

SRA: DESCRIPTION OF THE SOLOMON SYSTEM AS USED FOR MUC-5

Chinatsu Aone, Sharon Flank, Doug McKee, Paul Krause

Systems Research and Applications (SRA)
2000 15th Street North
Arlington, VA 22201
aonec@sra.com

BACKGROUND

SRA used a language-independent, domain-independent, multipurpose text understanding system as the core of the MUC-5 system for extraction from English and Japanese joint venture texts. SRA's NLP core system, SOLOMON, has been under development since 1986. It has been used for a variety of domains, and was aimed from the start to be language-independent, domain-independent, and application-independent. More recently, SOLOMON has been extended to be multilingual, beginning with Spanish in 1990 and Japanese in 1991. The Spanish-Japanese text understanding system that uses SOLOMON was developed for a domain very different from the MUC-5 joint venture domain (cf. Aone, et al. [2]).

SOLOMON's principal applications have been in data extraction, but it is also used in a prototype machine translation system (cf. Aone and McKee [5]). The domain areas in which SOLOMON applications have been developed are: financial, terrorism, medical, and the MUC-5 joint-venture domain. SRA has significantly enhanced its capability to add new domains and languages by developing new strategies for data acquisition using both statistical techniques and a variety of user-friendly tools.

MUC-5 SYSTEM ARCHITECTURE

SOLOMON employs a modular, data-driven architecture to achieve its language- and domain-independence. The MUC-5 system, which uses SOLOMON as a core engine, consists of seven processing modules and corresponding data modules, as shown in Figure 1, which will be described in the following sections.

Message Zoner

The Message Zoner uploads the SGML-annotated text file into the data extraction system. Input files are assumed to have been preprocessed so that they contain only "rigorous markup" (cf. Goldfarb [8]) SGML tags and text; however, we do not require sentences or paragraphs to be tagged. Japanese text is assumed to be encoded in EUC, but tags must be ASCII.

All input, including tags, is tokenized using a simple, language-independent, regular expression recognizer. The (multi-word) tokens are parsed into sentences, paragraphs, headers and documents using a simple operator-precedence grammar (cf. Aho, Sethi and Ullman [1]) operating on punctuation and tags. The tokenizer and parser are written entirely in lex.

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE AUG 1993		2. REPORT TYPE		3. DATES COVERED 00-00-1993 to 00-00-1993	
4. TITLE AND SUBTITLE SRA: Description of the Solomon System as Used for MUC-5				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Systems Research and Applications (SRA), 2000 15th Street North, Arlington, VA, 22201				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES Proceedings of the Fifth Message Understanding Conference (MUC-5), 25-27 Aug 1993, Baltimore, MD. Sponsored by the Defense Advanced Research Projects Agency.					
14. ABSTRACT					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 14	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

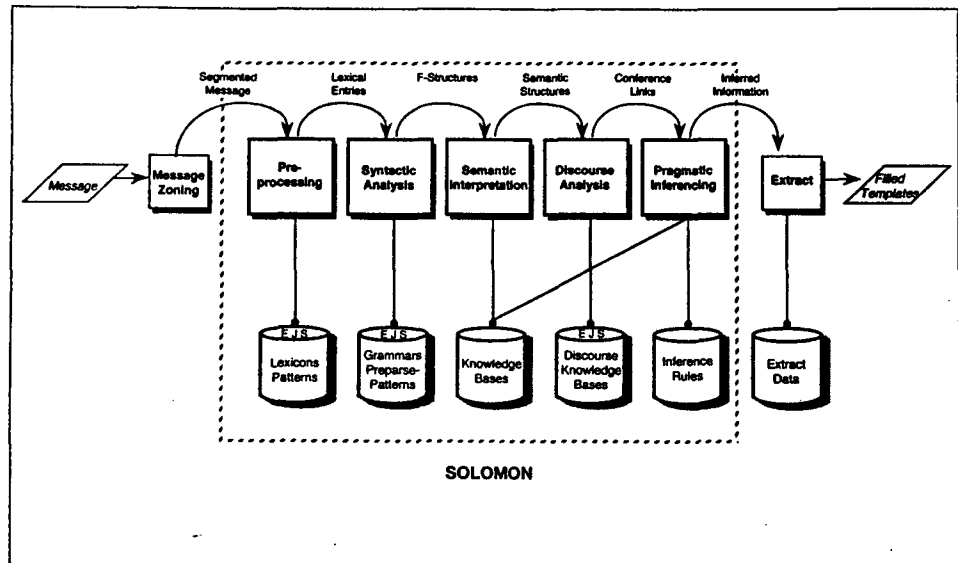


Figure 1: MUC-5 System Architecture

Sentence and paragraph boundaries are inferred using a conservative algorithm and marked as inferred. Inference is not performed if sentences and paragraphs are rigorously marked. The output is piped to a post-processor, which does a fast lookup of each word in a btree gazetteer, and includes entry information in the tokens of place names.

