

The CN2 Induction Algorithm

PETER CLARK

(PETE@TURING.AC.UK)

TIM NIBLETT

(TIM@TURING.AC.UK)

The Turing Institute, 36 North Hanover Street, Glasgow, G1 2AD, U.K.

(Received: June 25, 1987)

(Revised: October 25, 1988)

Keywords: Concept learning, rule induction, noise, comprehensibility

Abstract. Systems for inducing concept descriptions from examples are valuable tools for assisting in the task of knowledge acquisition for expert systems. This paper presents a description and empirical evaluation of a new induction system, CN2, designed for the efficient induction of simple, comprehensible production rules in domains where problems of poor description language and/or noise may be present. Implementations of the CN2, ID3, and AQ algorithms are compared on three medical classification tasks.

1. Introduction

In the task of constructing expert systems, methods for inducing concept descriptions from examples have proved useful in easing the bottleneck of knowledge acquisition (Mowforth, 1986). Two families of systems, based on the ID3 (Quinlan, 1983) and AQ (Michalski, 1969) algorithms, have been especially successful. These basic algorithms assume no noise in the domain, searching for a concept description that classifies training data perfectly. However, application to real-world domains requires methods for handling noisy data. In particular, one needs mechanisms that do not *overfit* the induced concept description to the data, and this requires relaxing the constraint that the induced description must classify the training data perfectly.

Fortunately, the ID3 algorithm lends itself to such modification by the nature of its general-to-specific search. Tree-pruning techniques (e.g., Quinlan, 1987a; Niblett, 1987), used for example in C4 (Quinlan, Compton, Horn, & Lazarus, 1987) and ASSISTANT (Kononenko, Bratko, & Roskar, 1984), have proved effective against overfitting. However, the AQ algorithm's dependence on specific training examples during search makes it less easy to modify. Existing implementations, such as AQ11 (Michalski & Larson, 1983) and AQ15 (Michalski, Mozetic, Hong, & Lavrac, 1986), leave the basic AQ algorithm intact and handle noise with pre-processing and post-processing techniques. Our objective in designing CN2 was to modify the AQ algorithm itself in ways that removed this dependence on specific examples and increased the

space of rules searched. This lets one apply statistical techniques, analogous to those used for tree pruning, in the generation of if-then rules, leading to a simpler induction algorithm.

One can identify several requirements that learning systems should meet if they are to prove useful in a variety of real-world situations:

- *Accurate classification.* The induced rules should be able to classify new examples accurately, even in the presence of noise.
- *Simple rules.* For the sake of comprehensibility, the induced rules should be as short as possible. However, when noise is present, overfitting can lead to long rules. Thus, to induce short rules, one must usually relax the requirement that the induced rules be consistent with all the training data. The choice of how much to relax this requirement involves a trade-off between accuracy and simplicity (Iba, Wogulis, & Langley, 1988).
- *Efficient rule generation.* If one expects to use large example sets, it is important that the algorithm scales up to complex situations. In practice, the time taken for rule generation should be linear in the size of the example set.

With these requirements in mind, this paper presents a description and empirical evaluation of CN2, a new induction algorithm. This system combines the efficiency and ability to cope with noisy data of ID3 with the if-then rule form and flexible search strategy of the AQ family. The representation for rules output by CN2 is an ordered set of if-then rules, also known as a *decision list* (Rivest, 1987). CN2 uses a heuristic function to terminate search during rule construction, based on an estimate of the noise present in the data. This results in rules that may not classify all the training examples correctly, but that perform well on new data.

In the following section we describe CN2 and three other systems used for our comparative study. In Section 3 we consider the time complexity of the various algorithms and in Section 4 we compare their performance on three medical diagnosis tasks. We also compare the performance of ASSISTANT and CN2 on two synthetic tasks. In Section 5 we discuss the significance of our results, and we follow this with some suggestions for future work in Section 6.

2. CN2 and related algorithms

CN2 incorporates ideas from both Michalski's (1969) AQ and Quinlan's (1983) ID3 algorithms. Thus we begin by describing Kononenko et al.'s (1984) ASSISTANT, a variant of ID3, and AQR, the authors' reconstruction of Michalski's method. After this, we present CN2 and discuss its relationship to these systems. We also describe a simple Bayesian classifier, which provides a reference for the performance of the other algorithms. We characterize each system along three dimensions:

- the representation language for the induced knowledge;
- the performance engine for executing the rules; and
- the learning algorithm and its associated search heuristics.

In all of our experiments, the example description language consisted of attributes, attribute values, and user-specified classes. This language was the same for each algorithm.

2.1 ASSISTANT

The ASSISTANT algorithm (Kononenko et al., 1984) is a descendant of Quinlan's ID3 (1983), and incorporates a tree-pruning mechanism for handling noisy data.

2.1.1 *Concept description and interpretation in ASSISTANT*

ASSISTANT represents acquired knowledge in the form of decision trees. An internal node of a tree specifies a test of an attribute, with each outgoing branch corresponding to a possible result of this test. Leaf nodes represent the classification to be assigned to an example.

To classify a new example, a path from the root of the decision tree to a leaf node is traced. At each internal node reached, one follows the branch corresponding to the value of the attribute tested at that node. The class at the leaf node represents the class prediction for that example.

2.1.2 *The ASSISTANT learning algorithm*

ASSISTANT induces a decision tree by repeatedly specializing leaf nodes of an initially single-node tree. The specialization operation involves replacing a leaf node with an attribute test, and adding new leaves to that node corresponding to the possible results of that test. Heuristics determine the attribute on which to test and when to stop specialization. Table 1 summarizes this algorithm.

2.1.3 *Heuristic functions in ASSISTANT*

ASSISTANT uses an entropy measure to guide the growth of the decision tree, as described by Quinlan (1983). This corresponds to the function IDM in Table 1. In addition, the algorithm can apply a tree-cutoff method based on an estimate of maximal classification precision. This technique estimates whether additional branching would reduce classification accuracy and, if so, terminates search (there are no user-controlled parameters in this calculation). This cutoff criterion corresponds to the function TE in Table 1. If ASSISTANT is to generate an 'unpruned' tree, the termination criterion TE(E) is satisfied if all the examples E have the same class value.

2.2 AQR

AQR is an induction system that uses the basic AQ algorithm (Michalski, 1969) to generate a set of classification rules. Many systems use this algorithm in a more sophisticated manner than AQR to improve predictive accuracy and rule simplicity; e.g., AQ11 (Michalski & Larson, 1983) uses a more complex method of rule interpretation that involves degrees of confirmation. AQR is a reconstruction of a straightforward AQ-based system.

