
New Measures of Effectiveness for Human Language Technology

■ Research in human language technology (HLT) has made great progress in the past few years. The general field of HLT encompasses a wide range of algorithms and applications dedicated to processing human speech and written communication. The two specialized HLT fields we address here, from the perspective of the needs of the Defense Department, are (1) machine translation systems, which convert text and speech files from one human language to another, for example, from Arabic into English, and (2) speech-to-text (STT) systems, which produce text transcripts when given audio files of human speech as input. Both processes are subject to machine errors and may produce varying levels of garbling in their output. Automatic STT systems are currently capable of producing English text transcripts of conversational telephone speech at a word-error rate of 15.2%, an error reduction of 53% over the past five years. Arabic-to-English machine translation systems are capable of producing English text output at a precision rate of 47% for a weighted word sequence recognition measure, an increase in performance of over 300% over the past three years. These measures of performance are defined in technology-centric terms to help guide research and development. As important as these measures are, they do not say what the technology can do for us. In a collaborative, interdisciplinary project involving Lincoln Laboratory, the MIT Department of Brain and Cognitive Sciences, and the Defense Language Institute Foreign Language Center, we are addressing the question of how these remarkable gains in machine translation and SST are reflected in measures of effectiveness. Our measures of effectiveness are designed to show the impact of HLT applications on the effectiveness of human users in accomplishing real-world language-understanding tasks. For both machine translation and SST, we have developed techniques to scientifically measure the effectiveness of these technologies when they are used by human subjects.

A MAJOR GOAL of current Department of Defense sponsored research in machine translation of human speech and text is to build automatic systems that translate Chinese and Arabic texts into English texts that can be used by English native readers to perform pertinent foreign language tasks. Figure 1 shows a fragment of an Arabic Level 3

text (moderately difficult for non-native speakers), rated according to the Interagency Language Roundtable skill levels, and sample questions that a competent Arabic reader could be expected to answer after reading this text. The translation on the left in the figure is produced by a state-of-the-art machine translation system; the translation on the right in the figure was

رأى الأهرام تبدو تصريحات وزير خارجية الحكومة العراقية المعينة التي هدد فيها بإمكان السماح للقوات الأمريكية التي تحتل العراق بشن هجمات علي البلدان المجاورة للعراق ردا علي دعمها لعمليات المقاومة العراقية.

Al-Ahram Al-view

Foreign minister's statements appear to be designated the Iraqi government, which threatened the American forces, which could be allowed to occupy Iraq to launch attacks on neighbouring countries to Iraq in retaliation for its support for Iraqi resistance.

Al-Ahram Opinion column

Declarations by the Minister of Foreign Affairs in the "appointed" Iraqi government threatening to allow the American forces that occupy Iraq to launch attacks on neighboring countries seems to be in response to their support of the Iraqi resistance.

FIGURE 1. Arabic-to-English translation: machine translation system (left) versus human translation (right). Two sample questions that a competent reader of Arabic could be expected to answer are "List one of the author's expectations for the new Iraqi regime," or "Why does the author consider the minister's remarks dangerous for the Iraqi government?"

produced by professional human translators. The machine translation contains a variety of disfluencies and mistakes. These errors impair our ability to understand facts explicitly stated in the Arabic article, and severely degrade our ability to infer ideas that are not specifically stated but that an educated reader would be expected to make after reading the article in the original language. The professional human translation on the right contains enough nuances for us to begin to understand the point of view of the original author and to make relevant inferences. The decision of whether to use machine translation or human translation depends upon factors such as cost, availability, and the level of quality needed for the task at hand.

Likewise, a major goal of Defense Department sponsored research in speech recognition is to provide high-quality automatic speech-to-text (STT) transcripts of news broadcasts and conversational telephone speech. Unlike dictation tasks, for which an STT system may be trained and tuned to a specific speaker's voice to reduce word recognition errors, these systems need to operate independently of the speaker. The goal is that the transcripts are good enough to use for tasks that would ordinarily be accomplished by listening to the speech broadcasts or recordings. Ordinary STT transcripts not only contain word recognition errors, but they are typically produced in single-

case, lack punctuation, and include, verbatim, every word or word fragment that is spoken, including fillers such as "um," "uh," repeats, false starts, and so on. Figure 2 illustrates three potential ways that we might view a telephone conversation: (a) as an audio signal, (b) as a transcript produced by an experimental STT system (bottom left), and (c) as a reference transcript (bottom right) produced by trained human readers cross-checked with quality controls (i.e., there are no transcription errors in the texts, standard punctuation and capitalization have been applied, and disfluencies have been removed).

