the assessment of individual long-term risks that we want to estimate using Case-based Reasoning. Risk-Pattern has two parts (Figure 1): 1) the macrocomplication risks, and 2) the microcomplication risks. For macrocomplications we want to assess both the global risk and the risk of three particular macrocomplications (namely stroke, infarct and amputation). Similarly, for microcomplications we want to assess both the global risk and the risk of polyneuropathy, nephropathy and retinopathy. The global risk represents the risk of a patient to vascular alterations that have not defined symptoms such as the intestine infarct.

There are two kinds of risk for complications: *development* risk and *progression* risk. The development risk has to do with patient's likelihood of developing a new complication in the future. The progression risk is when a patient already has a macrocomplication and thus the risk of further deterioration has to be assessed.

The next section shows in detail how DIRAS uses Case-based Reasoning to obtain a risk pattern for individual patients.

3. CBR Assessment Risk Pattern

The goal of DIRAS is to obtain an individual risk pattern for diabetic patients using *Case-based Reasoning* [2]. Case-based Reasoning (CBR) is an AI technique based on the human capability to solve new situations according to past experience. The core idea of CBR is that when a new situation is similar to one or several old situations, the decisions taken and the knowledge contained in old situations provide a starting point to interpret or solve the new situation. Each situation is called a *case* (or precedent) that may be reused to solve new problems. The collection of cases of a system is called the *case base*. The description of a new situation to be solved is called the *current case* or problem.

Given a case base and a problem, CBR methods perform three tasks [3]: 1) *retrieve*, that obtains past cases similar to the new case; 2) *select*, that decides which of the retrieved past cases is the most similar (i.e. the best precedent) to the current problem; and 3) *adapt*, that decides how to adapt the solution of the best precedent to solve the current problem.

DIRAS uses a case base where each case is a patient described as explained in the previous section, i.e. the patient's data plus the solution (the risk pattern) for that patient. The goal of DIRAS is to obtain the risk pattern for the current patient. Several features (see Figure 1) form the risk pattern and DIRAS obtains the risk for each feature in an independent way.

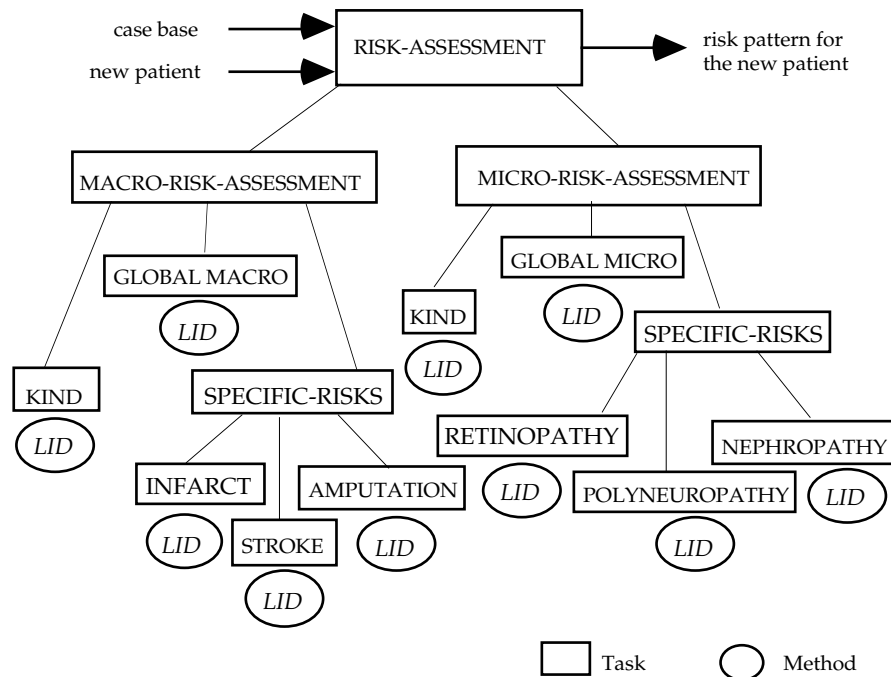


Figure 2. Task decomposition of the risk-assessment task that obtains the risk pattern of diabetic patients.