Table 1. The core of the ASSISTANT algorithm.

```

Let E be a set of classified examples.
Let A be a set of attributes for describing examples.
Let TE(E) be a termination criterion.
Let  $IDM(A_i, E)$  be an evaluation function, where  $A_i \in A$ .

Procedure Assistant(E)
  If E satisfies the termination criterion TE(E),
  Then return a leaf node for TREE, labelled with
    the most common class of examples in E.
  Else let  $A_{best} \in A$  be the attribute with the
    largest value of the function  $IDM(A_{best}, E)$ .
    For each value  $V_j$  of attribute  $A_{best}$ ,
      Generate subtrees using ASSISTANT( $E_j$ ),
        where  $E_j$  are those examples in E with
        value  $V_j$  for attribute  $A_{best}$ .
    Return a node labelled as a test on attribute  $A_{best}$ 
    with these subtrees attached.

```

2.2.1 Concept description and interpretation in AQR

AQR induces a set of decision rules, one for each class. Each rule is of the form 'if <cover> then predict <class>', where <cover> is a Boolean combination of attribute tests as we now describe. The basic test on an attribute is called a *selector*. For instance, $\langle \text{Cloudy} = \text{yes} \rangle$, $\langle \text{Weather} = \text{wet} \vee \text{stormy} \rangle$, and $\langle \text{Temperature} > 60 \rangle$ are all selectors. AQR allows tests in the set $\{=, \leq, >, \neq\}$. A conjunction of selectors is called a *complex*, and a disjunct of complexes is called a *cover*. We say that an expression (a selector, complex, or cover) *covers* an example if the expression is true of the example. Thus, the empty complex (a conjunct of zero attribute tests) covers all examples and the empty cover (a disjunct of zero complexes) covers no examples. A cover is stored along with an associated class value, representing the most common class of training examples that it covers.

In AQR, a new example is classified by finding which of the induced rules have their conditions satisfied by the example. If the example satisfies only one rule, then one assigns the class predicted by that rule to the example. If the example satisfies more than one rule, then one predicts the most common class of training examples that were covered by those rules. If the example is not covered by any rule, then it is assigned by default to the class that occurred most frequently in the training examples.

2.2.2 The AQR learning algorithm

The AQ rule-generation algorithm has been described elsewhere (Michalski & Larson, 1983; Michalski & Chilausky, 1980; O'Rorke, 1982), and the AQR system is an instance of this general algorithm. AQR generates a decision rule

Table 2. The AQR algorithm for generating a class cover.

Let POS be a set of positive examples of class C.
Let NEG be a set of negative examples of class C.
Procedure AQR(POS, NEG)
Let COVER be the empty cover.
While COVER does not cover all examples in POS,
Select a SEED (a positive example not covered by COVER).
Let STAR be STAR(SEED, NEG) (a set of complexes that
cover SEED but that cover no examples in NEG).
Let BEST be the best complex in STAR
according to user-defined criteria.
Add BEST as an extra disjunct to COVER.
Return COVER.
Procedure STAR(SEED, NEG)
Let STAR be the set containing the empty complex.
While any complex in STAR covers some negative examples in NEG,
Select a negative example E_{neg} covered by a complex in STAR.
Specialize complexes in STAR to exclude E_{neg} by:
Let EXTENSION be all selectors that cover SEED but not E_{neg} .
Let STAR be the set $\{x \wedge y x \in STAR, y \in EXTENSION\}$.
Remove all complexes in STAR subsumed by other complexes.
Repeat until size of STAR $\leq maxstar$ (a user-defined maximum):
Remove the worst complex from STAR.
Return STAR.

for each class in turn. Having chosen a class on which to focus, it forms a disjunct of complexes (the cover) to serve as the condition of the rule for that class. This process occurs in stages; each stage generates a single complex, and then removes the examples it covers from the training set. This step is repeated until enough complexes have been found to cover all the examples of the chosen class. The entire process is repeated for each class in turn. Table 2 summarizes the AQR algorithm.

2.2.3 Heuristic functions in AQR

The particular heuristic functions used by the AQ algorithm are implementation dependent. The heuristic used by AQR to choose the best complex is "maximize the number of positive examples covered." The heuristic used to trim the partial star during generation of a complex is "maximize the sum of positive examples covered and negative examples excluded." In the case of a tie for either heuristic, the system prefers complexes with fewer selectors. Seeds are chosen at random and negative examples nearest to the seed are picked first, where distance is the number of attributes with different values in the seed and negative example. In the case of contradictions (i.e., if the seed

and negative example have identical attribute values) the negative example is ignored and a different one is chosen, since the complex cannot be specialized to exclude it but still include the seed.

2.3 The CN2 algorithm

Now that we have reviewed ASSISTANT and AQR, we can turn to CN2, a new algorithm that combines aspects of both methods. We begin by describing how the general approach arises naturally from consideration of the decision-tree and AQ algorithms and then consider its details.

2.3.1. Relation to ID3 and AQ

ID3 can be easily adapted to handle noisy data by virtue of its top-down approach to tree generation. During induction, *all* possible attribute tests are considered when ‘growing’ a leaf node in the tree, and entropy is used to select the best one to place at that node. Overfitting of decision trees can thus be avoided by halting tree growth when no more significant information can be gained. We wish to apply a similar method to the induction of if-then rules.

The AQ algorithm, when generating a complex, also performs a general-to-specific search for the best complex. However, the method only considers specializations that exclude some particular covered negative example from the complex while ensuring some particular ‘seed’ positive example remains covered, iterating until all negative examples are excluded. As a result, AQ searches only the space of complexes that are completely consistent with the training data. The basic algorithm employs a beam search, which can be viewed as several hill-climbing searches in parallel.

For the CN2 algorithm, we have retained the beam search method of the AQ algorithm but removed its dependence on specific examples during search and extended its search space to include rules that do not perform perfectly on the training data. This is achieved by broadening the specialization process to examine all specializations of a complex, in much the same way that ID3 considers all attribute tests when growing a node in the tree. Indeed, with a beam width of one the CN2 algorithm behaves equivalently to ID3 growing a single tree branch. This top-down search for complexes lets one apply a cutoff method similar to decision-tree pruning to halt specialization when no further specializations are statistically significant.

Finally, we note that CN2 produces an ordered list of if-then rules, rather than an unordered set like that generated by AQ-based systems. Both representations have their respective advantages and disadvantages for comprehensibility. Order-independent rules require some additional mechanism to resolve any rule conflicts that may occur, thus detracting from a strict logical interpretation of the rules. Ordered rules also sacrifice in comprehensibility, in that the interpretation of a single rule is dependent on the other rules that precede it in the list.¹

¹One can make CN2 produce unordered if-then rules by appropriately changing the evaluation function; e.g., one can use the same evaluation function as AQR, then generate a rule set for each class in turn.

