

JANUS 93: TOWARDS SPONTANEOUS SPEECH TRANSLATION

M. Woszczyna N. Aoki-Wibel F.D. Buø N. Coccaro K. Horiguchi T. Kemp A. Lavie
A. McNair T. Polzin I. Rogina C.P. Rose T. Schultz B. Suhm M. Torita A. Wibel

Carnegie Mellon University — USA
 University of Karlsruhe — Germany

ABSTRACT

public domain. In spring 1993, independently, other speech translation systems [4, 5] have been presented, showing the growing interest in the field. Improvements include increasing coverage, robustness, generality and speed to begin extending our system to spontaneous human-to-human dialogs, however, improvements and changes along several dimensions of our earlier system are necessary to increase speed, robustness and coverage of the system in the face of ill-formed and ungrammatical spontaneous input. The Translation Engine have been upgraded to permit the introduction of spontaneous human input. In this paper, we will report on the state of these most recent efforts. To allow for development and evaluation of adequate data, a large database with human-to-human dialogs is being gathered for English.

2. THE SCHEDULING TASK DATABASE

We have started collecting a large database of human-to-human dialogs centered around the scenario of appointment scheduling. In each recording session, two subjects are each given a calendar (one of 15 scenarios) and asked to schedule a meeting. Initially, we had initially begun building the recording setup with the initial dialogues. The recording setup was designed to read dialogs of the conference and allow one person to speak at a time by way of a push-button. A vocabulary of around 100 words was used. Data is being collected in similar fashion for French, German, Italian, Japanese, Korean, Russian, Spanish and Swedish. The system is being developed for English, French, German, Italian, Japanese, Korean, Russian, Spanish and Swedish.

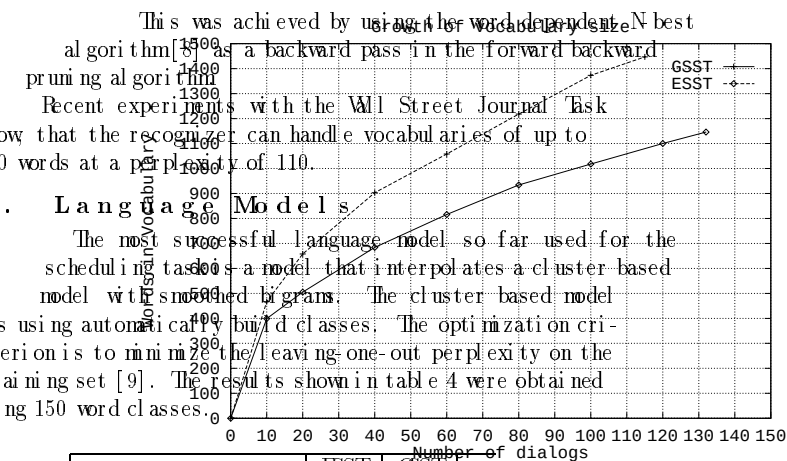
	recorded	transcribed
450	150	
200	140	
70	70	
75	30	

2.1. Scheduling Task Database

The Scheduling Task Database is a semantic parser that takes dialogs cover about 10-12 utterances long. The test set perplexities are about 10-12. The initial set of dialogs were collected from a set of 100 dialogs. The test set perplexities are about 40% larger than the training set. The training set has a larger number of inflections. The test set has a larger number of inflections.

The dialogs are transcribed following the standard transcription format in L2.

The dialogs are transcribed following the standard transcription format in L2. The dialogs are transcribed following the standard transcription format in L2. The dialogs are transcribed following the standard transcription format in L2.



	ESST	GSST
Smoothed Bigrams	46.8	88.2
Cluster	47.0	71.0
Interpolated	39.2	62.1

Figure 2. Development of Vocabulary Size

3. RECOGNITION ENGINE

3.1. Acoustic Modeling

Perplexity Reduction due to automatic clustering

For acoustic modeling, several alternative algorithms are being explored, including TDNN, MF-TDNN, MLP and IVO [6, 7]. In the main JANUS system, an IVO algorithm with context dependent phonemes is currently used in spontaneous speech compared to read speech, particularly for speaker independent recognition, for each phoneme, there is a context independent set of prototypical vectors.

