

Research Article

AI-Complete CAPTCHAs as Zero Knowledge Proofs of Access to an Artificially Intelligent System

Roman V. Yampolskiy

Computer Engineering and Computer Science, University of Louisville, Louisville, KY 40292, USA

Correspondence should be addressed to Roman V. Yampolskiy, rvyamp01@louisville.edu

Received 12 August 2011; Accepted 10 September 2011

Academic Editors: Y. Liu and P. Trunfio

Copyright © 2012 Roman V. Yampolskiy. This is an open access article distributed under the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

Experts predict that in the next 10 to 100 years scientists will succeed in creating human-level artificial general intelligence. While it is most likely that this task will be accomplished by a government agency or a large corporation, the possibility remains that it will be done by a single inventor or a small team of researchers. In this paper, we address the question of safeguarding a discovery which could without hesitation be said to be worth trillions of dollars. Specifically, we propose a method based on the combination of zero knowledge proofs and provably AI-complete CAPTCHA problems to show that a superintelligent system has been constructed without having to reveal the system itself.

1. Introduction and Motivation

Experts predict that in the next 10 to 100 years scientists will succeed in creating human-level artificial general intelligence (AGI) [1–5]. While it is most likely that AGI will be created by a government agency [6], such as DARPA, or a large corporation such as Google Inc., the possibility remains that it will be done by a single inventor or a small team of “garage inventors.” The history of computer science is the history of such inventors. Steve Jobs (Apple), Bill Gates (Microsoft), Mark Zuckerberg (Facebook), and Page and Brin (Google) to name just a few, all revolutionized the state of technology while they were independent inventors.

What is an inventor to do after a successful construction of an artificially intelligent system? Going public with such an invention may be dangerous as numerous powerful entities will try to steal the invention. Worse yet, they will also likely try to reduce inventors freedom and safety either to prevent leaking of information or to secure necessary assistance in understanding the invention. Potential nemeses include security agencies, government representatives, military complex, multinational corporations, competing scientists, foreign governments, and potentially anyone who understands the value of such an invention.

It has been said that a true AI is the last invention we will ever have to make [7], as it will make all the other inventions for us. Monetary value of a true AI system is hard to overestimate, but it is well known that billions have been spent on research already by governments and industry [8]. Its potential for military complex is unprecedented both in terms of smart weapons and human-free combat [9]. Even if the initial system has only human-level intelligence, such a machine would among other things be capable of designing the next generation of even smarter intelligent machines and it is generally assumed that an intelligence explosion will take place shortly after such a technological self-improvement cycle begins leading to creation of superintelligence. Possession of such a system would clearly put the inventor of the system in danger [7].

In this paper, we address the question of safeguarding a true AI, a discovery which could without hesitation be said to be worth trillions of dollars. Without going into details, we assume that the inventor through code obfuscation, encryption, anonymization, and location obscurity is able to prevent others from directly accessing the system but still wishes to prove that it was constructed. For this purpose, we propose a novel method based on combination of zero knowledge proofs and provably AI-complete CAPTCHA

problems to show that a superintelligent system has been constructed without having to reveal the design of the system.

Alternatively, our method could be used to convince a group of skeptics that in fact a true AI system has been invented without having to resort to time-consuming individual demonstrations. This would be useful if the inventor faces a skeptical reception from the general public and scientific community. In the past, exaggerated claims have been made [8] about some AI systems and so a skeptical reception would not be that surprising. The following sections provide an overview of zero knowledge proofs, CAPTCHAs, and the concept of AI-completeness, all of which are necessary to understand the proposed method.

2. Zero Knowledge Proof

Simply stated a zero knowledge proof (ZKP) is an interactive probabilistic protocol between two parties that gives, with a high degree of certainty, evidence that a theorem is true and that the prover knows a proof while providing not a single bit of information about the said proof to the verifier [10]. ZKP works by breaking up the proof into several pieces in such a way that [10]:

- (1) the verifier can tell whether any given piece of the proof is properly constructed;
- (2) the combination of all the pieces constitutes a valid proof;
- (3) revealing any single piece of the proof does not reveal any information about the proof.

