

Evolving Polynomials of the Inputs for Decision Tree Building

Chris J. Hinde and Anoud I. Bani-Hani

Loughborough University/Computer Science, Loughborough, UK

Email: {C.J.Hinde, A.I.Bani-Hani}@lboro.ac.uk

Thomas.W. Jackson

Loughborough University/Information Science, Loughborough, UK

Email: T.W.Jackson@lboro.ac.uk

Yen P. Cheung

Monash University/Clayton School of IT, Wellington, Australia

Email: Yen.Cheung@monash.edu.au

Abstract— The aim of this research is to extend the discrimination of a decision tree builder by adding polynomials of the base inputs to the inputs. The polynomials used to extend the inputs are evolved using the quality of the decision trees resulting from the extended inputs as a fitness function. Our approach generates a decision tree using the base inputs and compares it with a decision tree built using the extended input space. Results show substantial improvements. Rough set reducts are also employed and show no reduction in discrimination through the transformed space.

Index Terms—Decision tree building, Polynomials

I. INTRODUCTION

This paper addresses the well-known problem of data mining where given a set of data; the expected output is a set of rules. Decision trees using the ID3 approach [1], [2] are popular and in most cases successful in generating rules correctly. Extensions to ID3 such as C4.5 and CART are developed to cope with uncertain data. Fu et al [3] used C4.5 followed by a Genetic Algorithm (GA) to evolve better quality trees; in Fu's work C4.5 was used to seed a GA, which were then used as a basis for evolving better trees then using Genetic Programming (GP) techniques to cross over the trees. Many rule discovery techniques combining ID3 with other intelligent techniques such as genetic algorithms and genetic programming have also been suggested [4], [5], [6]. Generally, when using ID3 with genetic algorithms, individuals which are usually fixed length strings are used to represent decision trees and the algorithm evolves to find the optimal tree. When Genetic Programming is used to generate decision trees, individuals are variable length trees, which represent the decision tree. Variations in these approaches can be found in the gene encoding. One rule per individual as done in Greene [7], Freitas et al [8], [9] is a simple approach but the fitness of a single rule is not necessarily the best indicator of the quality of the discovered rule set. Encoding several rules in an individual requires longer and more complex operators

[10], [11]. In genetic programming, a program can be represented by a tree with rule conditions and/or attribute values in the leaf nodes and functions in the internal nodes. Here the tree can grow dynamically and pruning of the tree is necessary [12]. Papagelis & Kelles [13] used a gene to represent a decision tree and the GA then evolves to find the optimal tree, similar to Fu et al [3]. To further improve the quality of the trees, Eggermont et al [14] applied several fitness measures and ranked them according to their importance in to tackle uncertain data. Previous work has taken the input space as a given and used evolution to produce the trees. In this work, as we shall see, the trees are generated using a variant of C4.5 and the input space is evolved rather than the trees, in direct contrast to other workers.

A vast majority of the approaches use decision trees as a basis for the search in conjunction with either a GA or GP to further improve the quality of the trees. Our approach described in this paper addresses continuous data and adds polynomials of the input values to extend the input set. A GA is used to search the space for these polynomials based on the quality of the tree discovered using a version of C4.5.

II. ITERATIVE DISCRIMINATION

ID3, C4.5 and their derivatives proceed by selecting an attribute that results in an information gain with respect to the dependent variable. A simple data set with 2 continuous attributes that are linearly separable is shown in Fig. 1.

Applying C4.5 to the data set gives the result shown in Figure 2, which was first documented in [15]. If no errors are required over a large training set then the complexity of the decision tree grows with the size of the training set. This is unsatisfactory.

Anticipating the results of the proposed system a higher level discriminant of $x-y$ in addition to the two basic variables x and y would give the result shown in Fig. 3.


```

r2 > 8.125 : out (72.0)
r2 <= 8.125 :
|   r2 <= 0.625 : out (8.0)
|   r2 > 0.625 : in (48.0)

```

Evaluation on training data (128 items):
 Before Pruning After Pruning

Size	Errors	Size	Errors	Estimate
5	0(0.0%)	5	0(0.0%)	(3.1%)

Fig. 6. The decision tree produced by C4.5 from the augmented toroidal data. $r2$ is the sum of the squares of x and y .

