

Multi-Agent Fault Tolerance Inspired by a Computational Analysis of Cancer

Megan Olsen

Department of Computer Science
University of Massachusetts Amherst
Amherst, MA 01003

Overview

My thesis investigates fault tolerance for cooperative agent systems that have some equivalent of self-replication and self-death. Utilizing biologically-inspired mechanisms, I increase multi-agent system robustness for faulty agents when it is unknown exactly which agent is malfunctioning. It is important to determine new ways to increase robustness of a system, as otherwise it cannot be guaranteed to function in all situations and thus cannot be relied upon. Robustness of a system allows agents to recover from errors and thus function continuously, an increasingly important trait as agent systems are deployed in real world scenarios such as sensor networks or surveillance systems where faulty or malicious nodes could disrupt application performance. To achieve robustness, there must either be prevention of all errors, or a technique for recovering from errors after they have occurred. My thesis creates a new fault tolerance mechanism inspired by cancer biology to remove faulty agents, and then re-applies the developed technique to study the removal of biological cancer cells in simulation.

For a multi-agent system to function continuously it must adapt on-line to failures. The problem must be diagnosed, and a plan must be provided to react to the problem (Hamscher, Console, and de Kleer 1992). Diagnosis for pre- and post-failure analysis for causal tasks can allow the system to both prevent a failure and recover from it. It is argued that post-failure protocols are less domain dependent and thus more crucial for the design of robust systems (Toyama and Hager 1997). My work shows that it is possible to fix the problem without first diagnosing exactly which agents are malfunctioning, thus removing the need to decide who is wrong before reacting.

Two main approaches for dealing with agent failures are Survivalist and Citizen. The survivalist approach requires each agent to be capable of dealing with all problems as an individual following a prepared set of actions for each specific problem (Marin et al. 2001). The citizen approach utilizes an external system that is alerted when an agent dies and then reallocates tasks so that the overall system continues to function correctly (Klein, Rodriguez-Aguilar, and Dellarocas 2003).

My approach is a combination of these two techniques: it does not require all agents to deal with failures individually, but instead utilizes a systemic approach (Olsen, Siegelmann-Danieli, and Siegelmann 2008). Unlike the citizen approach that requires special monitor nodes to diagnose failures, my system has each agent monitor its neighbors to detect and eliminate anomalies. I accomplish this task by implementing a message passing paradigm inspired by inter-cellular communication within tissues. The first type of message (PLEASE DIE) is sent by a single agent to all neighboring agents when it notices a potential error in its area. Once any agent receives enough of these messages, based on proximity to sender and an internal threshold amount, it will choose to die. Before dying the agent will send the I'M DYING message to all of its neighbors, which also eventually causes death based on an internal threshold level. By using this combination of messages, we remove all malfunctioning agents in the system. Although some correctly functioning agents may also die, they may be replaced by the self-replication that is built into the system. In practice, the concept of an agent "replicating" or "dying" may take many forms including cloning, restarting a hardware node, re-installing software on a node, etc.

My thesis examines how agent robustness can be accomplished using these apparently simple mechanisms, as well as how they can be applied back to the biological fields that originally served as inspiration (Olsen, Siegelmann-Danieli, and Siegelmann 2009). The agents in my system can be viewed as tissue cells, and the faulty agents are akin to cancer cells. I thus also propose that my basic robustness mechanisms could represent how normal tissue cells are able to remove cancer cells without cancer treatment, as has been witnessed in human cancer. My work thus applies not only to multi-agent computer systems, but also to increasing understanding of cancer development in biological tissues.

Similarly, these mechanisms may relate to neurodegenerative diseases. I plan to test this theory on a neural network by implementing similar communication protocols. Success would indicate a new technique for use in neural networks, as well as a broader range of applicability in biological systems. I thus have multiple goals for my thesis:

- Develop a robustness mechanism for multi-agent systems in which agents may replicate and die

- Determine how these mechanisms provide insights into the body's natural cancer defenses or potential treatments
- Apply the robustness mechanism in terms of neurodegenerative diseases to investigate the ability to utilize the mechanisms in neural networks

Results So Far

I have primarily tested this system using a simulation of self-replicating, self-dying, and self-organizing agents in a 3D world. Correctly functioning agents will only be created in the neighborhood of another correctly functioning agent if there is room for it, and will only occur when the probability of replication allows it. Malfunctioning agents will however be created at a more frequent rate once a single agent has begun to malfunction. The system is tested starting at the point where an agent begins to malfunction.