Preprocessing

Preprocessing consists of two processors, the morphological analyzer and the pattern matcher, and associated data in the form of morphological data, lexicons, and patterns for each language. Its input is a tokenized message, and its output is a series of lexical entries with syntactic and semantic attributes.

Declarative morphological data for inflection-rich Japanese and Spanish is compiled into finite-state machines. The English domain lexicon was derived from development texts automatically, using a statistical technique (cf. McKee and Maloney [10]). This derived lexicon also contains automatically acquired domain-specific subcategorization frames and predicate-argument mapping rules called *situation types* (cf. Aone and McKee [3]), as shown in Figure 2.

Pattern recognition handles a wide range of phenomena, including multi-words, numbers, acronyms, money, date, person names, locations, and organizations. We extended the Pattern matcher to handle *multi-level* pattern recognition. The pattern data are divided into ordered multiple groups called *priority groups*, and the patterns in each group are fired sequentially, avoiding recursive applications as much as possible. This extension speeded up the performance of Preprocessing significantly.

Syntactic Analysis

The processor for Syntactic Analysis is a parser based on Tomita's algorithm (cf. Tomita [11]), with modifications for disambiguation during parsing. Syntactic Analysis data consist of X-bar based phrase structure grammars and prepare patterns for each of the three languages, English, Japanese, and Spanish. Syntactic Analysis outputs F-structures (grammatical relations), along the lines of Lexical-Functional Grammar (cf. Bresnan [7]), as shown in Figure 3. The Semantic Interpretation module is interleaved for disambiguation

```

(SWIM
  ((CATEGORY . V)
   (IDIOSYNCRACIES (THEME (MAPPING (LITERAL WITH)))) ; "swim with the big fish"
   (OCCS 11)
   (PREDICATE ANIMATE-OBJECT-ACTIVITY)
   (SITUATION-TYPE ACTIVITY)))
(STEP
  ((CATEGORY . V)
   (IDIOSYNCRACIES (SOURCE (MAPPING (LITERAL FROM)))
                    (GOAL (MAPPING (LITERAL ON INTO))))
   (OCCS 36)
   (PREDICATE CHANGING-EVENT)
   (PROB 8.1 . 1)
   (SITUATION-TYPE ACTIVITY)))
(Team
  ((CATEGORY . V)
   (IDIOSYNCRACIES (THEME (MAPPING (LITERAL WITH))))
   (OCCS 31)
   (PREDICATE ANIMATE-OBJECT-ACTIVITY)
   (SITUATION-TYPE PROCESS CAUSED-PROCESS)))
(SWITCH
  ((CATEGORY . V)
   (IDIOSYNCRACIES (SOURCE (MAPPING (LITERAL FROM))))
   (OCCS 161)
   (PREDICATE TURNKEY-CHANGE)
   (PROB 2.1 . 1)
   (SITUATION-TYPE CAUSED-PROCESS)))

```

Figure 2: Statistically Acquired Lexical Entries

of prepositional phrase attachment, conjunctions, and so on, by calling semantic functions, which are shared by all three languages, from inside the grammar.

Preparsing takes the burden off of main parsing and increases accuracy, by recognizing structures such as sentential complements, appositives, certain PP's, etc. by pattern matching, and sending these to the parser as chunks. These preparse chunks are parsed prior to main parsing using the *same* grammars, and their output consists of F-structures as well.

- Appositives: 業界最大手の東京海上 "industry's largest Tokyo Kaijou"
- Sentences with certain verb endings:
[五年間、資金を積み立て、] その後、一定期間、運用したうえで年金の形でカネを受け取る。]
- PP's: *start production [in january 1990] with production of 20,000 iron*

In order to test the progress of grammar development and pinpoint trouble spots, automatic evaluation of grammars was used. SRA adapted the community-wide program Parseval (cf. Black, et al. [6]) for use in Japanese in addition to English. Testing on Japanese was limited, since there are not many bracketed Japanese texts to use as answer keys.

Semantic Interpretation

Semantic Interpretation uses a language-independent processing module, and its data are predicate-argument mapping rules for each verb, plus both core and domain knowledge bases. Semantic Interpretation works

BRIDGESTONE SPORTS CO. SAID FRIDAY IT HAS SET UP A JOINT VENTURE IN TAIWAN WITH A LOCAL CONCERN AND A JAPANESE TRADING HOUSE ...

```
[ST: <S>
  SUBJECT: [ST: <NP>
    HEAD: BRIDGESTONE-SPORTS-CO.]
  ADJUNCTS: ([ST: <NP>
    HEAD: FRIDAY])
  PREDICATE: [ST: <VP>
    TENSE: PAST
    PREDICATE: (COMMUNICATE)
    ROOT: SAY
    SENT-COMP: [ST: <S>
      SUBJECT: [ST: <NP>
        HEAD: IT]
      PREDICATE: [ST: <VP>
        TENSE: PRESENT
        ASPECT: PERFECT
        PREDICATE: (CREATE)
        ROOT: SET
        VERB-PARTICLE: UP]
      OBJECT: [ST: <NP>
        HEAD: A-JOINT-VENTURE]
      PREP-ARGS: ([ST: <PP>
        MARKED: WITH
        HEAD: A-LOCAL-CONCERN-AND-A-JAPANESE-TRADING-HOUSE])
      ADJUNCTS: ([ST: <PP>
        MARKED: IN
        HEAD: TAIWAN])]]]
```