The complete risk pattern is obtained by solving the Risk-Assessment task that decomposes in two tasks (Fig. 2): the Macro-Risk-Assessment task and the Micro-Risk-Assessment task. The Macro-Risk-Assessment task decomposes, in turn, in three subtasks: the Kind task determines whether the risk of macrocomplications is progression or development; the Global-Macro task assesses the global risk of macrocomplications; and the Specific-Risks task assesses the risk of three specific macrocomplications, namely infarct, stroke and amputation. Similarly, the Micro-Risk-Assessment task is decomposed in three subtasks: Kind, Global-Micro and Specific-Risks. In particular the Specific-Risks task assesses the risk of three specific microcomplications: retinopathy, polyneuropathy and nephropathy. All these risks are inferred using LID (*Lazy Induction of Descriptions*), a Case-based Reasoning method that is explained in the following section.

3.1. The LID Method

In this section we introduce LID, the CBR method used by DIRAS to solve the risk assessment tasks described in Fig. 2. For each diabetic task LID searches the case base for the best precedent and infers the risk according to that precedent.

For a given collection of risk classes $R = \{\text{unknown}, \text{low}, \text{moderate}, \text{high}, \text{very-high}\}$, a diabetic complication C , and a problem p , the task of LID is to obtain the risk $R_i \in R$ of p concerning C . For each complication C , this can be seen as a classification task where the goal is to identify the class in R to which p belongs. DIRAS solves this classification task using LID.

Given a case base B containing diabetic patients classified into the collection of risk classes R for a diabetic complication C , and a problem p to be classified, LID obtains the class $R_i \in R$ to which p belongs. Intuitively, LID follows a top-down strategy to build a description D containing the most relevant features of p such that all features in D are satisfied by a subset of cases in B . In general, cases in this subset belong to different solution classes in R . LID adds relevant features to D until the subset of cases satisfying D belong to one unique solution class R_i . LID takes this class R_i as the solution for the current task, i.e. R_i is the risk of p concerning C .

```

D := ∅; R = {R1 ... Rn}
Function LID (B, p, D)
  SD := Discriminatory-set (D, B)
  if ∃ ei ∈ SD ⇒ ei ∈ Ri then return Ri
    else fd := Select-Feature (p, B, R)
      D' := Add-Feature (fd, D)
      LID (SD, p, D')
  end-if
end-function

```

Figure 3. The LID algorithm

The LID algorithm (Fig. 3) begins with the whole set of precedents B classified into the collection of risk classes R for a complication C , a problem p to be solved and the description $D = \emptyset$ (i.e. D has no features). In the following we will explain this algorithm using an example. See [4] for more detailed explanation and some results of the LID method on other domains of application.

Example 1. Let p be a patient with no macrocomplications (i.e. feature *macro-compl?* in Assessment has value *false*), high blood pressure and low albumin. In this example DIRAS has to determine the risk $R_i \in R$ for the macrocomplication $C = \text{stroke}$.

The set of cases $S_D \subseteq B$ that are subsumed by the description D is called *discriminatory set*. Intuitively, a case c is subsumed by a description D when all the information contained in D is also contained in c , although c can contain more information. See a formal description of subsumption in [5].

Initially D is an empty description, i.e. it is the most general description. Therefore D subsumes all the cases in B (i.e. $S_D = B$), and consequently D has to be specialized. The specialization of a description D is achieved by adding features to it. In particular, LID adds a feature f with the value v that this feature has in the current problem p . After that, the new description $D' = D + (f=v)$ has a smaller discriminatory set $S_{D'}$ formed by those cases subsumed by D' . Thus, specialization reduces the discriminatory set $S_{D'} \subseteq S_D$ at each step. As we explain below, LID uses a heuristic measure based on the López de Mántaras distance [6] to determine the feature to be added.

LID specializes D by selecting one feature f from all the features used in p in the following way. Each feature f_i in p induces a partition $P_i = \{S_{i1} \dots S_{in}\}$ in the set S_D such that each $S_{ik} \in P_i$ contains those patients in S_D having the same value v_k in the feature f_i . For instance, the presence or absence of macrocomplications will divide the set S_D (currently $S_D = B$) in two subsets: one containing those precedents having macrocomplications and the other one containing patients without macrocomplications. There is also a partition of S_D , called the *correct partition* P_c , that divides S_D according to the risk ($R_i \in R$) for the complication C . In the example, S_D is divided in subsets according to the values for the risk of stroke being *unknown*, *low*, *moderate*, *high*, and *very-high*.