With the data sets shown in Figs 7 and 11, C4.5 does not produce a tree at all. Of the two data sets presented a higher order combined attribute results in a concise tree where no tree is produced without the higher order attribute. In the case of the quadrant data set, a concise decision is possible with the unaugmented data set, one is not produced by C4.5.

A. Banded data set

This test shows a data set that does not project down onto 1 dimension. This 2 dimensional data set results in the following tree from C4.5, Fig. 8.

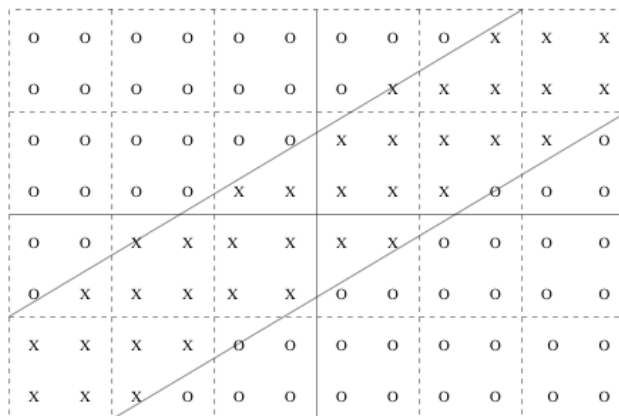


Fig. 7. A granularised version of banded linearly separable data.

out (128.0/52.0)

Evaluation on training data (128 items):

Before Pruning After Pruning

Size	Errors	Size	Errors	Estimate
1	52(40.6%)	1	52(40.6%)	(44.1%)

Fig. 8. The decision tree produced by C4.5 from the banded data.

The decision tree produced from the banded data, Figure 8, is shown in Fig. 8 and is almost useless. It does not reveal any useful information from the data. The most that can be gained from this data is that there are 52 elements in the concept and the rest are out. Adding the attribute $x - y$ gives the tree shown in Fig. 9, this is a good predictor and also makes the information held in the data clear.

B. Quadrant data set

This test shows a data set that cannot be discriminated by C4.5, however a decision tree does exist. It is shown in Fig. 10. This clear 2 dimensional data set results in the following tree from C4.5, Figure 12.

V. GENETIC ALGORITHM

The genetic algorithm attached to the front of c4.5 has a few special features. It follows most of the guidelines in [17], [18] and so has aspects designed to preserve inheritability and to ensure that no part of the genome has an inordinate effect on the phenotype. With this in mind the structure of the genome is made up from a set of integers, rather than a binary genome.

```

x-y <= -2 : out (38.0)
x-y > -2 :
|   x-y <= 1.5 : in (52.0)
|   x-y > 1.5 : out (38.0)

```

Evaluation on training data (128 items):
 Before Pruning After Pruning

Size	Errors	Size	Errors	Estimate
5	0(0.0%)	5	0(0.0%)	(3.2%)

Fig. 9. The decision tree produced by C4.5 from the banded data given the added input feature of $x-y$.

```

x <= 0.0 : (64.0)
|   y <= 0.0 : in(32.0)
|   y > 0.0 : out(32.0)
x > 0.0 : (64.0)
|   y <= 0.0 : out (32.0)
|   y > 0.0 : in(32.0)

```

Evaluation on training data (128 items)

Before Pruning After Pruning

Size	Errors	Size	Errors	Estimate
6	0(0.0%)	6	0(0.0%)	(3.2%)

Fig. 10. The decision tree, which could be used to discriminate the quadrant data, but cannot be produced by C4.5.