Figure 3: Simplified F-Structure Output by Syntactic Analysis

off of language-neutral F-structures in order to handle *all* the languages. It outputs semantic structures, i.e. predicate-argument and modification relations, as shown in Figure 4. The predicate-argument mapping rules (i.e. rules which map F-structures to semantic structures) are acquired automatically (cf. Aone and McKee [3]). Domain knowledge bases, on the other hand, were acquired manually. However, a new rapid knowledge acquisition tool called KATool was used to link a lexical entry to its corresponding semantic concept in the knowledge bases (cf. Figure 5).

If a full parse cannot be created, SOLOMON uses a fragment combination strategy. *Debris Parsing* and its subsequent process, *Debris Semantics*, work together to obtain the best interpretation from sentence fragments. They use as data the grammars and knowledge bases, and they output semantic structures just like when a full parse is created. Debris Parsing retrieves the largest and most preferred constituents from the parse stack. It then reparses the rest of the input, and creates debris F-structures with the best fragment constituents. Debris Semantics relies on the semantic interpreter to process each fragment, and then fits fragments together using semantic constraints on unfilled slots.

Discourse Analysis

Discourse Analysis, which was redesigned and implemented this year (cf. Aone and McKee [4]), performs reference resolution. Discourse Analysis uses a data-driven architecture to achieve language-independence, domain-independence, and extensibility. It employs a single language-independent, domain-independent processor, and several discourse knowledge bases, some of which are shared among different languages. The output of Discourse Analysis is a set of semantic structures with coreference links added, i.e. File Cards (cf. Heim [9]). Discourse phenomena handled for the joint venture domain include name anaphora (e.g.

BRIDGESTONE SPORTS CO. SAID FRIDAY IT HAS SET UP A JOINT VENTURE IN
TAIWAN WITH A LOCAL CONCERN AND A JAPANESE TRADING HOUSE ...

```
(COMMUNICATE-1176 (ISA (VALUE (COMMUNICATE)))
  (TIME (VALUE (FRIDAY-1178)))
  (AGENT (VALUE (COMPANY-1146)))
  (THEME (VALUE (CREATE-1163)))
  (TENSE (VALUE (PAST))))
(COMPANY-1146 (ISA (VALUE (COMPANY)))
  (QUANTITY (VALUE ((EXACT 1))))
  (UNIT (VALUE (NATURAL-UNIT)))
  (NAMES (VALUE ((BRIDGESTONE SPORTS CO)))))
(CREATE-1163 (ISA (VALUE (CREATE)))
  (LOCATION (VALUE (COUNTRY-1144)))
  (AGENT (VALUE (THING-1166)))
  (THEME (VALUE (TIE-UP-EVENT-1164)))
  (CO-THEME (VALUE (CONJOINED-COLLECTION COMPANY)-1172))
  (ASPECT (VALUE (PERFECT)))
  (TENSE (VALUE (PRESENT))))
((CONJOINED-COLLECTION COMPANY)-1172
  (ISA (VALUE ((AND CONJOINED-COLLECTION COMPANY))))
  (HAS-MEMBERS (VALUE (COMPANY-1170 COMPANY-1168))))
(COMPANY-1168 (ISA (VALUE (COMPANY)))
  (QUANTITY (VALUE ((EXACT 1))))
  (UNIT (VALUE (NATURAL-UNIT)))
  (LOCATION (TYPE (AND T PHYSICAL-LOCATION)) (VALUE (LOCAL))))
(COMPANY-1170 (ISA (VALUE (COMPANY)))
  (QUANTITY (VALUE ((EXACT 1))))
  (UNIT (VALUE (NATURAL-UNIT)))
  (NATIONALITY (VALUE (JAPAN))))
(COUNTRY-1144 (ISA (VALUE (COUNTRY)))
  (ENGLISH-GAZ-STRING (VALUE (Taiwan (COUNTRY)))))
```

Figure 4: Semantic (Predicate-Argument) Structure

Japanese Knowledge Acquisition Tool (V1.00)	
File: samples	
百貨店 (N)	ENTITY
バンク (COSF)	META-OBJECT
ファイナンス (COSF)	META-OBJECT
電機 (COSF)	META-OBJECT
日系 (PROPN)	META-OBJECT
公司 (COSF)	META-OBJECT
メガビット (CLS)	META-OBJECT
製鉄 (COSF)	META-OBJECT
財閥 (N)	ENTITY
キャピタル (COSF)	META-OBJECT
製紙 (COSF)	META-OBJECT
カンパニー (COSF)	META-OBJECT
物産 (COSF)	META-OBJECT
米穀 (N)	ENTITY
マリーナ (COSF)	META-OBJECT
リミテッド (COSF)	META-OBJECT
米系 (PROPN)	META-OBJECT
日本支社 (N)	ENTITY
興産 (COSF)	META-OBJECT
農林 (COSF)	META-OBJECT
本社 (N)	ENTITY