2.3.2 Concept description and interpretation in CN2

Rules induced by CN2 each have the form ‘if <complex> then predict <class>’, where <complex> has the same definition as for AQR, namely a conjunction of attribute tests. This ordered rule representation is a version of what Rivest (1987) has termed *decision lists*. The last rule in CN2’s list is a ‘default rule’, which simply predicts the most commonly occurring class in the training data for all new examples.

To use the induced rules to classify new examples, CN2 tries each rule in order until one is found whose conditions are satisfied by the example being classified. The class prediction of this rule is then assigned as the class of the example. Thus, the ordering of the rules is important. If no induced rules are satisfied, the final default rule assigns the most common class to the new example.

2.3.3 The CN2 learning algorithm

Table 3 presents a summary of the CN2 algorithm. This works in an iterative fashion, each iteration searching for a complex that covers a large number of examples of a single class *C* and few of other classes. The complex must be both predictive and reliable, as determined by CN2’s evaluation functions. Having found a good complex, the algorithm removes those examples it covers from the training set and adds the rule ‘if <complex> then predict *C*’ to the end of the rule list. This process iterates until no more satisfactory complexes can be found.

The system searches for complexes by carrying out a pruned general-to-specific search. At each stage in the search, CN2 retains a size-limited set or *star* *S* of ‘best complexes found so far’. The system examines only specializations of this set, carrying out a beam search of the space of complexes. A complex is specialized by either adding a new conjunctive term or removing a disjunctive element in one of its selectors. Each complex can be specialized in several ways, and CN2 generates and evaluates all such specializations. The star is trimmed after completion of this step by removing its lowest ranking elements as measured by an evaluation function that we will describe shortly.

Our implementation of the specialization step is to repeatedly *intersect*² the set of all possible selectors with the current star, eliminating all the null and unchanged elements in the resulting set of complexes. (A null complex is one that contains a pair of incompatible selectors, e.g., $\text{big} = y \wedge \text{big} = n$.) CN2 deals with continuous attributes in a manner similar to ASSISTANT – by dividing the range of values of each attribute into discrete subranges. Tests on such attributes examine whether a value is greater or less (or equal) than the values at subrange boundaries. The complete range of values and size of each subrange is provided by the user.

²The intersection of set *A* with set *B* is the set $\{x \wedge y | x \in A, y \in B\}$. For example, $\{a \wedge b, a \wedge c, b \wedge d\}$ intersected with $\{a, b, c, d\}$ is $\{a \wedge b, a \wedge b \wedge c, a \wedge b \wedge d, a \wedge c, a \wedge c \wedge d, b \wedge d, b \wedge c \wedge d\}$. If we now remove unchanged elements in this set, we obtain $\{a \wedge b \wedge c, a \wedge b \wedge d, a \wedge c \wedge d, b \wedge c \wedge d\}$.

Table 3. The CN2 induction algorithm.

```

Let E be a set of classified examples.
Let SELECTORS be the set of all possible selectors.

Procedure CN2(E)
  Let RULE_LIST be the empty list.
  Repeat until BEST_CPX is nil or E is empty:
    Let BEST_CPX be Find_Best_Complex(E).
    If BEST_CPX is not nil,
      Then let E' be the examples covered by BEST_CPX.
      Remove from E the examples E' covered by BEST_CPX.
      Let C be the most common class of examples in E'.
      Add the rule 'If BEST_CPX then the class is C'
        to the end of RULE_LIST.
  Return RULE_LIST.

Procedure Find_Best_Complex(E)
  Let STAR be the set containing the empty complex.
  Let BEST_CPX be nil.
  While STAR is not empty,
    Specialize all complexes in STAR as follows:
    Let NEWSTAR be the set {x ∧ y | x ∈ STAR, y ∈ SELECTORS}.
    Remove all complexes in NEWSTAR that are either in STAR (i.e.,
      the unspecialized ones) or null (e.g., big = y ∧ big = n).
    For every complex Ci in NEWSTAR:
      If Ci is statistically significant and better than
        BEST_CPX by user-defined criteria when tested on E,
      Then replace the current value of BEST_CPX by Ci.
    Repeat until size of NEWSTAR ≤ user-defined maximum:
      Remove the worst complex from NEWSTAR.
    Let STAR be NEWSTAR.
  Return BEST_CPX.

```

For dealing with unknown attribute values, CN2 uses the simple method of replacing unknown values with the most commonly occurring value for that attribute in the training data. In the case of numeric attributes, it uses the midvalue of the most commonly occurring subrange.

2.3.4 Heuristics in CN2

The CN2 algorithm must make two heuristic decisions during the learning process, and it employs two evaluation functions to aid in these decisions. First it must assess the quality of complexes, determining if a new complex should replace the 'best complex' found so far and also which complexes in the star S to discard if the maximum size is exceeded. Computing this involves first finding the set E' of examples which a complex *covers* (i.e., which satisfy all of

its selectors) and the probability distribution $P = (p_1, \dots, p_n)$ of examples in E' among classes (where n is the number of classes represented in the training data). CN2 then uses the information-theoretic entropy measure

$$Entropy = - \sum_i p_i \log_2(p_i)$$

to evaluate complex quality (the lower the entropy the better the complex). This function thus prefers complexes covering a large number of examples of a single class and few examples of other classes, and hence such complexes score well on the training data when used to predict the majority class covered.

The entropy function was used in preference to a simple 'percentage correct' measure, such as taking $max(P)$, for two reasons. First, entropy will distinguish probability distributions such as $P = (0.7, 0.1, 0.1, 0.1)$ and $P = (0.7, 0.3, 0, 0)$ in favor of the latter, whereas $max(P)$ will not. This is desirable, since there exist more ways of specializing the latter to a complex identifying only one class. If the examples of the majority class are excluded by specialization, the distributions become $P = (0, 0.33, 0.33, 0.33)$ and $P = (0, 1, 0, 0)$, respectively. Second, the entropy measure tends to direct the search in the direction of more significant rules; empirically, rules of low entropy also tend to have high significance.

The second evaluation function tests whether a complex is *significant*. By this we mean a complex that locates a regularity unlikely to have occurred by chance, and thus reflects a genuine correlation between attribute values and classes. To assess significance, CN2 compares the *observed* distribution among classes of examples satisfying the complex with the *expected* distribution that would result if the complex selected examples randomly. Some differences in these distributions will result from random variation. The issue is whether the observed differences are too great to be accounted for purely by chance. If so, CN2 assumes that the complex reflects a genuine correlation between attributes and classes.