For each partition P_i , LID computes the López de Mántaras (RLM) distance [6] to the correct partition P_c . Intuitively, the RLM distance assesses how similar a partition is with respect to a referent partition (i.e. the correct partition), in the sense that the lesser the distance the more similar they are. The RLM distance was introduced as an alternative to the Quinlan's Gain [7] used in the ID3 inductive learning algorithm. The Quinlan's Gain is a selection measure that selects the object feature providing the highest information gain. RLM distance shows that normalizing the Quinlan's Gain in an appropriate way, we obtain a distance between partitions.

Formally, given two partitions P_i and P_c of a set S_D , the RLM distance between them is computed as follows:

$$\mathbf{RLM}(P_i, P_c) = 2 - \frac{I(P_i) + I(P_c)}{I(P_i \cap P_c)} \quad \text{where } I(P_i) = -\sum_{j=1}^n p_j \log_2 p_j, p_j = \frac{|S_D \cap S_{ij}|}{|S_D|}$$

$$I(P_c) = -\sum_{k=1}^m p_k \log_2 p_k, p_k = \frac{|S_D \cap R_k|}{|S_D|}$$

$$I(P_i \cap P_c) = -\sum_{j=1}^n \sum_{k=1}^m p_{jk} \log_2 p_{jk}$$

$$p_{jk} = \frac{|S_D \cap R_k \cap S_{ij}|}{|S_D|}$$

where $I(P_i)$ measures the information contained in the partition P_i ; n is the number of possible values of the feature inducing P_i ; $m = \text{Card}(R)$; p_i is the probability of occurrence of class S_{ij} (R_j) i.e. the proportion of examples in S_D that belong to S_{ij} (R_j); $I(P_i \cap P_c)$ is the mutual information of two partitions; and p_{jk} is the probability of occurrence of the intersection $R_j \cap S_{ik}$, i.e. the proportion of examples in S_D that belong to R_j and to S_{ik} .

Using the RLM distance, we can define what it means for a feature to be more discriminatory than another.

Definition ["More discriminatory than" relation]. Let P_c be the correct partition (i.e. the partition that correctly classifies the examples), and P_j and P_k the partitions induced by features f_j and f_k respectively, we say that feature f_j is *more discriminatory than* feature f_k iff $\mathbf{RLM}(P_j, P_c) < \mathbf{RLM}(P_k, P_c)$

In other words, when a feature f_j is more discriminatory than another feature f_k the partition that f_j induces in S_D is closer to the correct partition P_c than the partition induced by f_k . Intuitively, the most discriminatory feature classifies the cases in S_D in a more similar way to the correct classification of cases (i.e. that determined by the risk of the complication C).

Thus, LID selects the feature f having minimum RLM distance to the correct partition as being the most discriminatory. Then, LID builds D' , a specialization of D , by adding to D the feature f with the same value that f takes in the current problem p . In the example, the most discriminatory feature is macro-comp? Thus the description D' will contain the feature macro-comp? with value *false* (since the current problem p corresponds to a patient that has no macrocomplications).

Let $S_{D'}$ be the subset of precedents that subsumed by D' . If all the precedents in $S_{D'}$ belong to only one risk class R_i then LID finishes the process and classifies p as belonging to R_i . Otherwise, D' needs to be further specialized. In Example 1, $S_{D'}$ contains those cases in $S_D = B$ with the feature macro-comp? being *false*. However, with respect to the correct partition P_c (induced from the risk of stroke) the cases in $S_{D'}$ belong to several of these solution classes in R . Because of this, D' need be specialized in order to reduce the discriminatory set $S_{D'}$. This specialization is made using LID with the set $S_{D'}$, the description D' and the patient p . Notice that LID may safely ignore the precedents

that do not belong to $S_{D'}$ because the precedents that are not subsumed by D' will not be subsumed by any specialization of D' .

The next step in the specialization of D' is the selection of the next most discriminatory feature. Now LID finds that the most discriminatory feature is the qualitative measure of the blood pressure (feature q-bp of Assessment in Figure 1). The patient p has value *high* in the q-bp feature, therefore the specialization D'' of D' contains two features: macro-comp? with value *false* (as before) and the newly added feature q-bp with value *high*.

Subsequently, LID considers the discriminatory set $S_{D''}$ containing those precedents in $S_{D'}$ subsumed by D'' and finds whether or not all cases in $S_{D''}$ belong only to one class R_i . In Example 1, all the cases contained in $S_{D''}$ belong to the class $R_i = \text{high}$ for the risk of stroke. Therefore, LID finishes inferring that the patient p has also a high risk of stroke.

