

Abstracts of Recent PhDs

Analysis of the Dynamics of Cognitive Processes

Candidate: Tibor Bosse

Institution: Department of Artificial Intelligence, Vrije Universiteit Amsterdam, The Netherlands

Supervisors: J. Treur and C. M. Jonker

Year awarded: 2005

URL: <http://www.cs.vu.nl/res/theses/bosse.html>

Abstract

In this thesis a novel approach to Cognitive Modelling is introduced, which focuses on analysis and simulation of the dynamics of cognitive processes. Dynamic aspects play an important role in most cognitive processes. For example, within human reasoning processes such dynamic aspects are posing reasoning goals, making assumptions and evaluating assumptions. Such cognitive processes cannot be understood, justified

or explained without taking into account these dynamic aspects. Therefore, the main goal of this thesis is to introduce a novel approach for the analysis of the dynamics of cognitive processes, and to explore its applicability in a variety of cognitive domains. This approach makes use of a hybrid language called TTL, which is based on a combination of symbolic and mathematical constructs.

Efficient Human Pose Estimation from Real World Images

Candidate: Timothy J. Roberts

Institution: School of Computing, University of Dundee, Scotland, UK

Supervisors: Stephen J. McKenna and Ian W. Ricketts

Year awarded: 2005

URL: <http://www.computing.dundee.ac.uk/staff/troberts/>

Abstract

Reliable, efficient human pose estimation from images is a precursor to many useful applications including advanced human computer interfaces, surveillance systems, image archive analysis and smart environments. Whilst progress has been made on human pose estimation, the research often makes strong assumptions about the appearance. In particular, assumptions are often made regarding the background, foreground, self and other object occlusion and number of viewpoints. These assumptions limit the application of computer human pose estimation systems. In contrast, the focus of this thesis is pose estimation from single real world images or monocular sequences of poorly constrained scenes. Furthermore, it aims to accomplish this efficiently. The body of the thesis is structured into three layers: formulation, likelihood and estimation.

First, the popular, generative, part-based approach is extended to allow pose hypotheses that have different numbers of parts to be compared. This *partial configuration* formulation allows pose estimation in the presence of other object occlusion, enables efficient estimation and automatic (re)initialisation and gives robustness to body parts with a non-contrasting or poorly modelled appearance. The problem of comparing partial configurations is stated as a Bayesian decision problem of discriminating between the class of people and of backgrounds. To describe the body part shape, a probabilistic model is learnt from manually segmented

and aligned training data of multiple subjects in various poses. In order to obtain a low-dimensional model, variations due to intra-person differences and clothing as well as difficult to observe degrees of freedom and differences between certain similar body parts are marginalised over. The resulting model allows uncertainty in measurements to be quantified as well as improving estimation efficiency. Finally, a prior is developed to encode inter-part constraints and it is shown that due to these constraints smaller configurations contain much of the information of larger configurations.

Although a strong likelihood model is critical in determining the success of human pose estimation many existing models have limitations in terms of discrimination and efficiency when applied to real world images. Therefore, two novel techniques are developed to discriminate people with complex, textured appearance from cluttered backgrounds. A boundary model is developed based upon the divergence between the appearance distribution of the foreground region and its adjacent background. In particular, the distribution of the divergence between the joint colour histograms of these regions is learnt for correct and incorrect configurations. In order to provide a quantitative empirical evaluation the statistics of intensity edges on and off human boundaries are also learnt. It is shown that the new boundary model is more discriminatory and searchable. This is particularly important as early

identification of body parts focuses the estimation. Next, a model is proposed that encodes the spatial structure of human appearance. In particular, the statistics of the similarity between regions on the surface of correct and incorrect configurations are learnt. Encoding inter-part similarity is important in discriminating larger incorrect configurations, and due to the combinatorial growth in the number of large configurations, is key to efficient estimation in real world images. In addition to these likelihood models a foreground model is developed that encodes the expectation of temporal consistency in appearance for use in human tracking applications. It builds upon previous techniques by matching feature distributions and using clothing structure to improve estimation of the adapting foreground appearance.