To begin, the prover hides each piece of the proof by applying a one-way function to it. After that, the verifier is allowed to request a decryption of any single piece of the proof. Since the verifier can select a specific piece at random, seeing that it is properly constructed provides probabilistic evidence that all pieces of the proof are properly constructed and so is the proof as the whole [10].

3. CAPTCHA

With the steady increase in popularity of services offered via the Internet, the problem of securing such services from automated attacks became apparent [11]. In order to protect limited computational resources against utilization by the growing number of human impersonating artificially intelligent systems, a methodology was necessary to discriminate between such systems and people [12]. In 1950 Turing published his best known paper “*Computing Machinery and Intelligence*” in which he proposed evaluating abilities of an artificially intelligent machine based on how closely it can mimic human behavior [13]. The test, which is now commonly known as the Turing test, is structured as a conversation and can be used to evaluate multiple behavioral parameters, such as agent’s knowledge, skills, preferences, and strategies [14]. In essence, it is the ultimate multimodal behavioral biometric, which was postulated to make it

possible to detect differences between man and machine [11].

The theoretical platform for an automated Turing test (ATT) was developed by Naor in 1996 [15]. The following properties were listed as desirable for the class of problems which can serve as an ATT:

- (i) many instances of a problem can be automatically generated together with their solutions;
- (ii) humans can solve any instance of a problem quickly and with a low error rate. The answer should be easy to provide either by a menu selection or via typing a few characters;
- (iii) the best known artificial intelligence (AI) programs for solving such problems fail a significant percentage of times, despite the full disclosure of how the test problem is generated;
- (iv) the test problem specification needs to be concise in terms of description and area used to present the test to the user.

Since the initial paper by Naor, a great deal of research has been performed in the area, with different researchers frequently inventing new names for the same concept of human/machine disambiguation [16, 17]. In addition to ATT, the developed procedures are known under such names as [11]: reversed Turing test (RTT) [18], human interactive proof (HIP) [19], mandatory human participation (MHP) [20], or Completely automated public Turing test to tell computers and humans apart (CAPTCHA) [21, 22].

As ongoing developments in AI research allow some tests to be broken [23–26], research continues on developing more secure and user friendly ways of telling machines and humans apart [27–32]. Such tests are always based on as-of-yet unsolved problem in AI [33]. Frequent examples include pattern recognition, in particular character recognition [34–40] or image recognition [41–43]; a number of CAPTCHAs are based on recognition of different biometrics such as faces [44–46], voice [47, 48] or handwriting [49, 50]. Additionally the following types of tests have been experimented with [11, 51] the following.

- (i) *Reading*: password displayed as a cluttered image.
- (ii) *Shape*: identification of complex shapes.
- (iii) *Spatial*: text image is rendered from a 3D model.
- (iv) *Quiz*: visual or audio puzzle or trivia question.
- (v) *Match*: common theme identification for a set of related images.
- (vi) *Virtual reality*: navigation in a 3D world.
- (vii) *Natural*: uses media files collected from the real world, particularly the web.
- (viii) *Implicit*: test is incorporated into the web page navigation system [52].

4. AI-Completeness

A somewhat general definition of the term included in the 1991 Jargon File [53] states

“AI-complete: [MIT, Stanford, by analogy with “NP-complete”] adj. Used to describe problems or subproblems in AI, to indicate that the solution presupposes a solution to the “strong AI problem” (i.e., the synthesis of a human-level intelligence). A problem that is AI-complete is, in other words, just too hard. Examples of AI-complete problems are “The Vision Problem”, building a system that can see as well as a human, and “The Natural Language Problem”, building a system that can understand and speak a natural language as well as a human. These may appear to be modular, but all attempts so far (1991) to solve them have foundered on the amount of context information and “intelligence” they seem to require.”

As such, the term “AI-complete” (or sometimes AI-hard) has been a part of the field for many years [54] and has been frequently brought up to express difficulty of a specific problem investigated by researchers (see [55–68]).