There is an abnormal stopping condition produced when there is no possible to specialize the current description. This situation occurs when LID has used all the features candidates to specialize a description but the current description D^n subsumes precedents belonging to several classes $R' \subseteq R$. In that situation, LID proposes as solution for the current problem the classes in R' .

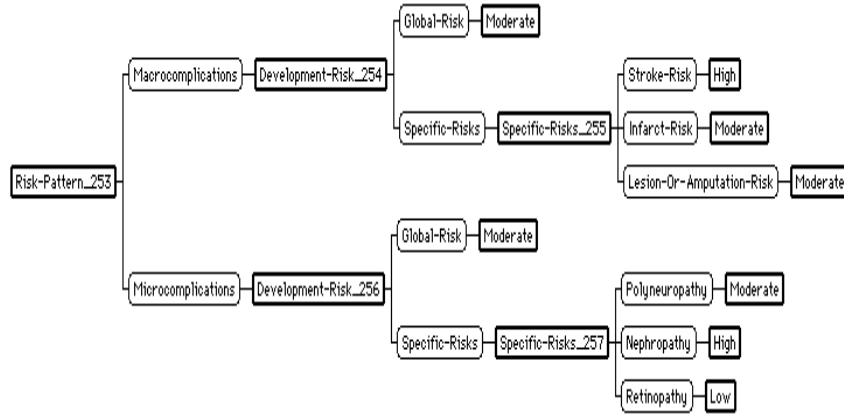


Figure 4. Browsing of a Risk-Pattern obtained by DIRAS in the example 1.

DIRAS uses LID to solve the tasks shown in Figure . That is to say, for macrocomplications LID is used to solve the tasks Kind, Global-Macro, Infarct, Stroke, and Amputation; and for microcomplications LID is used to solve the tasks Kind, Global-Micro, Retinopathy, Polyneuropathy and Nephropathy. Therefore, for each diabetic patient, DIRAS obtains a risk pattern as the one in Fig. 4 for example 1.

Another concern of LID is the interpretation of the descriptions as an explanation for the assessment of a specific risk. For instance, Fig. 5 shows the description obtained by LID for Example 1. This description considers that macro-comp? and q-bp are relevant to assess the risk of stroke. In particular, the expert diabetologist agrees with this explanation since it is consistent with his knowledge. For other examples, LID explains the moderate risk of stroke of a patient with no macrocomplications because of the moderate blood pressure. Nevertheless, when the patient has mac-

rocomplications and his blood pressure is moderate, LID assesses a high risk of stroke. The expert agrees with both explanations, since it is known that the risk of stroke directly depends on the level of the blood pressure and it is increased when the patient has macrocomplications. There also are examples to which DIRAS has assessed a high global progression risk of macrocomplications because the patient has high levels of both LDL-cholesterol and HDL-cholesterol. Moreover, DIRAS has obtained the same explanation to justify the high global progression risk of other cases. The expert agrees this explanation for these specific cases.

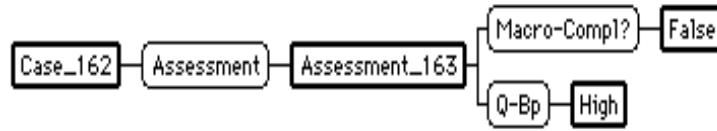


Figure 5. Explanation provided by LID of why a case has a high risk of stroke.

3.2. Discussion

Currently, the DIRAS system is fully implemented with a case base of 370 patient records. We are performing several qualitative and quantitative measures to estimate which parts of the system are stable and which parts need further refinement. In particular, we are interested in determining whether the case base of 370 patient compiled from just one hospital, is a good enough sample of the population or should be increased. We'll explain first the tests we have performed to estimate the current status and then we'll explain the plans the final validation of DIRAS.

The qualitative verification is performed by an expert diabetologist that focuses on two concerns: the risk assessments for a patient, and the explanation of each assessment. The goal of this process is to determine the stable parts of the system, the quality of the case base, and the kinds of expansion of the case base that could be needed. In order to perform a more quantitative estimate of the system performance our expert diabetologist constructed a *gold standard* consisting of a risk pattern for macrocomplications for all 370 cases. This gold standard is somewhat artificial, in that it gives a unique "correct" risk value for each complication and considers all other risk estimations "incorrect".