To test significance, the system uses the likelihood ratio statistic (Kalbfleish, 1979). This is given by

$$2 \sum_{i=1}^n f_i \log(f_i/e_i),$$

where the distribution $F = (f_1, \dots, f_n)$ is the observed frequency distribution of examples among classes satisfying a given complex and $E = (e_1, \dots, e_n)$ is the expected frequency distribution of the same number of examples under the assumption that the complex selects examples randomly. This is taken as the $N = \sum f_i$ covered examples distributed among classes with the same probability as that of examples in the entire training set. This statistic provides an information-theoretic measure of the (noncommutative) distance between the two distributions.³ Under suitable assumptions, one can show that this statistic is distributed approximately as χ^2 with $n - 1$ degrees of freedom.

³We assume that F is continuous with respect to E ; i.e., that the f_i are zero when the e_i are zero.

This provides a measure of indicates significance – the lower the score, the more likely that the apparent regularity is due to chance.

Thus these two functions – entropy and significance – serve to determine whether complexes found during search are both ‘good’ (have high accuracy when predicting the majority class covered) and ‘reliable’ (the high accuracy on training data is not just due to chance). CN2 uses these functions to repeatedly search for the ‘best’ complex that also passes some minimum threshold of reliability until no more reliable complexes can be found.

2.4 A Bayesian classifier

To establish a reference point, we also implemented a simple Bayesian classifier and compared its behavior to that of the other algorithms.

2.4.1 Bayesian concept description and interpretation

This classifier represents its ‘decision rule’ as a matrix of probabilities $p(v_j|C_k)$ specifying the probability of occurrence of each attribute value given each class. To classify a new example, one applies Bayes’ theorem⁴

$$p(C_i | \bigwedge v_j) = \frac{p(\bigwedge v_j | C_i) p(C_i)}{\sum_k p(\bigwedge v_j | C_k) p(C_k)},$$

where the summation is over the n classes and $p(C_i | \bigwedge v_j)$ denotes the probability that the example is of class C_i given v_j . One calculates this probability for every class, and then selects the class with the highest probability. The term $p(C_k)$ is estimated from the distribution of the training examples among classes. If one assumes independence of attributes, $p(\bigwedge v_j | C_k)$ can be calculated using

$$p(\bigwedge v_j | C_k) = \prod_j p(v_j | C_k)$$

and the values $p(v_j | C_k)$ from the probability matrix. Note that, unlike the other algorithms we have discussed, our implementation of the Bayesian classifier must examine the values of all attributes when making a prediction.

We should note that there also exist more sophisticated applications of the Bayes’ rule in which the attribute tests are ordered (Wald, 1947). Such a sequential technique adds the contribution of each test to a total; when this score exceeds a threshold, the algorithm exits with a class prediction. Such an interpretation may be more comprehensible to a user than the approach we have used, as well as limiting the tests required for classification.

2.4.2 The Bayesian learning algorithm

The Bayesian learning method constructs the matrix $p(v_j | C_k)$ from the training examples by examining the frequency of values in each class. One can compute this matrix either incrementally, incorporating one instance at a time, or nonincrementally, using all data at the outset.

⁴The $\bigwedge$ symbol for conjunction, $\bigwedge v_j$ denoting a conjunct of attribute values all occurring in an example.

2.4.3 Bayesian heuristics

Sometimes a value of zero is calculated from the training data for some elements of the $p(v_j|C_k)$ matrix. Like all elements of the matrix, this number is subject to error due to the finite training data available. However, as classification of new examples involves multiplying elements together, a zero element can have drastic effect, nullifying the effect of all other probabilities in the multiplication. To avoid this, we assume that zero elements in the matrix would, given more data, converge on a small, non-zero value and hence replace the zeros with some appropriate estimate. In our implementation a value of $p(C_k) \times (1/N)$ was used, where N is the number of training examples. The factor $1/N$ represents the increasing certainty that this element must have an almost-zero value with increasing size of training data.

2.5 The default rule

Finally, we examined a fifth ‘algorithm’ that simply assigns the most commonly occurring class to all new examples, with no reference to their attributes at all. As we will see in Section 4, this simple procedure produced comparable performance to that of the other algorithms in one of the domains, and thus provided another useful reference point.

3. Time complexity of the algorithms

The ASSISTANT, AQR and CN2 algorithms all search a very large space of concept descriptions, and all use heuristics to guide this search. Furthermore, all three algorithms attempt to produce structures that are both consistent with the training examples and as compact as possible. In the design of such algorithms, there is a tradeoff between execution speed and the size of the induced structures. In each case, the exhaustive search for a *smallest* set of structures, although desirable, is computationally infeasible.

A major application of these algorithms is to extract useful information from very large databases, perhaps with millions of examples. With this in mind, it is worth examining the complexity of each algorithm. To be practical for very large problems, their behavior should be linear, or at least near-linear, in the number of examples and attributes.

Since the overall complexity of each algorithm is domain-dependent, we instead provide upper bounds for the critical components of the algorithms. For example, we do not consider the complexity of the cutoff procedure used by ASSISTANT. In our treatment, we will use e to denote the size of the example set, a to stand for the number of attributes, and s to represent the maximum star size (for CN2 and AQR). We also assume that each attribute is binary valued and that there are two classes.⁵

⁵One might also consider the complexity as a function of the number of distinct attribute values and classes. We have not done this in our analysis.

3.1 Time complexity of ASSISTANT

The critical component in ASSISTANT is the process of selecting a test attribute on which to branch. Each such choice involves the following operations:

1. For each attribute, example counts are put in an array, indexed by class and attribute. This takes time $O(e \cdot a)$;
2. The entropy function is calculated for each attribute, taking time $O(a)$;
3. Once the best attribute is found, the examples are divided into two sets;⁶ this takes time $O(e)$.

Therefore, the overall time for a single attribute choice is $O(a \cdot e)$. The time taken to construct the complete tree depends very much on the structure of the tree. It seems reasonable to use the first figure only for comparative purposes, as argued above. Thus the amount of time taken by ASSISTANT for the basic attribute selection operation is a linear function of the number of examples, when the number of classes and attributes are held constant.

We should note that extensions to this algorithm that use real-valued attributes such as ACLS (Paterson & Niblett, 1982), must sort the examples by attribute value at the first stage. This increases the overall time bound to $O(a \cdot e \log e)$.

3.2 Time complexity of CN2

The basic operation in CN2 is the specialization of the complexes in the current star. The number of single-selector complexes without disjuncts is $2a$. The number of intermediate complexes generated is at most $a \cdot s$, and the time taken to evaluate an example against a complex is bounded by $O(a)$. Three steps are required for this specialization operation:

1. Multiplying each complex in the star by the set of single selector rules; this takes time $O(a \cdot s)$;
2. Evaluation of each complex, taking time $O(s \cdot e \cdot a)$;
3. Sorting the complexes by value and then trimming the star, which takes time $O(a \cdot s \log(a \cdot s))$.