Recent work has attempted to formalize the intuitive notion of AI-completeness, in particular [54].

In 2003, von Ahn et al. [69] attempted to formalize the notion of an AI-problem and the concept of AI-hardness in the context of computer security. An AI-problem was defined as a triple: $\mathcal{P} = (S, D, f)$, where S is a set of problem instances, D is a probability distribution over the problem set S , and $f: S \rightarrow \{0, 1\}^*$ answers the instances. Let $\delta \in 2(0; 1]$. We require that for an $\alpha > 0$ fraction of the humans H , $\Pr_{x \sim D}[H(x) = f(x)] > \delta \dots$ An AI problem $\mathcal{P}$ is said to be (δ, τ) -solved if there exists a program A , running in time at most τ on any input from S , such that $\Pr_{x \sim D, r}[A_r(x) = f(x)] \geq \delta$. (A is said to be a (δ, τ) solution to $\mathcal{P}$.) $\mathcal{P}$ is said to be a (δ, τ) -hard AI problem if no current program is a (δ, τ) solution to $\mathcal{P}$, and the AI community agrees it is hard to find such a solution.” f is a function mapping problem instances to set membership. In other words, it determines if a specific pattern has a property in question. It is necessary that a significant number of humans can compute function f . If the same could be accomplished by a program in efficient time, the problem is considered to be solved. It is interesting to observe that the proposed definition is in terms of democratic consensus by the AI community. If researchers say the problem is hard, it must be so. Also, time to solve the problem is not taken into account. The definition simply requires that some humans be able to solve the problem [69].

In 2007, Shahaf and Amir [70] have published their work on the theory of AI-completeness. Their paper presents the concept of the human-assisted Turing machine and formalizes the notion of different Human Oracles (see section on Human Oracles for technical details). Main contribution of the paper comes in the form of a method for classifying problems in terms of human-versus-machine effort required to find a solution. For some common problems such as

natural language understanding (NLU), the paper proposes a method of reductions allowing conversion from NLU to the problem of speech understanding via Text-To-Speech software.

In 2010, Demasi et al. [71] presented their work on problem classification for artificial general intelligence (AGI). The proposed framework groups the problem space into three sectors.

- (i) *Non-AGI-bound*: problems that are of no interest to AGI researchers.
- (ii) *AGI-bound*: problems that require human-level intelligence to be solved.
- (iii) *AGI-hard*: problems that are at least as hard as any AGI-bound problem.

The paper also formalizes the notion of human oracles and provides a number of definitions regarding their properties and valid operations.

In 2011, Yampolskiy [54] proposed the following formalization of AI-completeness.

Definition 1. A problem C is *AI-complete* if it has two properties:

- (1) it is in the set of AI problems (Human Oracle solvable)
- (2) any AI problem can be converted into C by some polynomial-time algorithm.

Yampolskiy [54] showed that the Turing test problem is an instance of an AI-complete problem and further showed certain other AI problems to be AI-complete (question answering, speech understanding) or AI-hard (Programming) by utilizing polynomial-time reductions.

Furthermore, according to the Encyclopedia of Artificial Intelligence [72] published in 1992, the following problems are all believed to be AI-complete [54, 72].

- (i) *Natural language understanding*—“Encyclopedic knowledge is required to understand natural language. Therefore, a complete Natural Language system will also be a complete Intelligent system.”
- (ii) *Problem solving*—“Since any area investigated by AI researchers may be seen as consisting of problems to be solved, all of AI may be seen as involving Problem Solving and Search”.
- (iii) *Knowledge representation and reasoning*—“...the intended use is to use explicitly stored knowledge to produce additional explicit knowledge. This is what reasoning is. Together Knowledge representation and Reasoning can be seen to be both necessary and sufficient for producing general intelligence—it is another AI-complete area.”
- (iv) *Vision or image understanding*—“If we take “interpreting” broadly enough, it is clear that general intelligence may be needed to do this interpretation, and that correct interpretation implies general intelligence, so this is another AI-complete area.”