ENTITY

- PRODUCT
- PHYSICAL-OBJECT
- PART
- COLLECTION
- PARTY
- ABSTRACT-ENTITY

Figure 5: Knowledge Acquisition Tool

```
DISCOURSE: Classified #<DISCOURSE-MARKER DISCOURSE-MARKER-181>("BRIDGESTONE SPORTS") as DP-NAME
DISCOURSE: Found an exact match,
  ante: #<DISCOURSE-MARKER DISCOURSE-MARKER-83>("BRIDGESTONE SPORTS CO.")
  ref: #<DISCOURSE-MARKER DISCOURSE-MARKER-181>("BRIDGESTONE SPORTS")

DISCOURSE: Classified #<DISCOURSE-MARKER DISCOURSE-MARKER-206>("BRIDGESTONE SPORTS") as DP-NAME
DISCOURSE: Found an exact match,
  ante: #<DISCOURSE-MARKER DISCOURSE-MARKER-181>("BRIDGESTONE SPORTS")
  ref: #<DISCOURSE-MARKER DISCOURSE-MARKER-206>("BRIDGESTONE SPORTS")
```

Figure 6: English Discourse Trace Example

業界最大手の東京海上 => 東京海上火災保険

```
DISCOURSE: Classified #<DISCOURSE-MARKER DISCOURSE-MARKER-511>("業界最大手の東京海上") as DP-NAME
DISCOURSE: Found an exact match,
  ante: #<DISCOURSE-MARKER DISCOURSE-MARKER-248>("東京海上火災保険")
  ref: #<DISCOURSE-MARKER DISCOURSE-MARKER-511>("業界最大手の東京海上")
```

東京海上 => 業界最大手の東京海上

```
DISCOURSE: Classified #<DISCOURSE-MARKER DISCOURSE-MARKER-573>("東京海上") as DP-NAME
DISCOURSE: Found an exact match,
  ante: #<DISCOURSE-MARKER DISCOURSE-MARKER-511>("業界最大手の東京海上")
  ref: #<DISCOURSE-MARKER DISCOURSE-MARKER-573>("東京海上")
```

Figure 7: Japanese Discourse Trace Example

"BRIDGESTONE SPORTS" for "BRIDGESTONE SPORTS CO.") and definite NP's such as "THE NEW COMPANY".

The system traces for English and Japanese walkthrough examples are shown in Figure 6 and Figure 7. In the English example, the two instances of name anaphora for "Bridgestone Sports Co." are recognized, while in the Japanese example, all the references to "Tokyo Kaijou Kasai Hoken," including appositives, are resolved.

Pragmatic Inferencing

Pragmatic Inferencing performs reasoning in order to derive implicit information from the text, using a forward chainer and inference rules. Pragmatic Inferencing outputs semantic structures, with inferred information added. It infers additional information from "literal" meanings as required for application domains. For instance, in the walkthrough example, in order to infer "THE TAIWAN UNIT" is a joint venture company from the phrase "THE ESTABLISHMENT OF THE TAIWAN UNIT" the following rule is used.

```
(defrule rule-0009 ((?event) (?event))
  :example ("PNI and SRA established a new company.")
  :if (and (establish ?event)
           (theme ?event ?x)
           (company ?x))
  :then (and (tie-up-event ?event)
            (joint-venture-company ?x)
            (joint-venture-company ?event ?x)
            (in-jv-event ?x ?event)))
```

It is easy for developers to add, change or remove inferred information due to the declarative nature of the inference rules. For instance, to get an additional tie-up from “Company A and Company B tied with Company C”, in jjv-0002, we just had to add another rule to infer that when companies “tie,” they form a tie-up.

```
(defrule rule-0017b ((?event) (?event))
  :example ("PNI tied with SRA")
  :if (and (tie-event ?event)
           (not (theme ?event ?z))
           (agent ?event ?x)
           (company ?x)
           (co-theme ?event ?y)
           (company ?y))
  :then (tie-up-event ?event))
```

Extract

The Extract module performs template generation, translating the domain-relevant portions of our language-independent semantic structures into database records. We maintain a strong distinction between processing and data even in template generation. Thus, we use the same processing module to output in different languages and to several database schemata, including to a flat template-style schema as in MUC-4 and to a more object-oriented schema as in MUC-5.

To do the actual template filling, we rely on Extract data made up of kb-object/slot to db-table/field mapping rules and conversion functions for the individual values (e.g. set fills, string fills). For example, the #nationality slot of an #ORGANIZATION object in our knowledge base corresponds to the Nationality field of the Entity object in the MUC-5 template.