Once the model and likelihood have been defined pose estimation can be performed. Two approaches to pose

estimation can be identified in the literature. The combinatorial approach identifies body part candidates in the image and then combines the results, for example using dynamic programming, to estimate the overall body pose. Whilst such methods are efficient they rely heavily upon body part detection, which is particularly difficult in the presence of occlusion and clutter. In contrast, the full state space approach searches for whole body configurations and thereby models the complex self-occluding appearance. However, due to the high dimensional space such methods use local rather than global sampling and require manual initialisation. By taking advantage of the *partial configuration* formulation and the strong likelihood model a straightforward deterministic search algorithm is able to recover many of the body parts and results of such a search to challenging scenes are presented.

Adaptive Learning Algorithms for Bayesian Network Classifiers

Candidate: Gladys Castillo

Institution: Department of Mathematics, University of Aveiro, Portugal

Supervisors: João Gama and Ana M. Breda

Year awarded: 2006

URL: <http://www2.mat.ua.pt/gladys/publications.htm>

Abstract

This thesis is concerned with adaptive learning algorithms for Bayesian network classifiers (BNCs) in a prequential (on-line) learning scenario. An efficient supervised learning algorithm in such a scenario must be able to improve its predictive accuracy while optimizing the cost of updating. However, if the process is not strictly stationary, the target concept could change over time. Hence, the predictive model should be adapted quickly to these changes. We integrated all the adaptive algorithms into the adaptive prequential framework for supervised learning, AdPreqFr4SL, which is aimed at handling the cost-performance trade-off and concept drift. The adaptive strategy is based upon bias management and gradual adaptation. Instead of choosing a particular class of BNCs (e.g. Naive Bayes (NB), TAN, etc.) and use it during all the learning process, we propose to use the Sahami's class of k -Dependence Bayesian Classifiers and start with its simpler class-model, the NB, by setting k

(the maximum number of allowable attribute dependencies) to 0. We then attempt to reduce the bias of the NB by gradually increasing the k value and then searching for new attribute dependences in the extended search space. This bias control leads to the selection of the optimal class-model for the current amount of training data (the optimal k value), thus avoiding the problems caused by either too much bias (underfitting) or too much variance (overfitting). Since updating the structure is a costly task, we use new data to primarily adapt the parameters. We adapt the structure only when is actually necessary. The method for handling concept drift is based on the Shewhart P-Chart. If during the monitoring process a concept drift is detected, the classifier is adapted accordingly to these changes. We evaluated our algorithms on artificial domains and benchmark problems and show their advantages and future applicability in real-world on-line learning systems.

Enhancing Semantic Web Data Access

Candidate: Li Ding

Institution: Department of Computer Science and Electronic Engineering, University of Maryland, Baltimore County, MD, USA

Supervisor: Tim Finin

Year awarded: 2006

URL: <http://www.ksl.stanford.edu/people/ding/>

Abstract

The Semantic Web was invented by Tim Berners-Lee in 1998 as a web of data for machine consumption. Its applicability in supporting real world applications on the World Wide Web, however, remains unclear to this day because most existing works treat the Semantic Web as one universal RDF graph and ignore the Web aspect. In fact, the Semantic Web is distributed on the Web as a web of belief: each piece of Semantic Web data is independently published on the Web as a certain agent's belief instead of the universal truth. Therefore, we enhance the current conceptual model of the Semantic Web to characterize both the content and the context of

Semantic Web data. A significant sample dataset is harvested to demonstrate the non-trivial presence and the global properties of the Semantic Web on the Web. Based on the enhanced conceptual model, we introduce a novel search and navigation model for the unique behaviors in Web-scale Semantic Web data access, and develop an enabling tool—the Swoogle Semantic Web search engine. To evaluate the data quality of Semantic Web data, we also (i) develop an explainable ranking schema that orders the popularity of Semantic Web documents and terms and (ii) introduce a new level of granularity of Semantic Web data—RDF molecule that supports

lossless RDF graph decomposition and effective provenance tracking. This dissertation systematically investigates the Web aspect of the Semantic Web. Its primary contributions

are the enhanced conceptual model of the Semantic Web, the novel Semantic Web search and navigation model, and the Swoogle Semantic Web search engine.