5. SuperCAPTCHA

In this section, we describe our SuperCAPTCHA method which combines ideas of ZKP, CAPTCHA and AI-completeness to create a proof of access to a superintelligent system.

Imagine a CAPTCHA based on a problem which has been proven to be AI-complete, meaning only a computer with human-level intelligence or a real human would be able to solve it. We call such a problem SuperCAPTCHA. If we knew for a fact that such a test was not solved by real humans, that would lead us to conclude that a human-level artificially intelligent system has been constructed and utilized. One simple way to eliminate humans as potential test solvers is to design a test which would require contribution of all humans many times over in order to solve the test in the allotted time. In other words, a test comprised of K instances of a SuperCAPTCHA, for large values of K .

We can estimate the current human population at 7 billion people, which is really a great overestimation since not all people have skills to solve even a simple CAPTCHA, much less an AI-complete one. If the developed SuperCAPTCHA test required 50 billion human-effort-hours to be solved and it was solved in 1 hour, we can conclusively state that it has not been done by utilizing real people. To arrive at our conclusion, without the loss of generality, we assume that any AI software could be run on progressively faster hardware until it exceeds the speed of any human by a desired constant factor.

Utilizing the existing AI-complete problems, we propose a few SuperCAPTCHA tests which if properly administered could serve to prove that an artificially intelligent system has been developed without revealing the design of the system. As long as each SuperCAPTCHA is solved an order of magnitude more times than the number of potential human solvers, the conclusion of an artificial origin of the solver will remain valid. Examples of some AI-Complete CAPTCHAs are as follows.

- (i) Provide a detailed description and explanation of a random image.
- (ii) Write a book indistinguishable in quality from those written by human authors.
- (iii) Write a computer program to simulate human-level intelligence (currently too hard for people).

So, suppose a SuperCAPTCHA was administered and was comprised of properly labeling and describing a random set of 100 billion images. Also suppose that it was accomplished in the amount of time in which all humans in the world working together would not be able to complete the task, for example, in 2 minutes. The next question is the evaluation of a claimed solution to a SuperCAPTCHA. Evaluating the complete solution is too complicated, so our proposed method relies on human graders who randomly decide on a piece of the total solution they would like to examine and compare performance of the AI system to that of human users. While the traditional Turing test is based on dialogues, the SuperCAPTCHAs are based on random

sampling and verification. The verification procedure itself has to be represented by an efficient algorithm performing in at most a polynomial time or in probabilistic polynomial time. In our example, if a randomly chosen image's labeling conforms to the expectation of labeling which a human being would have produced, this increases probabilistic evidence towards the belief that a truly artificially intelligent system has been developed. With each additional inspected piece of the solution, public's confidence in such an explanation will increase in a probabilistic fashion inspired by the ZKP protocol. Best of all is that partially solved SuperCAPTCHAs or even cheating attempts by humans to pass SuperCAPTCHA will result in beneficial labeling of large datasets.

With every additional piece of SuperCAPTCHA verified, public's confidence that a true AI has been invented will increase just like in a classical zero knowledge proof system. As additional problems get proven to be AI-complete, the repertoire of potential SuperCAPTCHAs will grow proportionally. It is also interesting to observe that the inventor of a truly intelligent artificial system may delegate design of SuperCAPTCHAs to the system itself.

6. Conclusions

In this paper, we addressed the question of safeguarding an invention of a truly artificially intelligent system from public disclosure while allowing the inventor to claim credit for its invention. Short of simply using the developed AI system covertly and claiming no credit for its invention, our approach is the safest route an inventor could take to obtain credit for the invention while keeping its design undisclosed. Our methodology relies on analysis of output from the system as opposed to the system itself. Specifically we proposed a method based on combination of zero knowledge proofs and provably AI-complete CAPTCHA problems to show that a superintelligent system has been constructed without having to reveal the system itself. The only way to break a SuperCAPTCHA is to construct a system capable of solving AI-complete problems, an artificial general intelligence.