REUSABILITY OF THE SYSTEM

SOLOMON is designed for reusability. Each *processing* module is data-driven and reusable in other languages and other domains, as well as in applications other than data extraction (e.g. machine translation, abstracting, summarization). A large portion of the *data* is also reusable in:

- Other languages and domains
 - Core knowledge bases
- Other domains
 - Morphological data
 - General lexicons
 - General pattern data (e.g. date, location, personal name, organization name)
 - Grammars
 - Some of the discourse knowledge sources
- Other languages
 - Domain knowledge bases

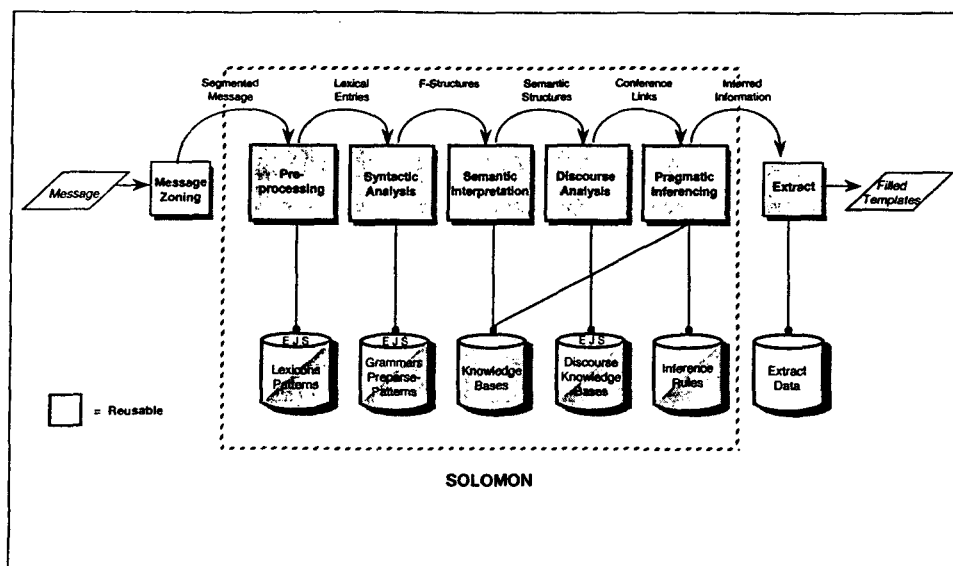


Figure 8: Reusability of SRA's MUC-5 System

- Some of the discourse knowledge sources
- Inference rules
- Extract (template generation) data

The data acquisition tools and techniques are also reusable in other languages and domains. The statistical techniques used to derive lexical information can be reused for other domains. LEXTool, the lexicon acquisition tool, is multilingual and relies on system data files for category and morphological information. KBTool, the knowledge base acquisition tool, is language-independent just as the knowledge bases are language-independent. KATool, the knowledge acquisition tool that links lexicon entries with the appropriate knowledge base concepts, is entirely data-driven as well, and is therefore completely reusable. Figure 8 summarizes the reusability of SRA's MUC-5 system.

TEST RESULTS AND ANALYSIS

Our MUC-5 results for the English and Japanese joint-venture domain task are shown in Table 1. We spent 10.55 person-months for this task, most of which were devoted to data development for both languages (see Table 2). The "other" category includes time spent on developing language-independent data such as a joint-venture domain knowledge base, pragmatic inference rules, and Extract data for template generation.

We believe that the results do not indicate the potential of our system, since the system performance for both languages was still improving after five months of development. Much of the work we did resulted in long-term improvements to our overall text understanding capability, all of which will ensure a stronger base system for future applications. This implies that although the development cycle for data extraction system using a text *understanding* system may be slower in its current maturity stage, the potential for such a system is still unknown and represents a most promising avenue for development. We are particularly pleased with the success of our Japanese system: no other Japanese MUC-5 site is using the full understanding approach, but we did as well and our performance continues to improve.¹

Staff time was the major limiting factor. We needed more time to perform more testing and evaluation

¹ In the 18-month Tipster evaluation, the highest JJV F-measure was about 40.

English						
	ERR	UND	OVG	SUB	REC	PRE
ALL OBJECTS	80	66	26	34	22	49
MATCHED ONLY	48	28	8	23	56	71
TEXT FILTERING	-	25	7	-	74	93
	P&R		2P&R		P&2R	
F-MEASURE	30.80		39.56		25.22	

Japanese						
	ERR	UND	OVG	SUB	REC	PRE
ALL OBJECTS	70	53	34	20	38	52
MATCHED ONLY	43	28	9	14	61	78
TEXT FILTERING	-	6	1	-	94	98
	P&R		2P&R		P&2R	
F-MEASURE	43.92		48.74		39.97	

Table 1: SRA's Scores for the English and Japanese Joint Venture Domain

<i>task</i>	<i>person-months</i>
EJV	3.2
JJV	2.2
Testing	1.5
Documentation	0.25
Other	3.4

Table 2: SRA's Time Expenditure for MUC-5

using the scoring program, and to finely tune Extract (template generation) mapping rules. We discovered we were hampered by formatting errors, and in addition considerable information was "understood" by the system all the way through, but was not extracted by the template generator. Since the discourse module was new, it would have been helpful to have additional time to test and expand it. In addition, we needed more time to fill the OWNERSHIP, REVENUE, and TIME objects, which we simply did not output.