Conditional log-likelihood MDL and Evolutionary MCMC

Candidate: Madalina-Mihaela Drugan

Institution: Computer Science Department, Utrecht University, The Netherlands

Supervisors: Linda C. van der Gaag and Dirk Thierens

Year awarded: 2006

URL: <http://igitur-archive.library.uu.nl/dissertations/2006-1130-200351/index.htm>

Abstract

In the current society there is an increasing interest in intelligent techniques that can automatically process, analyze and summarize the ever-growing amount of data. Artificial intelligence is a research field that studies intelligent algorithms to support people in making decisions. Algorithms that are able to induce knowledge from examples are researched in the field of machine learning. This thesis studies improvements of particular machine-learning algorithms. In the first part of this thesis we describe methods that are able to select useful attributes (or features) that can be used as inputs by a classification algorithm. We focus on selecting relevant attributes that should be used as inputs for a Bayesian network classifier. For our goal to construct selective Bayesian network classifiers, we propose and investigate a score function that can evaluate Bayesian network classifiers and that indicates the simplest and the most performant classifier. We theoretically and experimentally show that our proposed conditional

log-likelihood minimum description length function, MDL-FS, is well suited for constructing simple and well-performing Bayesian network classifiers. In the second part of the thesis we integrate some methods from evolutionary computation into a Markov chain Monte Carlo (MCMC) sampler. Sampling is related to optimization, but whereas in optimization we are only interested in the state with the highest fitness, in sampling we are interested in the overall probability distribution over states. To improve MCMC methods that are often used for sampling, we investigate the Evolutionary MCMC (EMCMC) framework, where population-based MCMCs exchange information between the individual states such that they are still MCMCs at population level. We investigate and propose various evolutionary techniques (e.g. recombination, selection), which we then integrate in the EMCMC framework. We experimentally show that our proposed EMCMCs can outperform the standard MCMC algorithms.

Dimensão Topológica e Mapas Auto Organizáveis de Kohonen (English translation: Topological Dimension and Kohonen's Self Organizing Maps)

Candidate: Sarajane Marques Peres

Institution: Department of Computer Engineering and Industrial Automation, School of Electrical and Computer Engineering, State University of Campinas, Brazil

Supervisor: Márcio Luiz de Andrade Netto

Year awarded: 2006

URL: www.usp.leste.usp.br/sarajane/SubPaginas/arquivos_formacao/TeseSMPeres.zip

Abstract

Kohonen's Self Organizing Map (SOM) is a tool for data analysis, used mainly for clustering analysis and dimension reduction. Although both research and use of this tool have already reached high levels of maturity, some questions remain unresolved. One of the difficulties concerned to the SOM is the specification of its parameters. There are some heuristics related to the decision process about: number of neurons, lattice on the topological output space, neighborhood function, weight initializing and stopping criteria to the learning process. However, the decision about the SOM's output space topological dimension is an opened question. Although empirical studies show that the two-dimensional space is enough, we show the utility of adjusting this parameter in accordance with each application.

In order to reach this goal, we have defined a new way to determine the output space topological dimension using knowledge about the explored dataset. This analysis

is carried out with the combined support of the Fractal Theory and the Fuzzy Approximated Reasoning, deriving a new fractal dimension measure: the Meaningful Fractal Fuzzy Dimension (named DFFS, from the name in Portuguese). The DFFS can provide a notion about the space portion filled by the dataset clustering distribution and hence infer the dimension where the datapoint topological relations are established. Therefore, it can be used as a support for the decision about the SOM output space topological dimension, which allows the reduction of topological distortions in specific cases.

The determination process of this new measure and its application as an inference to the SOM conception, have been both validated in this work. The former has been carried out through its application to the Clustering Tendency Analysis and the latter through the topological quality analysis of the SOM designed with such inference.

Learning Probabilistic Networks in Large and Structured Domains

Candidate: Massimiliano Mascherini

Institution: Department of Statistics, University of Florence, Italy

Supervisor: F. M. Stefanini

Year awarded: 2006

URL: <http://statind.jrc.it/mastino>

Abstract

The dissertation contributes to the development of the processes for learning Bayesian Networks in large and structured domain by focusing on the 'score + search' approach and presenting a new metric and a new search algorithm to discover the 'best' network. The manuscript is organized in four papers plus an appendix. Each paper is self-contained with bibliography, figures, tables and equation numbering. The original methods developed have been implemented in the R environment, and the user's manual for the suite of R functions created is included in the appendix.