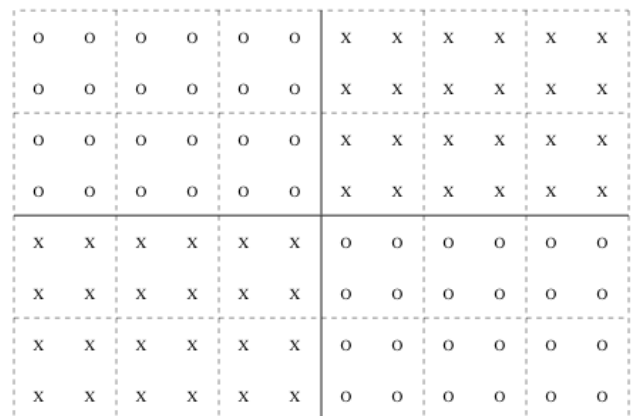


Fig. 11. A granularised version of quadrant data set.

out (128.0/52.0)

Evaluation on training data (128 items):

Before Pruning	After Pruning
Size Errors	Size Errors Estimate
1	52(40.6%)1 52(40.6%)(44.1%)

Fig. 12. The decision tree produced by C4.5 from the quadrant data.

```

x*y <= -2 : out (38.0)
x*y > -2 :
|       x*y <= 1.5 : in (52.0)
|       x*y > 1.5 : out (38.0)

```

Evaluation on training data (128 items):

Before Pruning	After Pruning
Size Errors	Size Errors Estimate
5 0(0.0%)	5 0(0.0%) (3.2%)

Fig. 13. The decision tree produced by C4.5 from the quadrant data given the added input feature of $x*y$.

A. Genetic Structure

The chromosome can deliver several genes corresponding to several combined attributes. The chromosome is a fixed maximum length and achieves a variable number of genes by an activation flag. Each gene delivers one new attribute and each variable is a linear combination of simpler variables.

1) *Variable.*: If the number in the variable slot is N and there are K basic continuous variables in the data set and M variables in the gene prior to this one then $N \bmod (K + M)$ refers to variable within those $K + M$ variables.

2) *Function.*: If the function is a monadic function then it is applied to variable 1, otherwise to both. The prototype system has a set of simple arithmetic functions, power, multiplication, division and subtraction. This is sufficient to extract all the decision trees we have considered.

3) *Number of genes and gene length.*: The variable length chromosome has disadvantages as the effect on the gene itself of the two fields that determine the length of the gene is considerably more than any other field and can be destructive. The variable length gene has similar disadvantages. The gene structure finally chosen for the system is shown in Fig. 14.

Active	
	Variable 1
	Function
	Variable 2

Fig. 14. This shows the basic structure of the gene adopted. The Active/Variable/Function/Variable is repeated up to the gene length.

This potentially has some of the properties of recessive genes that are attributed to diploid gene structures although no experiments have been conducted to determine this. An example gene is shown in Figure 15. This gene has 4 segments, 1 of which is active. Each segment has 2 attributes, some active and some not. The

function field is interpreted as 2 for plus, 3 for minus, 5 for multiply. No other function types are illustrated.

1		active
	1	x
	3	-
	2	y
0		inactive
	1	x
	5	*
	1	x
0		inactive
	2	y
	5	*
	2	y
0		inactive
	4	x^2
	2	-
	5	y^2

Fig. 15. An exemplar gene. x and y are variables number 1 and 2. The first new variable is $x - y$ and is variable number 3. As this is activated then it will be made available as an input to the decision tree generator. If variable 6 is activated then because it relies on variables 4 and 5 they will also be kept but not necessarily activated.

VI. EXEMPLAR DATA

The system described above was applied to some data sets taken from the Machine Learning Repository [19] in order to compare the capability of the system with other known decision tree generators.

The experiment compares the decision trees generated by C4.5 and the decision trees generated by C4.5 with the enhanced input space. The results consist of

- the percentage of correct results on the training set
- the percentage of correct results on the test set
- the of degrees of freedom for the decision space
- the probability that the result could not have arisen by chance
- the decision tree size

A. Experimental results

Each data set was split randomly into two sets, the training set which comprised 90% of the data and the test set, which comprised 10% of the data. The split was generated by choosing whether a particular data point was to be in the training set or the test set using a random number generator. This way any temporal aspects that may be in the data are accounted for. Notice the degrees of freedom are different for the training set and the test set, this is because there were no data elements belonging to one of the categories in the test set, where there were elements in the training set.