References

- [1] E. Yudkowsky, "Artificial intelligence as a positive and negative factor in global risk," in *Global Catastrophic Risks*, N. Bostrom and M. M. Cirkovic, Eds., pp. 308–345, Oxford University Press, Oxford, UK, 2008.
- [2] N. Bostrom, "Ethical Issues in advanced artificial intelligence," *Review of Contemporary Philosophy*, vol. 5, pp. 66–73, 2006.
- [3] B. Hibbard, "The Ethics and Politics of Super-Intelligent Machines," 2005, <http://www.ssec.wisc.edu/~billh/g/SI.ethics.politics.doc>.
- [4] D. J. Chalmers, "The singularity: a philosophical analysis," *Journal of Consciousness Studies*, vol. 17, no. 9–10, pp. 7–65, 2010.
- [5] J. S. Hall, "Ethics for Machines," 2000, <http://autogeny.org/ethics.html>.
- [6] C. M. Shulman, "Arms control and intelligence explosions," in *Proceedings of the 7th European Conference on Computing and Philosophy*, Barcelona, Spain, 2009.

- [7] I. J. Good, "Speculations concerning the first ultraintelligent machine," *Advances in Computers*, vol. 6, pp. 31–88, 1966.
- [8] S. Russell and P. Norvig, *Artificial Intelligence: A Modern Approach*, Prentice Hall, Upper Saddle River, NJ, USA, 2003.
- [9] R. Sparrow, "Killer Robots," *Journal of Applied Philosophy*, vol. 24, no. 1, pp. 62–77, 2007.
- [10] M. Blum, "How to prove a theorem so no one else can claim it," in *Proceedings of the International Congress of Mathematicians*, pp. 1444–1451, Berkeley, Calif, USA, August 1986.
- [11] R. V. Yampolskiy and V. Govindaraju, "Embedded non-interactive continuous bot detection," *ACM Computers in Entertainment*, vol. 5, no. 4, article 7, 2008.
- [12] C. Pope and K. Kaur, "Is it human or computer? Defending e-commerce with captchas," *IT Professional*, vol. 7, no. 2, pp. 43–49, 2005.
- [13] A. Turing, "Computing machinery and intelligence," *Mind*, vol. 59, no. 236, pp. 433–460, 1950.
- [14] R. M. French, "The turing test: the first 50 years," *Trends in Cognitive Sciences*, vol. 4, no. 3, pp. 115–122, 2000.
- [15] M. Naor, "Verification of a human in the loop or identification via the turing test," 1996, http://www.wisdom.weizmann.ac.il/~naor/PAPERS/human_abs.html.
- [16] H. S. Baird and K. Popat, "Human interactive proofs and document image analysis," in *Proceedings of the 5th International Workshop on Document Analysis Systems*, pp. 507–518, August 2002.
- [17] R. M. Sampson, "Reverse Turing Tests and Their Applications," 2006, <http://www-users.cs.umn.edu/~sampra/research/ReverseTuringTest.PDF>.
- [18] A. L. Coates, H. S. Baird, and R. J. Fateman, "Pessimistic print: a reverse turing test," in *Proceedings of the 6th International Conference on Document Analysis and Recognition*, pp. 1154–1158, Seattle, Wash, USA, September 2001.
- [19] K. Chellapilla, K. Larson, P. Simard, and M. Czerwinski, "Designing human friendly Human Interaction Proofs (HIPs)," in *Proceedings of the ACM Conference on Human Factors in Computing Systems*, pp. 711–720, ACM, New York, NY, USA, April 2005.
- [20] J. Xu, R. Lipton, I. Essa, M. Sung, and Y. Zhu, "Mandatory human participation: a new authentication scheme for building secure systems," in *Proceedings of the 2th International Conference on Computer Communications and Networks*, pp. 547–552, October 2003.
- [21] L. von Ahn, M. Blum, and J. Langford, "How Lazy Cryptographers do AI," *Communications of the ACM*, vol. 47, no. 2, 2007.
- [22] L. V. Ahn, *Utilizing the power of human cycles*, thesis proposal, Carnegie Mellon University, 2004.
- [23] K. Chellapilla and P. Simard, "Using machine learning to break visual human interaction proofs (HIPs)," in *Proceedings of the 17th Advances in Neural Information Processing Systems Conference (NIPS '04)*, MIT Press, Vancouver, Canada, December 2004.
- [24] G. Mori and J. Malik, "Recognizing objects in adversarial clutter: breaking a visual CAPTCHA," in *Proceedings of the 2003 IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, pp. I-134–I-144, June 2003.