The analysis takes its moves from the problem of extracting information from huge amounts of data, such as DNA microarray experiments for building genetic networks. The first step in the proposed learning process involves scoring the network, and this issue was addressed by developing a new quasi-Bayesian metric, the P-metric, which encodes also the prior information on structures. By exploiting weak prior knowledge on node ordering

and on graph topology, the P-metric allows for the elicitation of prior distribution on Direct Acyclic Graphs (DAGs). This feature is particularly useful in large and structured domains as the DNA microarray area, where the expert knowledge domain can be rich but not complete, and eliciting this information can noticeably improve the overall performance of the structural learning process. Then, the problem of the search of optimal Bayesian Network (BN) from a database of observations was analyzed proposing a new population-based algorithm, the M-GA. The algorithm permits to learn the structure of BNs without assuming any ordering of nodes and allowing for the presence of both discrete and continuous random variables. These new methods were combined and implemented in MASTINO, a suite of R functions written in R and coded on the top of the package DEAL. The package MASTINO can be downloaded from the Webpage of the author and may be used under the terms of the General Public License v2 (GNU).

An Exploratory Analysis of a Large Health Cohort Study Using Bayesian Networks

Candidate: Delin Shen

Institution: The Harvard-MIT Division of Health Sciences and Technology, Massachusetts Institute of Technology and Harvard Medical School, Cambridge, MA, USA

Supervisor: Peter Szolovits

Year awarded: 2006

Abstract

Large health cohort studies are among the most effective ways in studying the causes, treatments, and outcomes of diseases by systematically collecting a wide range of data over long periods. The wealth of data in such studies may yield important results in addition to the already numerous findings, especially when subjected to newer analytical methods. Bayesian Networks (BN) provide a relatively new method of representing uncertain relationships among variables, using the tools of probability and graph theory, and have been widely used in analyzing dependencies and the interplay between variables. We used BN to perform an exploratory analysis on a rich collection of data from one large health cohort study, the Nurses' Health Study (NHS), with the focus on breast cancer. We explored the data from the NHS using BN to look for breast cancer risk factors, including a group of Single Nucleotide Polymorphisms (SNP).

We found no association between the SNPs and breast cancer, but found a dependency between clomid and breast cancer. We evaluated clomid as a potential risk

factor after matching on age and number of children. Our results showed, for clomid, an increased risk for estrogen receptor positive cancer (odds ratio 1.52, 95 per cent CI 1.11–2.09) and a decreased risk of estrogen receptor negative breast cancer (odds ratio 0.46, 95 per cent CI 0.22–0.97). We developed breast cancer risk models using BN. We trained models on 75 per cent of the data, and evaluated them on the remaining. Because of the clinical importance of predicting risks for Estrogen Receptor positive and Progesterone Receptor positive breast cancer, we focused on this specific type of breast cancer to predict two-year, four-year, and six-year risks. The concordance statistics of the prediction results on test sets are 0.70 (95 per cent CI 0.67–0.74), 0.68 (95 per cent CI 0.65–0.72), and 0.66 (95 per cent CI 0.62–0.69) for two-, four-, and six-year models, respectively. We also evaluated the calibration performance of the models, and applied a filter to the output to improve the linear relationship between predicted and observed risks using Agglomerative Information Bottleneck clustering without sacrificing much discrimination performance.

Designing Invisible Handcuffs. Formal Investigations in Institutions and Organizations for Multi-agent Systems

Candidate: Davide Grossi

Institution: Institute of Information and Computing Sciences, Universiteit Utrecht, The Netherlands

Supervisors: John-Jules Ch. Meyer and Frank Dignum

Year awarded: 2007

URL: <http://www.cs.uu.nl/staff/davide.html>

Abstract

In human societies, institutions and organizations have developed as means for arranging social interaction in such a way that, on the one hand, agents remain autonomous and, on the other hand, that the society as a whole exhibits desirable properties. They set boundaries—handcuffs—for the activities of the individuals in the society but, like Smith's notorious 'invisible hand', they are something which just cannot be pointed at in the outside world. If invisible handcuffs need to be designed in order to coordinate software agents, then a formal theory of institutions and organizations needs to be found which can ground such design process. But if we aim at a formal theory of institutions and organizations, how should we think of them?

Or, to carry on the metaphor, how can we make them 'visible'?