CONCLUSION

Overall, the data-driven architecture in SOLOMON allowed for minimum work on processing modules when working on different languages and domains. We ported the system to Spanish in a week for the demonstration given at the MUC-5 conference.

Although we successfully acquired large amounts of domain data from domain texts in both languages, using both statistical methods and newly developed user-friendly knowledge acquisition tools, we recognize the need to move even more quickly to new domains and languages. We plan to continue our work on automatic acquisition of lexicons, knowledge bases, and links between them in multiple languages.

Tuning performance of each module (e.g. parsing, discourse analysis) as well as the performance of the whole system to a particular task more rapidly is another research issue we identified. We believe that developing automatic evaluation and training algorithms for such automated module/system tuning is crucial to develop a data extraction system that produces optimal results.

ACKNOWLEDGEMENTS

We are indebted to Rajeev Agarwal, Debbie Sanders, and Vera Zlatarski for their hard work and dedication in data development, module testing, and more. We also gratefully acknowledge the contributions of Scott Bennett, David Garfield, and Hatte Blejer to the MUC-5 process.

References

- [1] Alfred V. Aho, Revi Sethi, and Jeffry D. Ullman. *Compilers: Principles, Techniques and Tools*. Addison-Wesley, 1986.
- [2] Chinatsu Aone, Hatte Blejer, Sharon Flank, Douglas McKee, and Sandy Shinn. The Murasaki Project: Multilingual Natural Language Understanding. In *Proceedings of the ARPA Human Language Technology Workshop*, 1993.
- [3] Chinatsu Aone and Doug McKee. Acquiring Predicate-Argument Mapping Information from Multilingual Texts. In *Acquisition of Lexical Knowledge from Text: Proceedings of a Workshop Sponsored by the Special Interest Group on the Lexicon of the Association for Computational Linguistics*, 1993.
- [4] Chinatsu Aone and Doug McKee. Language-Independent Anaphora Resolution System for Understanding Multilingual Texts. In *Proceedings of 31st Annual Meeting of the ACL*, 1993.
- [5] Chinatsu Aone and Doug McKee. Three-Level Knowledge Representation of Predicate-Argument Mapping for Multilingual Lexicons. In *AAAI Spring Symposium Working Notes on Building Lexicons for Machine Translation*, 1993.
- [6] E. Black, S. Abney, D. Flickinger, C. Gdaniec, R. Grishman, P. Harrison, D. Hindle, R. Ingria, F. Jelinek, J. Klavans, M. Liberman, M. Marcus, S. Roukos, B. Santorini, and T. Strzalkowski. A Procedure for Quantitatively Comparing the Syntactic Coverage of English Grammars. In *Proceedings of the Fourth DARPA Speech and Natural Language Workshop*, 1991.
- [7] Joan Bresnan, editor. *The Mental Representation of Grammatical Relations*. MIT Press, 1982.
- [8] Charles F. Goldfarb. *The SGML Handbook*. Oxford, 1990.
- [9] Irene Heim. *The Semantics of Definite and Indefinite Noun Phrases*. PhD thesis, University of Massachusetts, 1982.
- [10] Doug McKee and John Maloney. Using Statistics Gained from Corpora in a Knowledge-Based NLP System. In *Proceedings of The AAAI Workshop on Statistically-Based NLP Techniques*, 1992.
- [11] Masaru Tomita. *Efficient Parsing for Natural Language*. Kluwer, Boston, 1986.