Therefore, the overall time for a single specialization step is bounded by $O(a \cdot s(e + \log(a \cdot s)))$. As with ASSISTANT, the time required is a linear function of the number of examples. If we restrict the size of the star to one, the time required has the same order as for ASSISTANT. In general, experience indicates that the time constants involved are somewhat less for ASSISTANT and other variants on ID3 than for CN2.

3.3 Time complexity of AQR

In AQR, the basic operation is the specialization of complexes in a star. This operation is similar to that of CN2, except that one only generates spe-

⁶With appropriate data structures, it may be possible to do much of this work in the first stage, but this does not affect the complexity class. Similarly, one can include any termination test that is linear in the number of examples.

cializations that cause a negative example to be uncovered by complexes in AQR's star. We show the complexity of this operation is the same as that of CN2. For each negative example, the following steps are performed:

1. A negative example is found by iterating through the negative set. We assume that the number of negative examples is not less than some fixed fraction of the entire example set. This takes time $O(e \cdot s)$;
2. The set of selectors that distinguish the negative example from the seed are found; this takes time $O(a)$;
3. Each complex in the star is specialized by intersection with this set of selectors, taking time $O(a \cdot s)$;
4. The resulting complexes are evaluated, which takes time $O(a \cdot s \cdot e)$;
5. The complexes are sorted and the star trimmed, taking time $O(a \cdot s \log(a \cdot s))$.

Thus, for each negative example the time is bounded by $O(a \cdot s(e + \log(a \cdot s)))$. This is the same figure as obtained for CN2. Observe that the number of iterations of this process (making the star disjoint from a negative example) is bounded by the number of attributes, not by the number of examples.

In practice, although the order of time taken by the algorithms for this particular operation of producing a new star is the same, CN2 is faster overall than AQR. This is because the number of iterations of this operation is lower in CN2 than in AQR, since CN2 may halt specialization of a complex before it performs perfectly on the training examples. Also, CN2 may halt the entire search for rules before all the training examples are covered if no further statistically significant rules can be found.

3.4 Time complexity of the Bayesian classifier

The time complexity of the Bayes' classifier for generating a probability matrix is $O(a \cdot e)$, where a is the number of attributes and e the number of examples. This learning algorithm was substantially faster than the other algorithms because the run time is independent of the decision 'rule' generated. In addition, this basic operation is performed only once, unlike the above algorithms in which the basic operation is repeatedly applied.

3.5 Summary and actual run times

We have shown that the time complexity of the basic learning step for all the algorithms tested is linear in the number of examples, with $O(a \cdot e)$ for ASSISTANT and the Bayes' classifier and $O(a \cdot e \cdot s)$ for AQR and CN2. This is an essential requirement for any algorithm that must work with very large data sets.

However, the time complexity of the entire induction process, requiring iteration of the basic learning steps, is also important. With ideal noise-tolerant algorithms, given a certain minimum number of examples, concept descriptions representing only the genuine regularities in the data should be induced. Additional examples should not cause the concept description to grow further and become overfitted, hence in this ideal case the above figures also represent

the time complexity of the overall learning task. When this ideal is not met, as when one seeks a concept description that classifies the training data perfectly, the complexity increases. In such cases, CN2 would at worst induce e rules of length a , giving an overall time complexity of $O(a^2 \cdot e^2 \cdot s)$ (Chan, 1988). ASSISTANT sorts a total of e examples among a attributes for each level of the tree, giving an overall time complexity of $O(a^2 \cdot e)$ as the tree depth is bounded by a . The worst-case time complexity for AQR is similar to that for CN2.

The actual run times are revealing, although it is difficult to make quantitative comparisons due to differences in implementation language and method. Run times for each algorithm were obtained for the lymphography domain (Section 4.2.1) using a four-megabyte Sun 3/75. ASSISTANT, implemented in about 5000 lines of Pascal, took one minute run time. CN2 and AQR, each implemented in about 400 lines of Prolog and with a value of fifteen for *maxstar*, took 15 and 170 minutes runtime respectively. The Bayesian classifier, implemented in 150 lines of Prolog, took a few seconds to calculate its probability matrix. Although it is difficult to draw conclusions from the absolute run times, it is our opinion that the ordering of these run times (Bayes fastest, followed by ASSISTANT, CN2 and AQR) is a fair reflection of the relative computation involved in using the algorithms. More detailed empirical comparisons of time and memory requirements of ID3 and the AQ-based system AQ11P have been conducted by O'Rorke (1982) and Jackson (1985) in the domain of chess end games.

4. Experiments with the algorithms

Other aspects of the systems' behaviors lend themselves more to experimental study than analysis. Below we describe the dependent measures used in our experiments with the algorithms. After this, we describe the results of our studies with three natural domains and two artificial domains.

4.1 Dependent measures

In addition to computational complexity, we are interested in two other aspects of the algorithms' behaviors – classificational accuracy and syntactic complexity of the acquired structure. This twofold evaluation is motivated by considering these systems as knowledge-acquisition tools for expert systems. A useful system should induce rules that are accurate, so that they perform well, and comprehensible, so that they can be validated by an expert and used for explanation.

We measure each algorithm's classification accuracy by splitting the data into a training set and a test set, presenting the algorithm with the training set to induce a concept description and then measuring the percentage of correct predictions made by that concept description on the test set. Quinlan (1983, 1987a) and others have taken a similar approach to measuring accuracy.

Cross-algorithm comparisons of the complexity of concept descriptions are difficult due to the differences in representation and the degree of subjectivity involved in judging complexity. Thus, we will only compare the gross features of the knowledge structures induced by the different algorithms. For

ASSISTANT's decision trees, we measure complexity by the number of nodes (including leaves) in the tree. For CN2 and AQR, we measure complexity by the number of selectors in the final rule list and rule set respectively. These measures reveal the gross features of the induced decision rules. More detailed measures of rule complexity have been made by O'Rorke (1982) but are not used here. We assign a complexity of one to the default rule, based on its equivalence to a decision tree with a single node.

Assessing the complexity of a Bayesian description is more difficult. One could count the number of elements in the $p(V_j|C_k)$ matrix. Thus, for a domain with n classes and a attributes, each with an average of v possible values, the complexity would be $a \times v \times n$. However, such a measure is independent of the training examples, and it ignores features of the matrix that may make it more comprehensible (e.g., a few elements may be very large and the rest small). Still, lacking any better measure, we provide the size of the matrix as a rough guide.