The transcription on the right in Figure 2 is the gold standard by which the automatic SST system output is scored. It serves as a target for technology research and development. The benefits of improving automatic speech transcripts fall into two categories: (1) making the transcripts more readable for human readers and (2) improving automatic downstream processes that use these transcripts as input. Our focus here is on the human readers.

For both technologies—machine translation and STT systems—it is clear that the human transcripts and translations are easier to read. What is not clear is the level at which these technologies can enable their intended consumers (e.g., analysts) to perform real tasks. In order to quantify the effectiveness of these

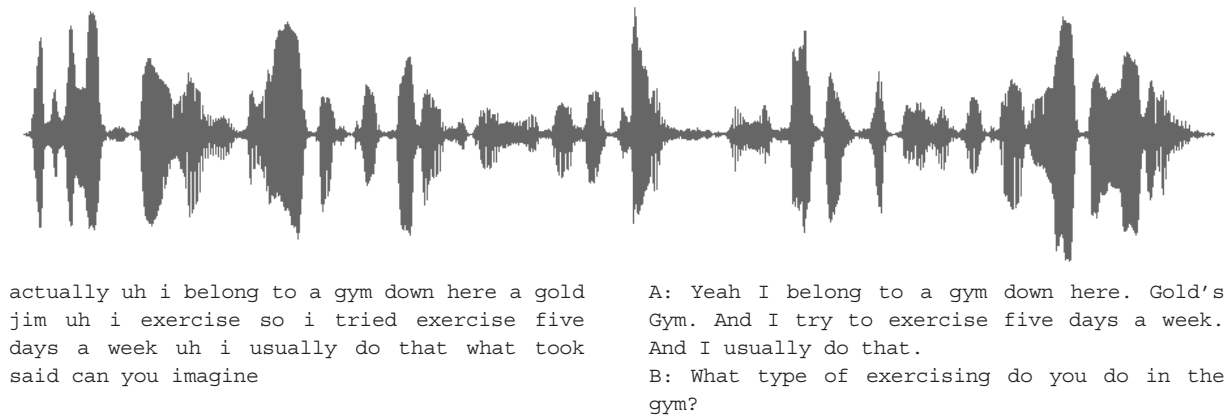


FIGURE 2. Speech-to-text (STT) transcripts from a telephone conversation: experimental system (left) versus human transcript (right). The human transcript is the gold standard by which the experimental results are measured.

technologies and provide feedback for research programs such as the Defense Advanced Research Projects Agency (DARPA) EARS Program (Effective, Affordable Reusable Speech-to-text) and the DARPA TIDES Program (Translingual Information Detection Extraction, and Summarization), we launched a related project that applies rigorous psycholinguistic experimentation and government-standard proficiency evaluation techniques to the methodology of these research efforts.

Figure 3 shows our general framework. Source materials exist in many forms: English audio signals from

broadcast news, conversational telephone speech, foreign language texts from newspapers, and so on. These materials are converted to standard English text (either by translation or transcription) and then presented to human subjects who are asked to read and answer comprehension questions. We present texts that were generated by machines and those generated by humans, and then we measure a human reader's ability to answer questions about the text and the speed at which that reader is able to process the text. With gold-standard texts originally processed by humans, we expect a high level of performance and faster read-