TABLE I
EXPERIMENTAL RESULTS FOR GLASS DATA SET

	C4.5	C4.5+GP
Train Correct	92.8	98.5
Test Correct	75	100
DOF Train	30	30
DOF Test	25	25
probability of not null	1.0	1.0
Tree size	43	61

The glass data set shows a considerable improvement for the enhanced input space, however the decision tree is larger.

The iris data set shows an improvement for the enhanced input space, but the improvement is marginal, however the decision tree is smaller.

TABLE II
EXPERIMENTAL RESULTS FOR IRIS DATA SET

	C4.5	C4.5+GP
Train Correct	98	100
Test Correct	100	100
DOF	6	6
probability of not null	0.99	0.99
Tree size	9	7

TABLE III
EXPERIMENTAL RESULTS FOR PIMA INDIANS DATA SET

	C4.5	C4.5+GP
Train Correct	82.7	98.3
Test Correct	74.5	84.2
DOF	2	2
probability of not null	1.0	1.0
Tree size	33	119

The Pima indians data set shows a considerable improvement for the enhanced input space. Both the training set and the test set show improvement. The enhanced decision tree is also considerably bigger, by a factor of nearly 4.

The experiments have shown that the enhanced system is able to significantly improve the quality of the decisions made, however this is often at the expense of a larger tree. The test on the iris data set indicates that the decision tree can be smaller, as shown by some of the demonstration data sets earlier in the paper.

VII. REDUCED DATA SETS

Work by Jensen and Shen on rough set theory aims to reduce the input set to a subset of attributes that have the same predictive value as the original set [20]. In this sense whereas the work reported here extends the input space by adding polynomials of the base features, Jensen's work reduces the input space. Using the glass data set as an example set of data Jensen's model removes the input value that measures the amount of Barium in the glass sample. Initial experiments show that the reduced set does reduce accuracy, but only marginally and not significantly, see Table IV.

TABLE IV
EXPERIMENTAL RESULTS FOR GLASS DATA SET COMPARING COMPLETE AND REDUCED DATA SETS WITH C4.5

	C4.5 Complete	C4.5 Reduced
Train Correct	92.8	92.3
Test Correct	75	70
DOF Train	30	30
DOF Test	25	25
probability of not null	1.0	1.0
Tree size	43	45

The interim conclusion is that no predictive accuracy has been lost; but it is also true that C4.5 alone does not extract everything from the data that it is possible to

extract. The next test evolves the reduced input space to extract as much predictive power as it can.

These preliminary tests indicate that reduction system of Jensen and Shen [20] does not remove useful information by eliminating input attributes, and coupled with the enhanced input space system reported here shows no loss of accuracy.

TABLE V
EXPERIMENTAL RESULTS FOR GLASS DATA SET COMPARING COMPLETE AND REDUCED DATA SETS WITH C4.5 AND C4.5 + GP

	C4.5 Complete	C4.5 Reduced	C4.5+GP Complete	C4.5+GP Reduced
Train Correct	92.8	92.3	98.5	99.0
Test Correct	75	70	100	100
DOF Train	30	30	30	30
DOF Test	25	25	25	25
probability of not null	1.0	1.0	1.0	1.0
Tree size	43	45	61	59

VIII. CONCLUSIONS

This paper has extended the capability of decision tree induction systems where the independent variables are continuous. The incremental decision process has been shown to be inadequate in explaining the structure of several sets of data without enhancement. The paper has shown that introducing variables based on higher order and higher dimensional combinations of the original variables can result in significantly better decision trees. This can all be accomplished by introducing these variables at the start of the decision tree generation and a suitable method for generating these would be a genetic algorithm. A fitness function for a genetic programming system has been introduced and serves to discover structure in the continuous domain. Although the work of [20] shows how to reduce the input set without losing any discriminating power they did not achieve all the predictive power that the input space could provide further work on a variety of different data sets should be performed to confirm this.