For the quantitative estimation we have built 15 test sets with our 370 patients case base, where each test set has 319 cases randomly chosen as training set. The results of DIRAS upon the remaining 70 cases for each test set where compared with the gold standard and averaged for each macrocomplication. The results of DIRAS are 100% correct in determining the kind of risk (progression or development) and the risk of stroke, 90% correct in determining amputation risk, and 72.45% correct in determining the global risk and the risk of infarct. These two last risks are in practice the same, which is the reason they have the same accuracy. In fact, the incorrect risk assessments here fail only by one degree (e.g. high risk vs. very-high risk) in an 81.69% of the cases.

The estimation of pure accuracy allows us to distinguish the parts of DIRAS that require revisions (like the one related to infarct, that probably requires an enlargement of the case base) but it's not a good way to validate a medical recommendation system. The final validation will follow the process and use the criteria developed at our Institute for validating the PneumonIA expert system [8]. This process involves several expert diabetologists to which the data of the patients in the case base are shown. Each expert independently assesses the risk pattern of each patient, and in this way we find out areas of consensus and areas where risk assessment may have some admissible variation. Finally, the goal of this validation process is to establish whether the performance of DIRAS is *indistinguishable* or not from that of the human experts. Specifically, we will perform a ranking of the set of composed of the experts plus DIRAS. If the system is ranked among the experts this means that its performance is indistinguishable from that of the best experts, while if it ranks below them it is distinguishable [8].

3.3. Report of a Diabetic Patient

In addition to the risk pattern of an individual patient, DIRAS produces a report that can be useful for nonexpert physicians to manage diabetic patients. This report (Fig. 6) is formed by four sections: 1. personal data of the patient (e.g. age, diabetes type, year of the diabetes diagnostic), 2. assessments about the measures, 3. Information about macrocomplications, and 4. Information about microcomplications. Section 2 of the report contains, for each measure M , the patient's value for M , the range of normal values for M , an assessment on whether the patient's value for M is acceptable or not. Section 3 of the report is divided in two parts. The first part contains the report of the patient's macrocomplications detailing some aspects of the patient's state (for instance the foot state). The second part shows the risk of each macrocomplication and the factors used to determine this risk. The section 4 of the report has the same structure that the section 3 but has to do with the patient's state and risks concerning microcomplications.

1- PATIENT DATA

Number: 3408301200246854

Man, 76 years old, with a diabetes type 2 diagnosed 25 years ago

2- GENERAL ASSESSMENT

	Value	Good Value	Assessment
HbA1c	9.10	<= 6.5	unacceptable
BP (SYS)	190.00	<= 130	unacceptable
BP (DIA)	90.00	<= 80	acceptable
Cholesterol	166.00	<= 180	correct
LDL Chol	—	<= 100 (macro)/ <= 130	—
HDL Chol	—	=> 45	—
Triglycerids	94.00	<= 150	correct
Albumin	289.00	<= 20	unacceptable
Creatinine	101.00	<= 106	correct

Good stability because he/she had no hypoglycemies nor hyperglycemies. No hospitalizations.

3- MACROVASCULAR COMPLICATIONS

Ischemic heart : No

Infarct (or coronary bypass or angioplasty): No

Anginal chest pain : No

Stroke : No

Low extremities vasculopathy : Yes

Amputation above ankle : No

Amputation below ankle : No

Leg claudication : No

Right foot

Left foot

Bypass or angioplasty :

No

No

Feet pulse present :

No

No

Healed ulcer :

No

No

Acute ulcer :

No

No

Global risk progression: HIGH because the value of total cholesterol is high

Risk for Specific Macrocomplications :

Stroke : VERY-HIGH because the blood pressure is very high and the patient has macrocomplications

Infarct : HIGH because the value of total cholesterol is high

Lesion/amp. : HIGH, because the patient has polyneuropathy with normal sensitivity and vasculopathy

4- MICROVASCULAR COMPLICATIONS

Polyneuropathy : Yes

Neuropathy symptoms : true

Right foot

Left foot

Pulse present :

No

No

Pin prich sensitivity :

Abnormal

Abnormal

Vibration sensitivity :

Abnormal

Abnormal

Nephropathy : Yes

Phase I: heavy MAU (MAU : 289)

Renal insuficiency : No (creatinine : 101)

Retinopathy : Yes

Right Eye

Left Eye

Retinopathy type :

PREPROLIFERATIVE

PREPROLIFERATIVE

Visual Acuity :

Unknown

Unknown

Risk for specific microcomplications :

Polyneuropathy : HIGH progression risk because the HbA1c is high

Nephropathy : VERY HIGH progression risk because the albumin is correct, the HbA1c is high, the blood pressure is not low and the patient does not follow nephropathy treatment

Retinopathy : HIGH progression risk because the HbA1c is high

Visual Ac. Dec. : MODERATE progression risk because the maculopathy has been photocoagulated

Figure 6. Report of complications for a diabetic patient.