I have tested my system with hundreds of agent parameters in combination with hundreds of robustness protocol parameter sets. I focus on worst case scenario parameters, in which faulty agents malfunction in such a way as to increase their likelihood of overrunning the system. I consider success to be represented by removing all malfunctioning agents without removing the majority of correctly functioning agents in the majority of experiments run for a specific set of parameters. My most recent results (not yet published) show that in the majority of cases we are able to remove the malfunctioning agents completely without also removing a majority of the correctly functioning agents. This can be accomplished whether the rate of malfunctioning agent creation is close to the rate of normal agent creation or significantly higher. I include in my experiments parameter sets in which malfunctioning agents need to receive many more PLEASE DIE and I'M DYING messages to die. This demonstrates that my robustness mechanism is able to effectively prevent system failure in a range of scenarios, even when faulty agents are partially resistant to the recovery protocols.

In addition, I have incorporated a possibility of the robustness protocols failing. In these cases, either malfunctioning agents may not always send I'M DYING messages before they die from the protocols, or they may sometimes ignore received messages. As seen in Figure 1, even when malfunctioning agents fail to acknowledge my robustness mechanisms the majority of the time, my mechanisms can still succeed in removing them from the system before all normal agents are also destroyed.

I have also analyzed the basic system mathematically using differential equations, as well as from the view of a cellular automata. These forms of analysis are currently most useful for the analogous cancer model, and we have found similar promising results when using biologically plausible data with our robustness mechanisms. It is possible that the mechanisms we propose will help cancer researchers understand related biological phenomenon that are currently unexplained.

Work Plan

I plan to continue the multi-agent robustness investigation, as well as increasing the work on the biological analogies

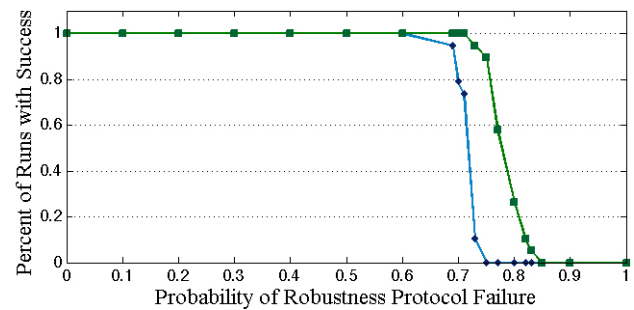


Figure 1: When agents may ignore robustness messages (diamond) or fail to send I'M DYING messages before dying (square), we still see high success rates. For up to 60% failure rates for either failure, we still see 100% success in systems where the correctly functioning robustness protocol also succeeded 100% of the time. The system cannot handle either failure occurring more than 80% of the time. However, these results show that my robustness mechanism is robust to failures within itself.

to be applied to human cells in regards to cancer and neurodegenerative diseases:

1. Add additional biological details such as blood vessels to increase the biological plausibility of the cancer model.
2. Continue developing a mechanism that can be easily applied to many computer systems.
3. Develop the neural network model of neuro-degenerative diseases. This project is a focus of the spring 2010 semester, and should be finished by the DC.

References

- Hamscher, W.; Console, L.; and de Kleer, J. 1992. *Readings in Model-Based Diagnosis*. Morgan Kaufmann Publishers Inc.
- Klein, M.; Rodriguez-Aguilar, J.; and Dellarocas, C. 2003. Using domain-independent exception handling services to enable robust open multi-agent systems: The case of agent death. *Autonomous Agents and Multi-Agent Systems* 7:179–189.
- Marin, O.; Sens, P.; Briot, J.; Guessoum, Z.; and Cedex, B. L. H. 2001. Towards adaptive fault tolerance for distributed multi-agent systems. In *Proceedings of European Research Seminar on Advances in Distributed Systems*.
- Olsen, M.; Siegelmann-Danieli, N.; and Siegelmann, H. 2008. Robust artificial life via artificial programmed death. *Artificial Intelligence* 172:884–898.
- Olsen, M.; Siegelmann-Danieli, N.; and Siegelmann, H. 2009. Computational modeling reveals the crucial role of cellular citizenship in selective tumor apoptosis. In *Systems Biology and Human Disease*.
- Toyama, K., and Hager, G. D. 1997. If at first you don't succeed... In *Proceedings of the 14th National Conference on Artificial Intelligence*.