APPENDIX

A ejv-0592 SRA's Original Response

```
<TEMPLATE-0592-1> :=  
  DOC NR: 0592  
  DOC DATE: 241189  
  DOCUMENT SOURCE: "Jiji Press Ltd.;"  
  CONTENT: <TIE_UP_RELATIONSHIP-0592-3>
```

```

<TIE_UP_RELATIONSHIP-0592-2>
<TIE_UP_RELATIONSHIP-0592-2> :=
  TIE-UP STATUS: EXISTING
  ENTITY: <ENTITY-0592-6>
    <ENTITY-0592-5>
    JOINT VENTURE CO: <ENTITY-0592-7>
    ACTIVITY: <ACTIVITY-0592-8>
<ACTIVITY-0592-8> :=
  INDUSTRY: <INDUSTRY-0592-9>
  ACTIVITY-SITE: (Taiwan (COUNTRY) <ENTITY-0592-10>)
<INDUSTRY-0592-9> :=
  INDUSTRY-TYPE: PRODUCTION
  PRODUCT/SERVICE: (67 "A JOINT VENTURE")
<ENTITY-0592-5> :=
  NAME: Taga CO
  TYPE: COMPANY
  ENTITY RELATIONSHIP: <ENTITY_RELATIONSHIP-0592-11>
<ENTITY_RELATIONSHIP-0592-11> :=
  ENTITY1: <ENTITY-0592-5>
    <ENTITY-0592-6>
  ENTITY2: <ENTITY-0592-7>
  REL OF ENTITY2 TO ENTITY1: CHILD
  STATUS: CURRENT
<ENTITY-0592-6> :=
  NAME: Union Precision Casting CO
  ALIASES: "Union Precision Casting"
  TYPE: COMPANY
  ENTITY RELATIONSHIP: <ENTITY_RELATIONSHIP-0592-11>
<ENTITY-0592-7> :=
  NATIONALITY: Taiwan (COUNTRY)
  TYPE: COMPANY
  ENTITY RELATIONSHIP: <ENTITY_RELATIONSHIP-0592-11>
<TIE_UP_RELATIONSHIP-0592-3> :=
  TIE-UP STATUS: EXISTING
  ENTITY: <ENTITY-0592-14>
    <ENTITY-0592-13>
  ACTIVITY: <ACTIVITY-0592-8>
<ENTITY-0592-13> :=
  NAME: Bridgestone Sports CO
  ALIASES: "Bridgestone Sports"
  TYPE: COMPANY
  ENTITY RELATIONSHIP: <ENTITY_RELATIONSHIP-0592-15>
<ENTITY_RELATIONSHIP-0592-15> :=
  ENTITY1: <ENTITY-0592-13>
    <ENTITY-0592-14>
  REL OF ENTITY2 TO ENTITY1: PARTNER
  STATUS: CURRENT
<ENTITY-0592-14> :=
  TYPE: COMPANY
  ENTITY RELATIONSHIP: <ENTITY_RELATIONSHIP-0592-15>

```

B ejv-0592 SRA's Corrected Response

```

<TEMPLATE-0592-1> :=
  DOC NR: 0592
  DOC DATE: 241189
  DOCUMENT SOURCE: "Jiji Press Ltd.;"
  CONTENT: <TIE_UP_RELATIONSHIP-0592-4>
    <TIE_UP_RELATIONSHIP-0592-3>
    <TIE_UP_RELATIONSHIP-0592-2>
<TIE_UP_RELATIONSHIP-0592-2> :=
  TIE-UP STATUS: EXISTING
  ENTITY: <ENTITY-0592-7>
    <ENTITY-0592-6>

```

JOINT VENTURE CO: <ENTITY-0592-8>
 ACTIVITY: <ACTIVITY-0592-9>
 <ACTIVITY-0592-9> :=
 INDUSTRY: <INDUSTRY-0592-10>
 ACTIVITY-SITE: (- <ENTITY-0592-11>)
 <INDUSTRY-0592-10> :=
 INDUSTRY-TYPE: PRODUCTION
 PRODUCT/SERVICE: (67 "A JOINT VENTURE")
 <ENTITY-0592-6> :=
 NAME: Taga CO
 TYPE: COMPANY
 ENTITY RELATIONSHIP: <ENTITY_RELATIONSHIP-0592-12>
 <ENTITY_RELATIONSHIP-0592-12> :=
 ENTITY1: <ENTITY-0592-6>
 <ENTITY-0592-7>
 ENTITY2: <ENTITY-0592-8>
 REL OF ENTITY2 TO ENTITY1: CHILD
 STATUS: CURRENT
 <ENTITY-0592-7> :=
 NAME: Bridgestone Sports CO
 Bridgestone Sports
 TYPE: COMPANY
 ENTITY RELATIONSHIP: <ENTITY_RELATIONSHIP-0592-12>
 <ENTITY-0592-8> :=
 NAME: Bridgestone Sports Taiwan CO
 ALIASES: "Bridgestone Sports CO"
 "Bridgestone Sports"
 NATIONALITY: Taiwan (COUNTRY)
 TYPE: COMPANY
 ENTITY RELATIONSHIP: <ENTITY_RELATIONSHIP-0592-12>
 <TIE_UP_RELATIONSHIP-0592-3> :=
 TIE-UP STATUS: EXISTING
 ENTITY: <ENTITY-0592-16>
 <ENTITY-0592-15>
 JOINT VENTURE CO: <ENTITY-0592-17>
 ACTIVITY: <ACTIVITY-0592-9>
 <ENTITY-0592-15> :=
 NATIONALITY: Taiwan (COUNTRY)
 TYPE: COMPANY
 ENTITY RELATIONSHIP: <ENTITY_RELATIONSHIP-0592-18>
 <ENTITY_RELATIONSHIP-0592-18> :=
 ENTITY1: <ENTITY-0592-15>
 <ENTITY-0592-16>
 ENTITY2: <ENTITY-0592-17>
 REL OF ENTITY2 TO ENTITY1: CHILD
 STATUS: CURRENT
 <ENTITY-0592-16> :=
 NAME: Union Precision Casting CO
 ALIASES: "Union Precision Casting"
 TYPE: COMPANY
 ENTITY RELATIONSHIP: <ENTITY_RELATIONSHIP-0592-18>
 <ENTITY-0592-17> :=
 NATIONALITY: Taiwan (COUNTRY)
 TYPE: COMPANY
 ENTITY RELATIONSHIP: <ENTITY_RELATIONSHIP-0592-18>
 <TIE_UP_RELATIONSHIP-0592-4> :=
 TIE-UP STATUS: EXISTING
 ENTITY: <ENTITY-0592-22>
 <ENTITY-0592-21>
 ACTIVITY: <ACTIVITY-0592-9>
 <ENTITY-0592-21> :=
 TYPE: COMPANY
 ENTITY RELATIONSHIP: <ENTITY_RELATIONSHIP-0592-23>
 <ENTITY_RELATIONSHIP-0592-23> :=
 ENTITY1: <ENTITY-0592-21>
 <ENTITY-0592-22>
 REL OF ENTITY2 TO ENTITY1: PARTNER
 STATUS: CURRENT