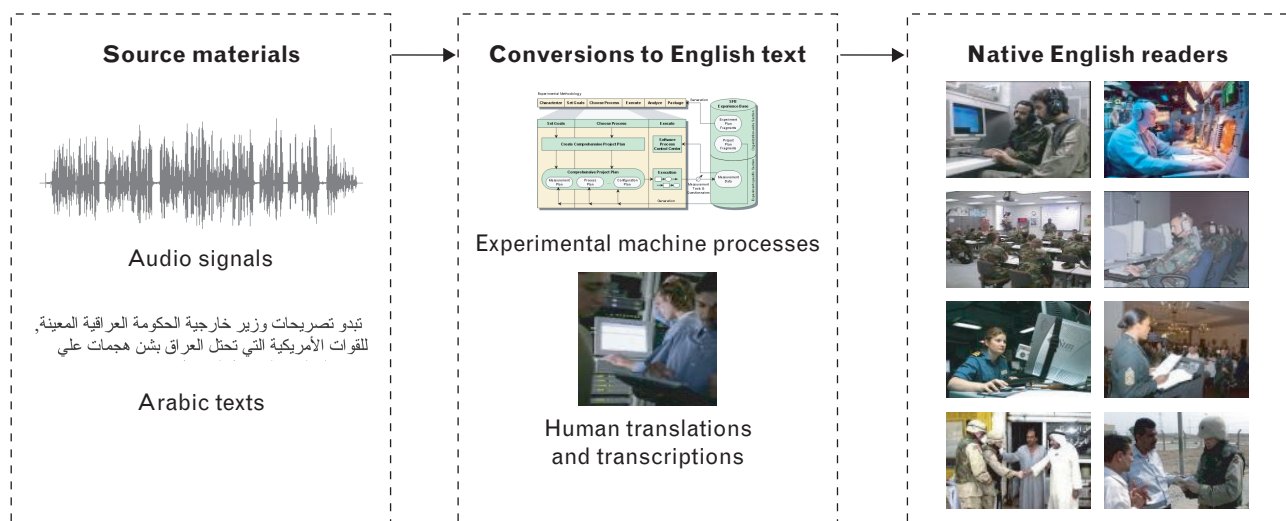


FIGURE 3. Measuring human language technology output effectiveness for English readers. We contrast the ability of people to process the output of experimental language processing algorithms with baselines that use manual processing of human language data.

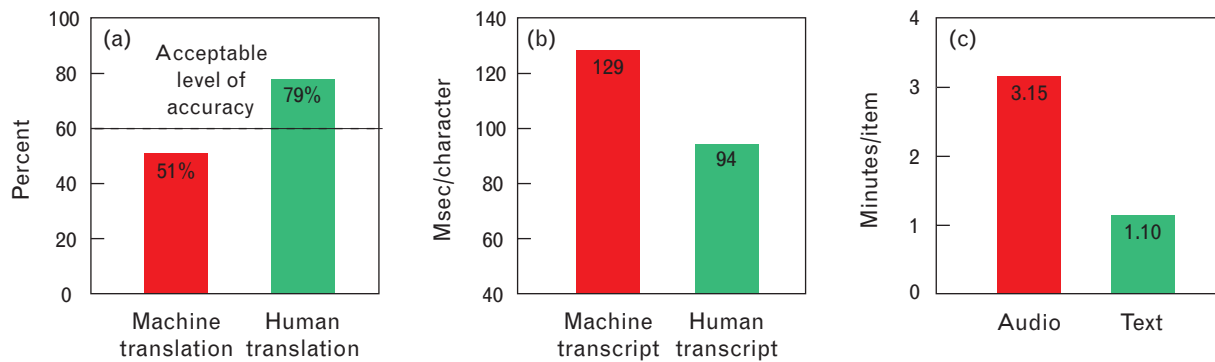


FIGURE 4. Quantitative effects of comprehension for machine translation and human translation. Three types of experiments are shown, defined in terms of (a) translation comprehension based on the Defense Language Proficiency Test, (b) speech transcript reading time, and (c) speech transcript scanning time.

ing times from our human subjects. With experimental machine system output, however, errors can occur, yielding texts that are garbled in places, resulting in a lower level of performance and higher reading times.

By using results from both machine- and human-generated output we can quantify the degradation that these automatic translation and transcription techniques introduce, in a variety of test conditions (such as ranges of errors, types of transformations, and so on). Figure 4 shows results from three different experiments. The top graph shows the results of a machine-translation experiment in which machine-translation output fails to enable people to pass a utility threshold (in this case, 70%) in question-answering accuracy on a standardized test. The middle graph shows that speech transcripts full of errors are read more slowly (129 msec/char versus 94 msec/char, a slow-down of 37%). The graph on the right shows that a simple text-classification experiment can be accomplished approximately three times faster when readers skim a gold-standard transcript, compared with when they listen to the audio files.