- [25] E. F. Aboufadel, J. Olsen, and J. Windle, "Breaking the holiday inn priority club CAPTCHA," *The College Mathematics Journal*, vol. 36, no. 2, 2005.
- [26] G. Moy, N. Jones, C. Harkless, and R. Potter, "Distortion estimation techniques in solving visual CAPTCHAs," in *Proceedings of the 2004 IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR '04)*, pp. II-23–II-28, July 2004.
- [27] Y. Rui, Z. Liu, S. Kallin, G. Janke, and C. Paya, "Characters or Faces: a user study on ease of use for HIPs," in *Proceedings of the 2nd International Workshop on Human Interactive Proofs*, Lehigh University, Bethlehem, Pa, USA, May 2005.
- [28] K. Chellapilla, K. Larson, P. Y. Simard, and M. Czerwinski, "Computers beat Humans at Single Character Recognition in Reading based Human Interaction Proofs (HIPs)," in *Proceedings of the 2nd Conference on Email and Anti-Spam*, Stanford University, California, USA, July 2005.
- [29] K. Chellapilla, K. Larson, P. Y. Simard, and M. Czerwinski, "Building segmentation based human-friendly human interaction proofs (HIPs)," in *Proceedings of the 2nd International Workshop of Human Interactive Proofs (HIP '05)*, H. S. Baird and D. P. Lopresti, Eds., pp. 1–26, Springer, Bethlehem, Pa, USA, May 2005.
- [30] S.-Y. Wang, H. S. Baird, and J. L. Bentley, "CAPTCHA challenge tradeoffs: familiarity of strings versus degradation of images," in *Proceedings of the 18th International Conference on Pattern Recognition (ICPR '06)*, pp. 164–167, August 2006.
- [31] M. May, "Inaccessibility of CAPTCHA. Alternatives to Visual Turing Tests on the Web, W3C Working Group Not," 2005, <http://www.w3.org/TR/turingtest/>.
- [32] D. Lopresti, "Leveraging the CAPTCHA problem," in *Proceedings of the 2nd International Workshop on Human Interactive Proofs (HIP '05)*, pp. 97–110, May 2005.
- [33] L. von Ahn, M. Blum, N. J. Hopper, and J. Langford, "CAPTCHA: using hard AI problems for security," *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, vol. 2656, pp. 294–311, 2003.
- [34] J. Bentley and C. L. Mallows, "CAPTCHA challenge strings: problems and improvements," in *Proceedings of the Document Recognition & Retrieval*, pp. 141–147, San Jose, Calif, USA, January 2006.
- [35] H. S. Baird and T. Riopka, "ScatterType: a reading CAPTCHA resistant to segmentation attack," in *The 12th IS&T/SPIE Conference on Document Recognition and Retrieval (DRR '05)*, Proceedings of SPIE, pp. 197–207, San Jose, Calif, USA, January 2005.
- [36] H. S. Baird, M. A. Moll, and S.-Y. Wang, "A highly legible CAPTCHA that resists segmentation attacks," in *Proceedings of the 2nd International Workshop on Human Interactive Proofs (HIP '05)*, vol. 3517, pp. 27–41, Bethlehem, Pa, USA, May 2005.
- [37] H. S. Baird, M. A. Moll, and S.-Y. Wang, "ScatterType: a legible but hard-to-segment CAPTCHA," in *Proceedings of the 8th International Conference on Document Analysis and Recognition*, pp. 935–939, September 2005.
- [38] M. Chew and H. S. Baird, "A human interactive proof," in *Proceedings of the 10th SPIE-IS&T Electronic Imaging, Document Recognition and Retrieval*, pp. 305–316, January 2003.
- [39] P. Y. Simard, R. Szeliski, J. Benaloh, J. Couvreur, and I. Calinov, "Using character recognition and segmentation to tell computer from humans," in *Proceedings of the 7th International Conference on Document Analysis and Recognition*, August 2003.
- [40] W.-H. Liao and C.-C. Chang, "Embedding information within dynamic visual patterns," in *Proceedings of the IEEE International Conference on Multimedia and Expo (ICME '04)*, pp. 895–898, June 2004.
- [41] M. Chew and J. D. Tygar, "Image recognition captchas," in *Proceedings of the 1th International Information Security Conference*, pp. 268–279, Springer, September 2004.