The thesis conceives of institutions as systems of constitutive rules. Following Searle, constitutive rules are

statements of the type 'X counts as Y in context C' (counts-as statements) and they underlie the whole construction of institutional reality. Obviously, many different institutions coexist which might disagree on the way they look at the same domain. This motivates the formal analysis of the notion of context and the related one of contextual terminology. In a nutshell, institutions are viewed as contextual terminologies, and counts-as statements as their basic building blocks. As to organizations, the thesis focuses on their structural dimension analyzing them as multi-graphs, to which properties of transition systems, expressed in logical formulae, are attached, which specify the impact of the structure on the activities of the agents in the organization. This perspective provides both quantitative methods, based on graph-theory, to compare different organizations from a structural point of view, and qualitative ones, based on logic, to address the types of interaction that different organizations put in place.

Modeling of Change in Multi-Agent Organizations

Candidate: Mark Hoogendoorn

Institution: Department of Artificial Intelligence, Vrije Universiteit Amsterdam, The Netherlands

Supervisors: Catholijn M. Jonker and Jan Treur

Year awarded: 2007

URL: <http://www.cs.vu.nl/mhoogen>

Abstract

Organizational changes occur in a wide variety, from reorganizations within companies to reduce their costs, to disaster prevention organizations that are created to save human lives. In all cases, it is essential that the success rate of such changes is optimal. During organizational changes however, often the goals set for the change are not reached. Research has shown that this holds for over 70 per cent of the changes within companies. Over the last decades, techniques within the field of computational and mathematical organization theory have been developed to describe organizations in a formal way. In this thesis, such a formal organizational modeling approach,

addressing both the structure as well as the behavior of an organization, is used to model change processes. Such models can be used by organizational change experts to optimize change processes, and improve their success rate. Three main parts within organizational change processes are addressed in this thesis. First of all, the process of analyzing the current organizational situation and creation of improvements thereof is modeled. Furthermore, the process of moving from the current organization to this new organization is addressed. Finally, approaches to evaluate the effectiveness of such changes are presented.

Algorithms for Coalition Formation in Multi-Agent Systems

Candidate: Talal Rahwan

Institution: School of Electronics and Computer Science, University of Southampton, UK

Supervisor: Nicholas R. Jennings

Year awarded: 2007

URL: <http://www.ecs.soton.ac.uk/people/tr03r>

Abstract

Coalition formation is a fundamental form of interaction that allows the creation of coherent groupings of distinct, autonomous agents in order to efficiently achieve their individual or collective goals. Forming effective coalitions is a major research challenge in the field of multi-agent systems. Central to this endeavour is the problem of determining which of the possible coalitions to form in order to achieve some goal. This usually requires calculating a value for every possible coalition, known as the coalition value, which indicates how beneficial that coalition would be if it were formed. Now since the number of possible coalitions grows exponentially with the number of agents involved, then, instead of having a single agent calculate all these values, it would be more efficient to distribute this calculation among all agents, thus, exploiting all computational resources that are

available to the system, and preventing the existence of a single point of failure.

Against this background, we develop a novel algorithm for distributing the value calculation among the cooperative agents. Specifically, by using our algorithm, each agent is assigned some part of the calculation such that the agents' shares are exhaustive and disjoint. Moreover, the algorithm is decentralized, requires no communication between the agents, has minimal memory requirements, and can reflect variations in the computational speeds of the agents. To evaluate the effectiveness of our algorithm we compare it with the only other algorithm available in the literature for distributing the coalitional value calculations (due to Shehory and Kraus). This shows that for the case of 25 agents, the distribution process of our algorithm took less than

0.02 per cent of the time, the values were calculated using 0.000006 per cent of the memory, the calculation redundancy was reduced from 383 229 848 to 0, and the total number of bytes sent between the agents dropped from 1146 989 648 to 0. Note that for larger numbers of agents, these improvements become exponentially better.

Once the coalitional values are calculated, the agents usually need to find a combination of coalitions in which every agent belongs to exactly one coalition, and by which the overall outcome of the system is maximized. This problem, which is widely known as the coalition structure generation problem, is extremely challenging due to the number of possible combinations which grows very quickly as the number of agents increases, making it impossible to go through the entire search space, even for small numbers of agents. Given this, many algorithms have been proposed to solve this problem using different techniques, ranging from dynamic programming, to integer programming, to stochastic search, all of which suffer from major limitations relating to execution time, solution quality, and memory requirements.