4. Related Work and Conclusions

There are two aspects of DIRAS that can be compared with other works: the application domain and the methodology. Concerning to the domain application, there are several applications used in the management of diabetic patients [9, 10]. These applications are oriented to determine the insulin dosage for a diabetic patient of type 1. Basically the goal is to determine a management plan for each patient according to both his particular lifestyle (i.e. diet and physical exercise) and his metabolic state (i.e. glucose levels). Instead, the goal of DIRAS is to assess the risk of long term complications for individual diabetic patients (either type 1 or type 2 diabetes).

Concerning the methodology, DIRAS uses LID, a Case-based Reasoning method that builds a discriminant explanation of the assessed risk of complication using an heuristic based on the RLM distance. This heuristic has been used in induction of decision trees [7]. The only similitude here is in the use of a heuristic for selecting an attribute as more discriminant than others, but the structure that is built in decision trees and in LID are different. Decision trees build a structure that classifies a training set of examples and uses the heuristic to select the branching criteria of that tree structure. The case-based method LID is a problem-centered technique that builds a structure that discriminates the new problem with respect to classes in the training set and the heuristic is used to decide the attribute that is more useful in discriminating this new example with respect to the training examples.

One of the advantages of Case-based Reasoning is that the solution to a problem provided by a CBR system can be justified showing the user the precedent case(s) used to support such a decision. This form of justification is supported by DIRAS by showing for each risk assessment the cases in the discriminatory set. Moreover, DIRAS constructs a symbolic explanation with the features that are relevant in classifying a patient complication in a risk class. This symbolic explanation is close to the justification that can be provided by an expert for the same problem and may allow the user to focus on the critical features for a particular patient.

Acknowledgment

The authors thank Ana M. Monteiro for her collaboration in the acquisition of the diabetes domain knowledge. This work has been developed in the context of the SMASH project supported by the Spanish Project CICYT TIC96-1038-C04-01

References

1. Lucas PJF, Abu-Hanna A. Prognostic methods in medicine. *Artificial Intelligence in Medicine* 1999; 15: 105-109.
2. Kolodner JL. *Case-based Reasoning*. Morgan Kauffman, 1993

3. Armengol E, Plaza E. A knowledge level model of Case-based Reasoning. In Topics in Case-based Reasoning. Lecture Notes in Artificial Intelligence, 837. S. Wess K.D Althoff and M. M. Richter (Eds). 1994: 53-64.
4. Armengol, E. A framework for integrated learning and problem solving. Monografies de l'IIIA. Vol 5. Institut d'Investigació en Intel·ligència Artificial Ed. Barcelona. 1997. <http://www.iiia.csic.es/Publications/monographs.html/>
5. Armengol E, Plaza E. Bottom-up induction of feature terms. Machine Learning Journal. To be published.
6. López de Mántaras R. A Distance-based Attribute Selection Measure for Decision Tree Induction. Machine Learning 1991; 6: 81-92.
7. Quinlan J.R. Induction of decision tress. Machine Learning 1986;1:81-106.
8. Verdaguer A. Patak A. Sancho J.J. Sierra C. and Sanz F. Validation of the Medical Expert System PNEUMON-IA. J.H. van Bemmell and A.T. McCray (Eds.) Yearbook of Medical Informatics 1993. IMIA's Publications, 1993: 446-461
9. Lehmann E.E, Deutsch T, Carson E.R, and Sönksen P.H. Combining rule-based resoning and mathematical modelling in diabetes care. Artificial Intelligence in Medicine 1994; 6: 137-160.
10. Larizza C, Bellazzi R, Riva A. Temporal abstractions for diabetic patients management. Proceedings of AIME 97,E. Keravnou, C. Garbay, R. Baud and J. Wyatt eds. 1997: 319-330.