- [42] W.-H. Liao, "A CAPTCHA mechanism by exchanging image blocks," in *Proceedings of the 18th International Conference on Pattern Recognition (ICPR '06)*, pp. 1179–1183, August 2006.
- [43] M. Dailey and C. Namprempre, "A text-graphics character CAPTCHA for password authentication," in *Proceedings of the 2004 IEEE Region 10 Conference (TENCON '04)*, pp. 45–48, November 2004.
- [44] D. Misra and K. Gaj, "Face recognition CAPTCHAs," in *Proceedings of the Advanced International Conference on Telecommunications and International Conference on Internet and Web Applications and Services, (AICT-ICIW '06)*, p. 122, February 2006.
- [45] Y. Rui and Z. Liu, "ARTiFACIAL: automated reverse turing test using FACIAL features," in *Proceedings of the 11th ACM International Conference on Multimedia*, pp. 295–298, November 2003.
- [46] Y. Rui and Z. Liu, "Excuse me, but are you human?" in *Proceedings of the 11th ACM International Conference on Multimedia*, pp. 462–463, Berkeley, Calif, USA, November 2003.
- [47] G. Kochanski, D. Lopresti, and C. Shih, "A reverse turing test using speech," in *Proceedings of the International Conferences on Spoken Language Processing*, pp. 1357–1360, Denver, Colorado, 2002.
- [48] T.-Y. Chan, "Using a Text-to-Speech Synthesizer to generate a reverse Turing Test," in *Proceeding of the 15th IEEE International Conference on Tools with Artificial Intelligence (ICTAI '03)*, pp. 226–232, November 2003.
- [49] A. Rusu and V. Govindaraju, "A human interactive proof algorithm using handwriting recognition," in *Proceedings of the 8th International Conference on Document Analysis and Recognition*, pp. 967–970, September 2005.
- [50] A. Rusu and V. Govindaraju, "Handwritten CAPTCHA: using the difference in the abilities of humans and machines in reading handwritten words," in *Proceedings of the 9th International Workshop on Frontiers in Handwriting Recognition (IWFHR-9 '04)*, pp. 226–231, October 2004.
- [51] R. V. Hall, "CAPTCHA as a Web Security Control," 2006, <http://www.richhall.com/captcha/>.
- [52] H. S. Baird and J. L. Bentley, "Implicit CAPTCHAs," in *The 12th IS&T/SPIE Conference on Document Recognition and Retrieval (DRR '05)*, Proceedings of SPIE, pp. 191–196, San Jose, Calif, USA, January 2005.
- [53] E. S. Raymond, Jargon File Version 2.8.1, March 1991, http://www.wisdom.weizmann.ac.il/~naor/PAPERS/human_abs.html.
- [54] R. V. Yampolskiy, "AI-complete, AI-hard, or AI-easy: classification of problems in artificial intelligence," type 02, Computer Engineering and Computer Science. University of Louisville, Louisville, Ky, USA, 2011, <http://louisville.edu/speed/computer/tr/UL-CECS-02>.
- [55] E. T. Mueller, *Daydreaming and computation*, Ph.D. dissertation, University of California, Los Angeles, Calif, USA, 1987.
- [56] J. C. Mallery, "Thinking about foreign policy: finding an appropriate role for artificially intelligent computers," in *Proceedings of the Annual Meeting of the International Studies Association*, St. Louis, Mo, USA, 1998.
- [57] C. Gentry, Z. Ramzan, and S. Stubblebine, "Secure distributed human computation," in *Proceedings of the 6th ACM Conference on Electronic Commerce*, pp. 155–164, June 2005.