The output scores for each phoneme segment are computed as the euclidean distance using context dependent segment weights.

Changes include the introduction of noise modeling as the improvement of the training algorithms; results in table 2 were obtained using tri phone corrective training and feature weights. Further using cross word tri phones are possible but here.

	1991	1993
Registration	9.1 %	3.7 %
Argument	24 %	7.5 %

Comparison of error-rates

of the recognizer builds a sorted list of s. Speed and memory requirements have been reduced by (table 3).

	1991	1993
3-5 min	3-10 sec	
50Mbytes	7Mbytes	
500	10000	
	120	

Search Module.
Task.

GEST-TIME. Other special and of other task last utter and short
 TE-CONSTRAINTS, AFFIRM and REQUEST RESPONSE to 6 generalized noise
 the interlingua sources utilize familiar parsing module. The total relative
 applied in a module is introduced and a limited reduction of between word
 images cannot be added. It is not the English and spontaneous Schedul-
 also and task the separate evaluation of parser and
 by matching against and generating from a set
 interlingua representations.

4. THE MACHINE TRANSLATION (MT) ENGINE

The MT component that we have previously used has now
 been replaced by a system that can run several alter-
 nate processing methods with parallel modules translating spoken
 language from one language to another, the analysis of spo-
 ken utterances suffer from ill-formed input and recog-
 nition is certainly the hardest part. Based on
 the best hypotheses delivered by the recognition
 module, we can now attempt to select and analyze the most
 sentence hypotheses in view of producing an accu-
 rate translation.
 The central in this attempt is *high fidelity and accu-*
on wherever possible, and *robustness or grace-*
 face of ill-formed or mispronounced input.
 parallel modules attempt to address these
 parser 2) a semantic pattern based approach.
 The most useful
 is mapped onto a common In-
 tent, but domain-specific rep-

Figure 3. Example for Interlingua Output
 Parser

translation. Best separate parsing
 /Compiler. Our formal semantic parser combines frame based semantics
 For application to the grammar [12]. We use a frame based
 parser stack and found the parser used by Carbonell, et
 to process ill-formed text [11], and the MNS system
 slightly extended. Semantic information is rep-
 resenting frames and all frame contains a set of slots
 pieces of information. In order to fill the slots
 es, we use semantic fragment grammars. The
 phrase can be viewed as "phrase spotting".
 important representations are pursued simultane-
 ly in a frame with some of its slots
 sentence represented by a separate Recursive
 which specifies all ways of saying the
 the slot. The grammar is a seman-
 are semantic concepts instead

List Parsing

ionist parsing system PARSEC [13] is used as
 if the LR parser fails to analyze the in-
 tant aspect of the PARSEC system is that
 sentences from a corpus of training ex-
 es the very difficult work of writing
 aspect is that it has proven ro-
 utterances which frequently are
 restarts, repairs or ungram-
 integration with other infor-
 is easier.

Recent developments in PAR

4.5. The Generator

The generation of target language from an Interlingua representation involves two steps: First, with the translation formalism used in the analysis phase, Interlingua representation is mapped into syntactic structure of the target language. The FROSPER system then fed into sentence generation software called "GENIT" to produce a sentence in the target language. A grammar for GENIT is written in the same formalism as the Generalized LR Parser: Esting Generality in JAV/S: A Multi-Lingual Speech phrase structure rules augmented with pseudo unification equations.

As a first experiment for generation on the spontaneous scheduling task, Yoto S. Sugawara and M. Wszczyzna tried to generate 264 new sentences. The results are shown in table 6. 47% of the generated sentences were found to be good or acceptable (see table 6).

ID	Output quality	%
1	good	65.2
2	acceptable	11.0
3	bad	4.5
4	no output	19.3

Table 6. Quality of Japanese Generation

5. CONCLUSION

In this paper we have described a number of system improvements and extensions that have recently been introduced, to accommodate extension of the speech generation system to handle spontaneously spoken (not read) human-to-

human spontaneous negotiation dialogs is being conducted in Japanese and Spanish, and first results of the system's data have been reported.

ACKNOWLEDGEMENTS

We would like to acknowledge the collaboration with the Japanese Language Laboratory, NEC, and the support of the Japanese Ministry of Education, Science and Culture.