4.2 Experiments on natural domains

The above algorithms were tested on three sets of medical data, which we will describe shortly. These data were obtained from the Institute of Oncology at the University Medical Center in Ljubljana, Yugoslavia (Kononenko et al., 1984). In each test, 70% of the training examples were selected at random from the entire data set, and the remaining 30% of the data were used for testing. The algorithms were all run on the same training data and their induced knowledge structures tested using the same test data. Five such tests were performed for each of the three domains, and the results were averaged. These data are thus identical to those used to test AQ15 in Michalski et al. (1986), though the particular random 70% and 30% samples are different. Both CN2 and AQR were given a value of 15 for *maxstar* in all runs.

4.2.1 Three medical domains

Table 4 summarizes the characteristics of the three medical domains used in the experiments. The first of these involved lymphography. For patients with suspected cancer, it is important for physicians to distinguish between patients that are healthy and those with metastases or malignant lymphoma. Patient data relating to this task were collected from Ljubljana's Oncology Institute. These data were consistent; i.e., examples of any two classes were always different. All the tested algorithms produced fairly simple and accurate rules. Unlike the other two domains, this data set was not submitted to a detailed checking after its original compilation by the Medical Center, and thus may contain errors in attribute values.

The second domain involved predicting whether patients who have undergone breast cancer operations will experience recurrence of the illness within five years of the operation. The recurrence rate is about 30%, and hence such prognosis is important for determining post-operational treatment. These data were verified after collection, and thus are likely to be relatively free of errors.

Table 4. Description of the three medical domains.

DOMAIN PROPERTY	LYMPHO- GRAPHY	BREAST CANCER	PRIMARY TUMOR
NUMBER OF ATTRIBUTES	18	9	17
VALUES PER ATTRIBUTE			
MINIMUM	2	2	2
MAXIMUM	8	5	3
AVERAGE	3.3	2.8	2.2
NUMBER OF CLASSES	4	2	22
NUMBER OF EXAMPLES	148	286	339
DISTRIBUTION OF EXAMPLES AMONG THE CLASSES	2, 81, 61, 4	85, 201	84, 20, 9, 14, 39, 1, 14, 6, 0, 2, 28, 16, 7, 24, 2, 1, 10, 29, 6, 2, 1, 24

The final medical domain focused on predicting the location of a primary tumor. Physicians distinguish between 22 possible locations, predicted from data such as age, histologic type of carcinoma, and possible locations of detected metastases; this is also important in determining treatment of patients. These data were inconsistent; i.e., examples of different classes existed with identical attribute values. They were verified after collection, and thus are likely to be relatively error-free. The set of attributes is relatively incomplete, and thus not sufficient to induce high-quality rules.

4.2.2 Results with natural domains

Table 5 presents the results for each algorithm on these domains, averaged over five runs. In each case, we present the average accuracy on the test data and the average complexity of the resulting knowledge structures. CN2 was tested using three values of significance threshold and ASSISTANT was run with and without pruning. The other systems have no such user-variable parameters.

The table contains some interesting regularities. The most important is that the algorithms designed to reduce problems caused by noisy data achieve a lower complexity without damaging their predictive accuracy. For example, in the lymphography domain, the version of CN2 with the highest threshold achieved the same classification accuracy as the other algorithms by inducing (on average) only eight rules, each containing 1.6 selectors. The tree-pruning version of ASSISTANT produced similar results.

Both systems apply a similar technique to reducing complexity, namely sometimes halting specialization of concept descriptions before they classify the training examples perfectly. As a result, ASSISTANT and CN2 avoid overfitting their decision trees and rules to the training data. This contrasts with the AQR algorithm, which specializes its rule set until it achieves as nearly complete consistency with the training data as possible, resulting in an overfitted rule set. Table 6 illustrates this effect by comparing accuracy on the training and test data.

Table 5. Accuracy and complexity of knowledge structures acquired by the algorithms in three natural domains. (Complexity for the Bayes' classifier is the size of the probability matrix.)

ALGORITHM	LYMPHOGRAPHY		BREAST CANCER		PRIMARY TUMOR	
	ACCUR.	COMP.	ACCUR.	COMP.	ACCUR.	COMP.
DEFAULT RULE	56%	1	71%	1	26%	1
ASSISTANT						
NO PRUNING	79%	41	62%	112	40%	178
PRUNING	78%	36	68%	44	42%	52
BAYES	83%	240 [†]	65%	540 [†]	39%	465 [†]
AQR	76%	76	72%	208	35%	562
CN2						
90% THRESH.	78%	24	70%	28	37%	33
95% THRESH.	81%	22	70%	20	36%	42
99% THRESH.	82%	12	71%	4	36%	19

[†] See discussion in Section 4.1 about difficulties in measuring the complexity of Bayesian classifiers.

The results also show that the Bayesian classifier does well, performing comparably to the more sophisticated algorithms in all three domains and giving the highest accuracy in the lymphography domain. Table 6 shows that this method regularly overfits the training data, but that its performance on the test set is still good. Even more surprising is the behavior of the frequency-based default rule, which outperforms ASSISTANT and the Bayes' method on the breast cancer domain. This suggests that there are virtually no significant correlations between attributes and classes in these data. This is reflected by CN2's inability to find significant rules in this domain at 99% threshold, suggesting that, in this domain at least, the significance test has been effective in filtering out rules representing chance regularities.

In general, the differences in performance seem to be due less to the learning algorithms than to the nature of the domains; for example the best classification accuracy for lymphography was barely half as high as that for primary tumor. This suggests the need for additional studies to examine the role of domain regularity on learning.

4.3 Experiments on artificial domains

To better understand the effects of overfitting, we experimented with CN2 and ASSISTANT on two artificial domains that let us control the amount of noise in the data.

4.3.1 Two artificial domains

Both domains contained twelve attributes and 200 examples that were evenly distributed between two classes. They differed only in the number of values each attribute could take (two in the first domain and eight in the second).

Table 6. Accuracy of the different algorithms on training and test data. The reported version of ASSISTANT incorporated pruning and the version of CN2 used a 99% threshold.

ALGORITHM	LYMPHOGRAPHY		BREAST CANCER		PRIMARY TUMOR	
	TRAIN	TEST	TRAIN	TEST	TRAIN	TEST
DEFAULT RULE	54%	56%	70%	71%	23%	26%
ASSISTANT	98%	78%	85%	68%	53%	42%
BAYES	89%	83%	70%	65%	48%	39%
AQR	100%	76%	100%	72%	75%	35%
CN2	91%	82%	72%	71%	37%	36%

In both cases, the target concept for one class could be stated as a simple conjunctive rule of the form ‘if $(a = v_1) \wedge \dots \wedge (d = v_1)$ then class X’. Both algorithms can represent such a regularity compactly. The second class was simply the negation of the first. Half of the data was used for training, half for testing, and the results averaged over five trials.