Each experiment employs a different measure of effectiveness. The machine translation experiments are based on test results using a modified Defense Language Proficiency Test (DLPT) for Arabic. The DLPT is a standardized test used in the U.S. Defense Department to measure foreign language skills. It has been administered for decades to assess the suitability of personnel for missions requiring foreign language skills, and it has undergone rigorous scrutiny in its

design. We modified the DLPT test in the following way: instead of presenting the human subjects with the original Arabic test materials, we substituted English translations produced by machine translation systems (the test condition) and by professional translation services (the control condition). Our modified DLPT contained texts rated at Interagency Language Roundtable (ILR) levels 1, 2 and 3. At least 70% of the questions must be answered correctly to pass a given level. We found that subjects generally passed Levels 1 and 2, meaning that they demonstrated language survival skills and can understand basic facts, but they generally failed Level 3, meaning that they were not able to read between the lines and make inferences about the materials. The other experiments were based on more generic measures of effectiveness (how fast do subjects read the texts, how accurately do they answer general questions, and how strongly do they prefer materials in different conditions).

The next steps for our project include extending our measures of effectiveness to other areas of HLT evaluation (for example, how effectively can Mandarin Chinese and Arabic audio and text files be distilled for particular tasks) and establishing relationships between different modalities (for example, measures of effectiveness for speech-to-speech machine translation devices to be used to accomplish tasks that require interactive dialog between people who do not speak a common language).

This work is a joint, interdisciplinary effort between the Defense Language Institute Foreign Language

Center, Lincoln Laboratory's Information Systems Technology group, and the MIT Brain and Cognitive Sciences Department. The principal investigators include: Neil Granoien and Martha Herzog and their colleagues at the Defense Language Institute who have served as the principal architects for foreign language proficiency tests for the U.S. Defense Department; professor Edward Gibson, an established leader in the field of psycholinguistics and human sentence processing; and Douglas Jones and Wade Shen and other members of the technical staff at Lincoln Laboratory.

This research project represents an opportunity to bring the technology development community into better contact with the two neighboring fields of psycholinguistics and foreign language training. Since 2002 we have conducted experiments with the participation of hundreds of human subjects at MIT and the surrounding communities, and have written several papers describing our results [1–8].

view, June 2003

7. Erard, Michael, "Translation in the Age of Terror", Technology Review, March 2004
8. Papineni, K.A., Roukos, S., Ward, T., Zhu, W.J.: 2001. BLEU: a method for automatic evaluation of machine translation. Technical Report RC22176 (W0109-022), IBM Research Division, Thomas J. Watson Research Center.

Contributed by Douglas Jones, Wade Shen, and Clifford Weinstein of Lincoln Laboratory; Edward Gibson, Florian Wolf, Rina Patel, and Charlene Chuang of the MIT Department of Brain and Cognitive Sciences; and Neil Granoien, Ray Clifford, and Martha Herzog of the Defense Language Institute.

REFERENCES

1. Jones, Douglas A., Wade Shen, Neil Granoien, Martha Herzog, Clifford Weinstein. 2005. Measuring Translation Quality by Testing English Speakers with a New Defense Language Proficiency Test for Arabic. In proceedings of 2005 International Conference on Intelligence Analysis, May 2-6, 2005, McLean, VA.
2. Jones, Douglas A., Edward Gibson, Wade Shen, Neil Granoien, Martha Herzog, Douglas Reynolds, Clifford Weinstein. Measuring Human Readability of Machine Generated Text: Studies in Speech Recognition and Machine Translation. In proceedings of 2005 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP), March 18-23, 2005, Philadelphia, PA, USA. Special Session on Human Language Technology: Applications and Challenge of Speech Processing.
3. Jones, Douglas A., Wade Shen, Elizabeth Shriberg, Andreas Stolcke, Teresa Kamm, Douglas Reynolds. Two Experiments Comparing Reading with Listening for Human Processing of Conversational Telephone Speech. In Proceedings of Interspeech 2005, 9th European Conference on Speech Communication and Technology. Sept 2-4, Lisbon, Portugal.
4. Jones, Douglas A., Wolf, Florian, Gibson, Edward, Williams, Elliott, Fedorenko, Evelina, Reynolds, Douglas A., Zissman, Marc, "Measuring the readability of automatic speech-to-text transcripts", EUROSPEECH-2003, 1585-1588.
5. Ray Clifford, Neil Granoien, Douglas Jones, Wade Shen, Clifford Weinstein. "The Effect of Text Difficulty on Machine Translation Performance -- A Pilot Study with ILR-Rated texts in Spanish, Farsi, Arabic, Russian and Korean". Proc. of Language Resources & Evaluation Conference, Lisbon 2004.
6. Walter, Chip, "The Translation Challenge", Technology Re-