

MacGyver challenges Turing with a new benchmark for 'creativity' in AI

Originally published June 08, 2017 at <https://www.intelligentautomation.network/artificial-intelligence/articles/macgyver-challenges-turing-with-a-new-benchmark>



What can the example of a 1980s TV super-agent offer the field of developmental AI research..?

If there is a quality that defines the border between scripting/branching and a more visionary definition of artificial intelligence, 'creativity' would seem to qualify – as anyone can attest who has ever been stuck in the cul-de-sac of three or four non-applicable options in an automated customer service system. Is such a system genuinely uninterested in non-templated problems, or simply too unsophisticated to think laterally?

In a [new paper](#), researchers from Tufts University in Massachusetts identify the quality of creative problem-solving as a cardinal virtue for machine intelligence, and - with due respect to Alan Turing - propose a new human avatar to represent it: geek super-agent MacGyver, as portrayed by Richard Dean Anderson in the 1980s TV show which ran for seven seasons on America's ABC network.

MacGyver's unlikely ability to defy fate with [paperclips](#), string, or other handy debris forms the basis of a new evaluation framework from research leads Vasanth Sarathy and Matthias Scheutz: the MacGyver Test (MT), which Sarathy and Scheutz intend 'to answer the question whether embodied machines can generate, execute and learn strategies for identifying and solving seemingly unsolvable real-world problems.'

The MacGyver Test presents an agent with a challenge that cannot be resolved by pre-existing knowledge or subsequent reference to analogous situations, such as case studies in data set resources. The test proposes criteria for evaluating if the machine's behaviour can be defined as genuinely creative, suggesting the resourcefulness [displayed](#) by NASA, its outsourcers and the Apollo 13 astronauts as an acme for this capacity to 'make do and mend'.

Learning to succeed - or give up

'[If] the agent can think outside of its current context, take some exploratory actions, and incorporate relevant environmental cues and learned knowledge to make the problem tractable (or at least computable) then the agent has the general ability to solve open-world problems more effectively.'

The researchers define the type of challenge suitable for the test as a 'MacGyver Problem' (MGP). The AI entity cannot know in advance whether the challenge facing it is actually this type of problem, making this classification part of its (presumably) tentative experiments out of its base domain and customary parameters.

Practical examples envisioned by the researchers include tightening a screw in the absence of a screwdriver (but in the presence of other objects, only some of which might be made to accommodate the task, alone or in combination) and moving both a block and a towel from one location to another (when each movement potentially causes conflict in the task end-result).

In the case of creating a makeshift screwdriver out of a dime, the machine needs particularly lateral creativity, since it is also supplied with a pair of pliers which, though unable to tighten the screw effectively by itself, could grip the dime and complete the task.

In some proposed cases the MT might involve a scenario that is actually not resolvable with the materials supplied, with the AI system under test expected to iterate through all possible combinations and finally exit the exploratory loop and accept defeat - internalising and cataloguing the non-effective procedures in its database for reference in future challenges.

The MacGyver Test is intended for sophisticated machine entities which share the human capability to run heuristics on an environment and filter out pragmatic approaches from the 'noise' of choice in a diverse and detailed scenario.

'Finally,' the researchers add 'the agent must also be able to remember this knowledge and be able to, more efficiently, solve future instances of similar problems.'

The work is intended to build on the various evaluative schemas which have evolved from the original Turing Test, such as Steven Harnad's [Total Turing Test](#), and Paul Schweizer's [Truly Total Turing Test](#), both of which attempt to move the field on from the notion of AI as effective simulacra,

towards the creation of an entity that is genuinely evolutionary – a concept that remains the subject of cultural, political, legal and scientific apprehension, but which perhaps should best be proven before it is feared.

Illustration: [Wikimedia Commons](#)