- [58] P. J. Phillips and J. R. Beveridge, "An introduction to biometric-completeness: the equivalence of matching and quality," in *Proceedings of the IEEE 3rd International Conference on Biometrics: Theory, Applications and Systems (BTAS '09)*, pp. 1–5, Washington, DC, USA, September 2009.
- [59] R. Bergmair, "Natural Language Steganography and an "AI-complete" Security Primitive," in *Proceedings of the 21st Chaos Communication Congress*, Berlin, December 2004.
- [60] N. Ide and J. Véronis, "Introduction to the special issue on word sense disambiguation: the state of the art," *Computational Linguistics*, vol. 24, no. 1, pp. 1–40, 1998.
- [61] R. Navigli and P. Velardi, "Structural semantic interconnections: a knowledge-based approach to word sense disambiguation," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 27, no. 7, pp. 1075–1086, 2005.
- [62] A. S. Nejad, *A framework for analyzing changes in health care lexicons and nomenclatures*, Ph.D. dissertation, Concordia University, Montreal, Quebec, Canada, 2010.
- [63] J. Chen, J. Liu, W. Yu, and P. Wu, "Combining lexical stability and improved lexical chain for unsupervised word sense disambiguation," in *Proceedings of the 2009 2nd International Symposium on Knowledge Acquisition and Modeling (KAM '09)*, pp. 430–433, December 2009.
- [64] J. P. McIntire, P. R. Havig, and L. K. McIntire, "Ideas on authenticating humanness in collaborative systems using AI-hard problems in perception and cognition," in *Proceedings of the IEEE 2009 National Aerospace and Electronics Conference (NAECON '09)*, pp. 50–55, July 2009.
- [65] J. P. McIntire, L. K. McIntire, and P. R. Havig, "A variety of automated turing tests for network security: using AI-hard problems in perception and cognition to ensure secure collaborations," in *Proceedings of the 2009 International Symposium on Collaborative Technologies and Systems (CTS '09)*, pp. 155–162, May 2009.
- [66] E. Mert and C. Dalkilic, "Word sense disambiguation for Turkish," in *Proceedings of the 24th International Symposium on Computer and Information Sciences (ISCIS '09)*, pp. 205–210, Guzelyurt, Turkey, September 2009.
- [67] J. Hendler, "We've come a long way, maybe," *IEEE Intelligent Systems*, vol. 23, no. 5, pp. 2–3, 2008.
- [68] L. Leahu, P. Sengers, and M. Mateas, "Interactionist AI and the promise of ubicomp, or, how to put your box in the world without putting the world in your box," in *Proceedings of the 10th International Conference on Ubiquitous Computing*, pp. 1–10, Seoul, South Korea, September 2008.
- [69] L. von Ahn, M. Blum, N. J. Hopper, and J. Langford, "CAPTCHA: using hard AI problems for security," *Lecture Notes in Computer Science*, vol. 2656, pp. 294–311, 2003.
- [70] D. Shahaf and E. Amir, "Towards a theory of AI completeness," in *Proceedings of the 8th International Symposium on Logical Formalizations of Commonsense Reasoning (Commonsense '07)*, Stanford University, Stanford, Calif, USA, March 2007.
- [71] P. Demasi, J. L. Szwarcfiter, and A. J. O. Cruz, "A theoretical framework to formalize AGI-Hard problems," in *Proceedings of the 3rd Conference on Artificial General Intelligence (AGI '10)*, pp. 178–179, Lugano, Switzerland, March 2010.
- [72] S. C. Shapiro, "Artificial intelligence," in *Encyclopedia of Artificial Intelligence*, S. C. Shapiro, Ed., pp. 54–57, John Wiley, New York, NY, USA, 1992.