REFERENCES

- [1] J. Quinlan, "Discovering rules from large collections of examples: a case study," in *Expert systems in the microelectronic age*, D. Michie, Ed. Edinburgh University Press, 1979.
- [2] J. Quinlan, "Induction of decision trees," *Machine Learning*, vol. 1, no. 1, 1986.
- [3] Z. Fu, B. Golden, and S. Lele, "A GA based approach for building accurate decision trees," *INFORMS Journal on Computing*, vol. 15, no. 5, pp. 3–23, 2003.
- [4] M. Ryan and V. Rayward-Smith, "The evolution of decision trees," in *Proceedings of the Third Annual Conference on Genetic Programming*, J. Koza, Ed. San Francisco, CA.: Morgan Kaufmann, 1998, pp. 350–358.
- [5] R. Marmelstein and G. Lamont, "Pattern classification using a hybrid genetic program-decision tree approach," in *Proceedings of the Third Annual Conference on Genetic Programming*, J. Koza, Ed. San Francisco, CA 94104, USA: Morgan Kaufmann, 1998.
- [6] G. Folino, C. Puzzuti, and G. Spezzano, "Genetic programming and simulated annealing: a hybrid method to evolve decision trees," in *Proceedings of the Third*

European Conference on Genetic Programming, R. Poli, W. Banzhaf, W. Langdon, J. Miler, P. Nordin, and T. Fogarty, Eds. Edinburgh, Scotland, UK: Springer-Verlag, 2000, pp. 294–303.

- [7] D. Greene and S. Smith, "Competition-based induction of decision models from examples," *Machine Learning*, vol. 13, pp. 229–257, 1993.
- [8] A. Freitas, "A GA for generalised rule induction," in *Advances in Soft Computing, Engineering Design and Manufacturing*. Berlin: Springer, 1999, pp. 340–353.
- [9] D. Carvalho and A. Freitas, "A genetic-algorithm based solution for the problem of small conjuncts," in *Principles of Data Mining and Knowledge Discovery (Proc. 4th European Conf., PKDD-2000, Lyon France)*, ser. Lecture Notes in Artificial Intelligence, vol. 1910. Springer-Verlag, 2000, pp. 345–352.
- [10] K. De Jong, W. Spears, and D. Gordon, "Using a genetic algorithm for concept learning," *Machine Learning*, vol. 13, pp. 161–188, 1993.
- [11] C. Janikow, "A knowledge intensive GA for supervised learning," *Machine Learning*, vol. 13, pp. 189–228, 1993.
- [12] Y. Hu, "A genetic programming approach to constructive induction," in *Genetic Programming, Proceedings 3rd Annual Conference*, San Mateo, California, 1998, pp. 146–151.
- [13] A. Papagelis and D. Kalles, "GA tree: Genetically evolved decision trees," in *Tools with AI, ICTAI Proceedings 12th IEEE International Conference*, 13–15 Nov 2000, pp. 203–206.
- [14] J. Eggermont, J. Kok, and W. Kusters, "Genetic programming for data classification: Partitioning the search space," in *ACM Symposium on Applied Computing*, 2004, pp. 1001–1005.
- [15] D. Michie and R. Chambers, "Boxes: an experiment in adaptive control," in *Machine Intelligence 2*, E. Dale and D. Michie, Eds. Edinburgh: Oliver and Boyd, 1968.
- [16] A. Konstam, "Linear discriminant analysis using GA," in *Proceedings Symposium on Applied Computing*, Indianapolis, IN, 1993.
- [17] M. Withall, "Evolution of complete software systems," Ph.D. dissertation, Computer Science, Loughborough University, Loughborough, UK, 2003.
- [18] M. Withall, C. Hinde, and R. Stone, "An improved representation for evolving programs," *Genetic Programming and Evolvable Machines*, vol. 10, no. 1, pp. 37–70, 2009.
- [19] "Machine learning repository," <http://archive.ics.uci.edu/ml/datasets.html>, 2010.
- [20] R. Jensen and Q. Shen, "New approaches to fuzzy-rough feature selection," *IEEE Transactions on Fuzzy Systems*, vol. 17, no. 4, pp. 824–838.