For each domain, we varied the amount of noise in the training data and measured the effect on complexity and on accuracy on the test data. Table 7 reports the results for the first artificial domain, with two values per attribute, and Table 8 for the second, with eight values per attribute.

The percentage of noise added indicates the proportion of attribute values in the training examples that have been randomized, where attributes chosen for randomization have equal chance of taking any of the possible values for that attribute. For the purposes of randomization, the class was treated simply as an additional attribute in the example description. Note that no noise was introduced into the test data.

4.3.2 Results with artificial domains

By experimenting with artificial domains, we can examine several features of the algorithms relating to their ability to handle noise. First, we can see the degradation of accuracy and simplicity of concept descriptions as noise levels are increased. Second, for CN2, it is also interesting to examine how the accuracy and simplicity of individual rules (as well as that of the rule set as a whole) is affected by noise in the data.

The results reveal some surprising features about both CN2 and ASSISTANT. Comparing classification accuracy alone, ASSISTANT performed better than CN2 in these particular domains. However, comparing complexity of concept description, CN2 produced simpler concept descriptions than ASSISTANT except at high levels of noise in the first domain.

Ideally, as the level of noise approaches 100%, both algorithms should fail to find any significant regularities in the data and thus converge on a concept description of complexity one (for CN2’s default rule alone or a single-node decision tree). However for CN2, this occurred only in the second of the two domains tested and did not occur in either domain for ASSISTANT. Indeed, in

Table 7. Results in artificial domain A1, with 12 attributes and 2 values per attribute.

NOISE LEVEL	CN2 (99% THRESHOLD)			ASSISTANT			
	TOTAL	NONDEF. [†]	COMP.	UNPRUNED		PRUNED	
	ACCUR.	ACCUR.		ACCUR.	COMP.	ACCUR.	COMP.
0%	95%	100%	3	99%	8	99%	8
2%	88%	99%	5	96%	16	98%	11
5%	88%	95%	10	91%	32	95%	16
10%	82%	95%	15	86%	45	91%	24
20%	73%	86%	20	76%	60	84%	27
40%	67%	76%	25	65%	74	76%	23
60%	56%	64%	26	62%	75	67%	23
100%	45%	49%	28	46%	85	43%	12

[†] This refers to the accuracy of those CN2 rules found by search; i.e., *excluding* the extra default rule ('everything is class X') at the end of the rule list. See discussion in Section 4.3.2.

the second domain ASSISTANT's tree-pruning mechanism did not prune the tree at all. CN2's generation of rules in the first domain, even at 100% noise level, probably results from a combination of the large number of rules searched (e.g., there are $12 \times 11 \times 10 = 1320$ rules of length three in the space) and the high coverage of these rules (each length three rule covers on average $100/2^3 = 12$ examples). Enough rules are searched so that, even with 99% significance test, some chance coverage of the 12 (average) examples will appear significant. This did not occur in the second domain, as the coverage of rules was considerably less; each length three rule covers on average $100/8^3 \approx 0.5$ examples, too few for the significance test to succeed.

These behaviors as the noise level approaches 100% suggest that the thresholding methods used in both CN2 and ASSISTANT need to be more sensitive to the properties of the application domain. Research on improvements to CN2's significance test (Chan, 1988) and ASSISTANT's pruning mechanism (Cestnik, Kononenko, & Bratko, 1987) is currently being conducted.

We also measured the accuracy of CN2's individual rules, as opposed to that of the entire rule set. Tables 7 and 8 include columns for 'non-default accuracy', which show the accuracy of CN2's rules excluding cases in which the default rule fires. These suggest that the rule list consists of high-accuracy rules plus a low-accuracy (50% in this domain) default rule at the end. This is a desirable property of the rule list if it is to be used for helping an expert articulate his or her knowledge, as each individual rule (apart from the default rule) represents a strong regularity in the training data. The decision-tree equivalent would involve examining the individual branches generated and their use in assisting an expert. Quinlan (1987b) has recently conducted work along these lines.

Table 8. Results in artificial domain A2, with 12 attributes and 8 values per attribute.

NOISE ADDED	CN2 (99% THRESHOLD)			ASSISTANT			
	TOTAL	NONDEF. [†]	COMP.	UNPRUNED		PRUNED	
	ACCUR.	ACCUR.		ACCUR.	COMP.	ACCUR.	COMP.
0%	93%	98%	8	99%	6	99%	6
2%	83%	99%	10	97%	12	97%	12
5%	86%	94%	13	96%	15	96%	15
10%	80%	98%	10	93%	22	93%	22
20%	73%	88%	15	85%	27	85%	27
40%	68%	82%	5	75%	33	75%	33
60%	63%	90%	4	66%	40	66%	40
100%	50%	58%	1	55%	43	55%	43

[†] This refers to the accuracy of those CN2 rules found by search; i.e., *excluding* the extra default rule ('everything is class X') at the end of the rule list. See discussion in Section 4.3.2.

5. Discussion

The results on the natural domains indicate that different methods of halting the rule specialization process, besides having the effect of reducing rule complexity, do not greatly affect predictive accuracy. This effect has been reported in a number of papers (Kononenko et al., 1984; Michalski et al., 1986; Niblett & Bratko, 1987). Indeed, it may be that this effect will occur with any technique, providing one does not exceed a certain maximum level of pruning. If so, then one should prefer the algorithm that most closely estimates this maximum level.

The results in Table 6 suggest that the 99% threshold for the CN2 algorithm is appropriate for the three natural domains. The accuracy on training data is close to that on test data, indicating that, in these domains at least, the algorithm is not overfitting the data. Additionally, high accuracy is maintained, indicating that the concept description is not underfitted either.

However, the results of the tests on the artificial domains, in particular the tests with 100% noise, indicate that the current measure of significance used by CN2 could be improved. As the noise level reaches 100%, the algorithm should ideally find no rules. The fact that this only occurred in one of the two artificial domains suggests that the significance measure should be more sensitive to properties of the domain in question.

In many ways the comparisons with the AQR system are unfair, as the AQ algorithm was never intended to deal alone with noisy data. It was included in these experiments to examine the basic AQ algorithm's sensitivity to noise. In practice it is rarely used on its own, instead being enhanced by a number of pre- and post-generation techniques. Experiments with the AQ15

system (Michalski et al., 1986) show that with post-pruning of the rules and a probability-based or 'flexible matching' method for rule application, one can achieve results similar to those of CN2 and ASSISTANT in terms of accuracy and complexity.