With this in mind, we develop a novel, anytime algorithm for solving the coalition structure generation problem. Specifically, the algorithm can generate solutions by partitioning the space of all potential coalition structures into sub-spaces containing coalition structures that are similar, according to some criterion, such that these sub-spaces can be pruned by identifying their bounds. Using this

representation, the algorithm can then search through the selected sub-space(s) very efficiently using a branch-and-bound technique. We empirically show that we are able to find solutions that are optimal in 0.082 per cent of the time required by the fastest available algorithm in the literature (for 27 agents), and that is using only 33 per cent of the memory required by that algorithm. Moreover, our algorithm is the first to be able to solve the coalition structure generation problem for numbers of agents bigger than 27, in reasonable time (less than 90 minutes for 30 agents as opposed to around 2 months for the current state of the art). The algorithm is anytime, and if interrupted before it would have normally terminated, it can still provide a solution that is guaranteed to be within a bound from the optimal one. Moreover, the guarantees we provide on the quality of the solution are significantly better than those provided by the previous state of the art algorithms designed for this purpose. For example, given 21 agents, and after only 0.0000002 per cent of the search space has been searched, our algorithm usually guarantees that the solution quality is no worse than 91 per cent of optimal value, while previous algorithms only guarantees 9.52 per cent. Moreover, our guarantee usually reaches 100 per cent after 0.0000019 per cent of the space has been searched, while the guarantee provided by other algorithms can never go beyond 50 per cent until the whole space has been searched. Again note that these improvements become exponentially better given larger numbers of agents.

Answer-Set Programming for the Semantic Web

Candidate: Roman Schindlauer

Institution: Knowledge-Based Systems Group 184/3, Institute of Information Systems, Vienna University of Technology, Vienna, Austria

Supervisor: Thomas Eiter

Year awarded: 2007

URL: <http://www.kr.tuwien.ac.at/staff/roman/>

Abstract

This thesis makes a contribution to the research efforts of integrating rule-based inference methods with current knowledge representation formalisms in the Semantic Web. Ontology languages such as OWL and RDF Schema seem to be widely accepted and successfully used for semantically enriching knowledge on the Web and thus prepare it for machine-readability. However, these languages are of restricted expressivity if it comes to inferring new from existing knowledge. On the other side, rule formalisms have a long tradition in logic programming, being a common and intuitive tool for problem specifications. It is evident that the Semantic Web needs a powerful rule language complementing its ontology formalisms in order to facilitate sophisticated reasoning tasks.

Answer-set programming (ASP) is one of the most prominent and successful semantics for non-monotonic logic programs. The specific treatment of default negation under ASP allows for the generation of multiple models for a single program, which in this respect can be seen as the encoding of a problem specification. Highly efficient reasoners for ASP are available, each extending the core language by various sophisticated features such as aggregates or weak constraints.

In the first part of this thesis, we propose a combination of logic programming under the answer-set semantics with the description logics *SHIF(D)* and *SHOIN(D)*, which underlie the Web ontology languages OWL Lite and OWL DL, respectively. This combination allows for building rules on top of ontologies but also, to a limited extent, building ontologies on top of rules. We introduce *description logic programs (dl-programs)*, which

consist of a description logic knowledge base *L* and a finite set of *description logic rules (dl-rules)* *P*. Such rules are similar to usual rules in logic programs with negation as failure, but may also contain *queries to L*, possibly default-negated, in their bodies. We give a precise picture of the complexity of deciding answer-set existence for a dl-program, and of brave, cautious, and well-founded reasoning. We lay out possible applications of dl-programs and present the implementation of a prototype reasoner.

In the second part of the thesis, we generalize our approach to *hex-programs*, which are non-monotonic logic programs under the answer-set semantics admitting *higher-order atoms* as well as *external atoms*. Higher-order features are widely acknowledged as useful for performing meta-reasoning, among other tasks. Furthermore, the possibility to exchange knowledge with external sources in a fully declarative framework, such as ASP, is particularly important in view of applications in the Semantic Web area. Through external atoms, hex-programs can model some important extensions to ASP, and are a useful KR tool for expressing various applications. We define syntax and semantics of hex-programs and show how they can be deployed in the context of the Semantic Web. We give a picture of the computation method of hex-programs based on the theory of splitting sets, followed by a discussion on complexity. Then, the implementation of a prototype reasoner for hex-programs is outlined, along with a description how to extend this framework by custom modules. Eventually, we show the usability and versatility of hex-programs and our prototype implementation on the basis of concrete, real-world scenarios.