Chris J. Hinde is Professor of Computational Intelligence in the Department of Computer Science at Loughborough University. His interests are in various areas of Computational Intelligence including fuzzy systems, evolutionary computing, neural networks and data mining. In particular he has been working on contradictory and inconsistent logics with a view to using them for data mining. A recently completed project was concerned with railway scheduling using an evolutionary system. He has been funded by various research bodies, including EPSRC, for most of his career and is a member of the EPSRC peer review college. Amongst other activities he has examined over 100 PhDs.



Anoud I. Bani-Hani is an EngD research Student at Loughborough University, researching into Knowledge Management in SMEs in the UK, with specific focus on implementing an ERP system into a low-tech SME. Prior to joining the EngD scheme Anoud was a Lecturer at Jordan University of Science and Technology and holds an undergraduate degree in Computer science and information technology system from the same university and a Master degree in Multimedia and Internet computing from Loughborough University.



Thomas W. Jackson - is a Senior Lecturer in the Department of Information Science at Loughborough University. Nicknamed 'Dr. Email' by the media Tom and his research team work in two main research areas, Electronic Communication and Information Retrieval within the Workplace, and Applied and Theory based Knowledge Management. He has published more than 70 papers in peer reviewed journals and conferences. He is on a number of editorial boards for international journals and reviews for many more. He has given a number of keynote talks throughout the world. In both research fields Tom has, and continues to work closely with both private and public sector organisations throughout the world and over the last few years he has brought in over £1M in research funding from research councils, including EPSRC. He is currently working on many research projects, including ones with The National Archives and the Welsh Assembly Government surrounding information management issues.



Yen P. Cheung has an honours degree in Data Processing from Loughborough University of Technology (UK) in 1986 and completed her doctorate in Engineering at the Warwick Manufacturing Group (WMG), University of Warwick in 1991. She worked initially as a Teaching Fellow and then as a Senior Teaching Fellow at WMG from 1988 1994 where she was responsible for the AI and IT courses in the M.Sc. programs in Engineering Business Management. She also ran IT courses for major companies such as British Aerospace, Royal Ordnance, London Electricity and British Airways in UK and for Universiti Teknologi Malaysia in Malaysia. Besides teaching, she also supervised a large number of industry based projects. After moving to Australia in 1994, she joined the former School of Business Systems at Monash University. Currently she is a Senior Lecturer at the Clayton School of IT at Monash University, Australia where she developed and delivered subjects in the area of business information systems, systems development, process design, modelling and simulation. Her current research interests are in the areas of collaborative networks such as eMarketplaces particularly for SMEs, intelligent algorithms for business systems, applications and data mining of social media. She publishes regularly in both international conferences and journals in these areas of research.

Call for Papers and Special Issues

Aims and Scope

Journal of Emerging Technologies in Web Intelligence (JETWI, ISSN 1798-0461) is a peer reviewed and indexed international journal, aims at gathering the latest advances of various topics in web intelligence and reporting how organizations can gain competitive advantages by applying the different emergent techniques in the real-world scenarios. Papers and studies which couple the intelligence techniques and theories with specific web technology problems are mainly targeted. Survey and tutorial articles that emphasize the research and application of web intelligence in a particular domain are also welcomed. These areas include, but are not limited to, the following:

- Web 3.0
- Enterprise Mashup
- Ambient Intelligence (Aml)
- Situational Applications
- Emerging Web-based Systems
- Ambient Awareness
- Ambient and Ubiquitous Learning
- Ambient Assisted Living
- Telepresence
- Lifelong Integrated Learning
- Smart Environments
- Web 2.0 and Social intelligence
- Context Aware Ubiquitous Computing
- Intelligent Brokers and Mediators
- Web Mining and Farming
- Wisdom Web
- Web Security
- Web Information Filtering and Access Control Models
- Web Services and Semantic Web
- Human-Web Interaction
- Web Technologies and Protocols
- Web Agents and Agent-based Systems
- Agent Self-organization, Learning, and Adaptation
- Agent-based Knowledge Discovery
- Agent-mediated Markets
- Knowledge Grid and Grid intelligence
- Knowledge Management, Networks, and Communities
- Agent Infrastructure and Architecture
- Agent-mediated Markets
- Cooperative Problem Solving
- Distributed Intelligence and Emergent Behavior
- Information Ecology
- Mediators and Middlewares
- Granular Computing for the Web
- Ontology Engineering
- Personalization Techniques
- Semantic Web
- Web based Support Systems
- Web based Information Retrieval Support Systems
- Web Services, Services Discovery & Composition
- Ubiquitous Imaging and Multimedia
- Wearable, Wireless and Mobile e-interfacing
- E-Applications
- Cloud Computing
- Web-Oriented Architectures

Special Issue Guidelines

Special issues feature specifically aimed and targeted topics of interest contributed by authors responding to a particular Call for Papers or by invitation, edited by guest editor(s). We encourage you to submit proposals for creating special issues in areas that are of interest to the Journal. Preference will be given to proposals that cover some unique aspect of the technology and ones that include subjects that are timely and useful to the readers of the Journal. A Special Issue is typically made of 10 to 15 papers, with each paper 8 to 12 pages of length.

The following information should be included as part of the proposal:

- Proposed title for the Special Issue
- Description of the topic area to be focused upon and justification
- Review process for the selection and rejection of papers.
- Name, contact, position, affiliation, and biography of the Guest Editor(s)
- List of potential reviewers
- Potential authors to the issue
- Tentative time-table for the call for papers and reviews

If a proposal is accepted, the guest editor will be responsible for:

- Preparing the “Call for Papers” to be included on the Journal’s Web site.
- Distribution of the Call for Papers broadly to various mailing lists and sites.
- Getting submissions, arranging review process, making decisions, and carrying out all correspondence with the authors. Authors should be informed the Instructions for Authors.
- Providing us the completed and approved final versions of the papers formatted in the Journal’s style, together with all authors’ contact information.
- Writing a one- or two-page introductory editorial to be published in the Special Issue.

Special Issue for a Conference/Workshop

A special issue for a Conference/Workshop is usually released in association with the committee members of the Conference/Workshop like general chairs and/or program chairs who are appointed as the Guest Editors of the Special Issue. Special Issue for a Conference/Workshop is typically made of 10 to 15 papers, with each paper 8 to 12 pages of length.

Guest Editors are involved in the following steps in guest-editing a Special Issue based on a Conference/Workshop:

- Selecting a Title for the Special Issue, e.g. “Special Issue: Selected Best Papers of XYZ Conference”.
- Sending us a formal “Letter of Intent” for the Special Issue.
- Creating a “Call for Papers” for the Special Issue, posting it on the conference web site, and publicizing it to the conference attendees. Information about the Journal and Academy Publisher can be included in the Call for Papers.
- Establishing criteria for paper selection/rejections. The papers can be nominated based on multiple criteria, e.g. rank in review process plus the evaluation from the Session Chairs and the feedback from the Conference attendees.
- Selecting and inviting submissions, arranging review process, making decisions, and carrying out all correspondence with the authors. Authors should be informed the Author Instructions. Usually, the Proceedings manuscripts should be expanded and enhanced.
- Providing us the completed and approved final versions of the papers formatted in the Journal’s style, together with all authors’ contact information.
- Writing a one- or two-page introductory editorial to be published in the Special Issue.

More information is available on the web site at <http://www.academpublisher.com/jetwi/>.

The Developing Economies' Banks Branches Operational Strategy in the Era of E-Banking: The Case of Jordan <i>Yazan Khalid Abed-Allah Migdadi</i>	189
Evolving Polynomials of the Inputs for Decision Tree Building <i>Chris J. Hinde, Anoud I. Bani-Hani, Thomas.W. Jackson, and Yen P. Cheung</i>	198