The principal advantage of CN2 over AQR is that the former algorithm supports a cutoff mechanism – it does not restrict its search to only those rules that are consistent with the training data. CN2 demonstrates that one can successfully control the search through the larger space of inconsistent rules with the use of judiciously chosen search heuristics. Second, by including a mechanism for handling noise in the algorithm itself, we have achieved a simple method for generating noise tolerant if-then rules that is easy to reproduce and analyze. In addition, interactive approaches to induction, in which the user interacts with the system during and after rule generation, introduce additional requirements, such as the need for good explanation facilities. In such cases, the logical rule interpretation used by CN2 should have practical advantages over the more complex probabilistic rule interpretation needed to apply order-independent rules (such as those generated by AQR) in which conflicts may occur.

Another result of interest is the high performance of the Bayesian classifier. Although the independence assumption of the classifier may be unjustified in the domains tested, it did not perform significantly worse in terms of accuracy than other algorithms, and it remains an open question as to how sensitive Bayesian methods are to violated independence assumptions. Although the probability matrices produced by the tested classifier are difficult to comprehend, the experiments suggest that variants of the Bayes' classifier which produce more comprehensible decision procedures would be worthy of further investigation.

6. Conclusions

In this paper we have demonstrated CN2, an induction algorithm that combines the best features of the ID3 and AQ algorithms, allowing the application of statistical methods similar to tree pruning in the generation of if-then rules. The CN2 system is similar to ASSISTANT in its efficiency and ability to handle noisy data, whereas it partially shares the representation language and flexible search strategy of AQR. By incorporating a mechanism for handling noise into the algorithm itself, a method for inducing if-then rules has been achieved that is noise-tolerant, simple to analyze, and easy to reproduce.

The experiments we have conducted show that, in noisy domains, the CN2 algorithm has comparable performance to that of ASSISTANT. By inducing concept descriptions based on if-then rules, CN2 provides a tool for assisting in the construction of knowledge-based systems where one desires classification procedures based on rules rather than decision trees. The most obvious improvement to the algorithm, suggested by the results on artificial domains, is an improvement to the significance measure used.

Acknowledgements

We thank Donald Michie and Claude Sammut for their careful reading and valuable comments on earlier drafts of this paper. We also thank Pat Langley and the reviewers for their detailed comments and suggestions about the presentation. We are grateful to G. Klanjšček, M. Soklič, and M. Zwitter of the University Medical Center, Ljubljana for the use of the medical data and to I. Kononenko for its conversion to a form suitable for the induction algorithms. This work was supported by the Office of Naval Research under contract N00014-85-G-0243 as part of the Cognitive Science Research Program.

References

- Cestnik, B., Kononenko, I., & Bratko, I. (1987). ASSISTANT 86: A knowledge-elicitation tool for sophisticated users. *Proceedings of the Second European Working Session on Learning* (pp. 31–45). Bled, Yugoslavia: Sigma Press.
- Chan, P. K. (1988). *A critical review of CN2: A polythetic classifier system* (Technical Report CS-88-09). Nashville, TN: Vanderbilt University, Department of Computer Science.
- Iba, W., Wogulis, J., & Langley, P. (1988). Trading off simplicity and coverage in incremental concept learning. *Proceedings of the Fifth International Conference on Machine Learning* (pp. 73–79). Ann Arbor, MI: Morgan Kaufmann.
- Jackson, J. (1985). *Economics of automatic generation of rules from examples in a chess end-game* (Technical Report UIUCDCS-F 85-932). Urbana: University of Illinois, Computer Science Department.
- Kalbfleish, J. (1979). *Probability and statistical inference* (Vol. 2). New York: Springer-Verlag.
- Kononenko, I., Bratko, I., & Roskar, E. (1984). *Experiments in automatic learning of medical diagnostic rules* (Technical Report). Ljubljana, Yugoslavia: E. Kardelj University, Faculty of Electrical Engineering.
- Michalski, R. S. (1969). On the quasi-minimal solution of the general covering problem. *Proceedings of the Fifth International Symposium on Information Processing* (pp. 125–128). Bled, Yugoslavia.
- Michalski, R. S., & Chilausky, R. (1980). Learning by being told and learning from examples: An experimental comparison of the two methods of knowledge acquisition in the context of developing an expert system for soybean disease diagnosis. *International Journal of Policy Analysis and Information Systems*, 4, 125–160.
- Michalski, R. S., & Larson, J. (1983). *Incremental generation of VL_1 hypotheses: The underlying methodology and the description of the program AQ11* (Technical Report ISG 83-5). Urbana: University of Illinois, Computer Science Department.
- Michalski, R. S., Mozetic, I., Hong, J., & Lavrac, N. (1986). The multi-purpose incremental learning system AQ15 and its testing application to three medical domains. *Proceedings of the Fifth National Conference on Artificial Intelligence* (pp. 1041–1045). Philadelphia: Morgan Kaufmann.

- Mowforth, P. (1986). *Some applications with inductive expert system shells* (TIOP 86-002). Glasgow, Scotland: Turing Institute.
- Niblett, T. (1987). Constructing decision trees in noisy domains. *Proceedings of the Second European Working Session on Learning* (pp. 67–78). Bled, Yugoslavia: Sigma Press.
- Niblett, T., & Bratko, I. (1987). Learning decision rules in noisy domains. In M. A. Bramer (Ed.), *Research and development in expert systems* (Vol. 3). Cambridge: Cambridge University Press.
- O'Rorke, P. (1982). *A comparative study of inductive learning systems AQ11P and ID3 using a chess end-game test problem* (Technical Report ISG 82-2). Urbana: University of Illinois, Computer Science Department.
- Paterson, A., & Niblett, T. (1982). *ACLS manual, Version 1* (Technical Report). Glasgow, Scotland: Intelligent Terminals Limited.
- Quinlan, J. R. (1983). Learning efficient classification procedures and their application to chess end games. In R. S. Michalski, J. G. Carbonell, & T. M. Mitchell (Eds.), *Machine learning: An artificial intelligence approach*. Los Altos, CA: Morgan Kaufmann.
- Quinlan, J. R. (1987a). Simplifying decision trees. *International Journal of Man-Machine Studies*, 27, 221–234.
- Quinlan, J. R. (1987b). Generating production rules from decision trees. *Proceedings of the Tenth International Joint Conference on Artificial Intelligence* (pp. 304–307). Milan, Italy: Morgan Kaufmann.
- Quinlan, J. R., Compton, P. J., Horn, K. A., & Lazarus, L. (1987). Inductive knowledge acquisition: A case study. *Applications of expert systems*. Wokingham, England: Addison-Wesley.
- Rivest, R. L. (1987). Learning decision lists. *Machine Learning*, 2, 229–246.
- Wald, A. (1947). *Sequential analysis*. New York: Wiley.