<ENTITY-0592-22> :=
NATIONALITY: Japan (COUNTRY)
TYPE: COMPANY
ENTITY RELATIONSHIP: <ENTITY_RELATIONSHIP-0592-23>

C jjv-0002 SRA's Original Response

<テンプレート-0002-1> :=
記事符号: 0002
発行年月日: 850108
ニュース出所: "朝日新聞 朝刊"
内容: <提携-0002-3>
 <提携-0002-2>
<提携-0002-2> :=
提携状況: 現行
エンティティ: <エンティティ-0002-4>
経済活動: <経済活動-0002-5>
<経済活動-0002-5> :=
業種: <業種-0002-6>
場所: (- <エンティティ-0002-7>)
<業種-0002-6> :=
業種別: サービス
<エンティティ-0002-4> :=
エンティティ名: 大和証券
エンティティ別: 企業
エンティティ関係: <エンティティ関係-0002-8>
<エンティティ関係-0002-8> :=
エンティティ乙: <エンティティ-0002-4>
甲対乙関係: パートナー
状況: 現在
<提携-0002-3> :=
提携状況: 現行
エンティティ: <エンティティ-0002-10>
経済活動: <経済活動-0002-5>
<エンティティ-0002-10> :=
エンティティ名: 東京海上
エンティティ別: 企業
エンティティ関係: <エンティティ関係-0002-11>
<エンティティ関係-0002-11> :=
エンティティ乙: <エンティティ-0002-10>
甲対乙関係: パートナー
状況: 現在

D jjv-0002 SRA's Corrected Response

<テンプレート-0002-1> :=
記事符号: 0002
発行年月日: 850108
ニュース出所: "朝日新聞 朝刊"
内容: <提携-0002-4>
 <提携-0002-3>
 <提携-0002-2>
<提携-0002-2> :=
提携状況: 現行
エンティティ: <エンティティ-0002-6>
 <エンティティ-0002-5>
経済活動: <経済活動-0002-7>
<経済活動-0002-7> :=
業種: <業種-0002-8>
場所: (- <エンティティ-0002-9>)
<業種-0002-8> :=
業種別: サービス
<エンティティ-0002-5> :=
エンティティ名: 山一証券
エンティティ別: 企業
エンティティ関係: <エンティティ関係-0002-10>
<エンティティ関係-0002-10> :=
エンティティ乙: <エンティティ-0002-5>
 <エンティティ-0002-6>
甲対乙関係: パートナー
状況: 現在
<エンティティ-0002-6> :=
エンティティ名: 日新火災海上保険
エンティティ別: 企業
エンティティ関係: <エンティティ関係-0002-10>
<提携-0002-3> :=
提携状況: 現行
エンティティ: <エンティティ-0002-12>
経済活動: <経済活動-0002-7>
<エンティティ-0002-12> :=
エンティティ名: 大和証券
エンティティ別: 企業
エンティティ関係: <エンティティ関係-0002-13>
<エンティティ関係-0002-13> :=
エンティティ乙: <エンティティ-0002-12>
甲対乙関係: パートナー
状況: 現在
<提携-0002-4> :=
提携状況: 現行
エンティティ: <エンティティ-0002-15>
経済活動: <経済活動-0002-7>
<エンティティ-0002-15> :=
エンティティ名: 東京海上火災保険
別名: "東京海上"
エンティティ別: 企業
エンティティ関係: <エンティティ関係-0002-16>
<エンティティ関係-0002-16> :=
エンティティ乙: <エンティティ-0002-15>
甲対乙関係: パートナー
状況: 現在