

# **Automatically Improved Category Labels for Syntax-Based Statistical Machine Translation**

**Greg Hanneman**

Language Technologies Institute

Ph.D. Thesis Proposal

January 18, 2011

**Thesis Committee**

Alon Lavie (chair), Stephan Vogel,  
Noah Smith, and David Chiang (USC)



**Carnegie Mellon**

# Syntax-Based Statistical MT

- In this thesis, means MT systems that
  - Acquire syntax in a supervised manner from constituency parse trees
  - Model it with a synchronous context-free grammar (SCFG)

$NP \rightarrow DET\ N\ ADJ$  (“la voiture bleue”)

$NP \rightarrow DT\ JJ\ NN$  (“the blue car”)

$NP::NP \rightarrow [DET^1\ N^2\ ADJ^3]::[DT^1\ JJ^3\ NN^2]$

# SCFG Labeling Schemes

$$X::X \rightarrow [X^1 \text{ de } X^2]::[X^2 X^1]$$

[Chiang 2005]

Target  
Labels



Hiero

Source  
Labels

# SCFG Labeling Schemes

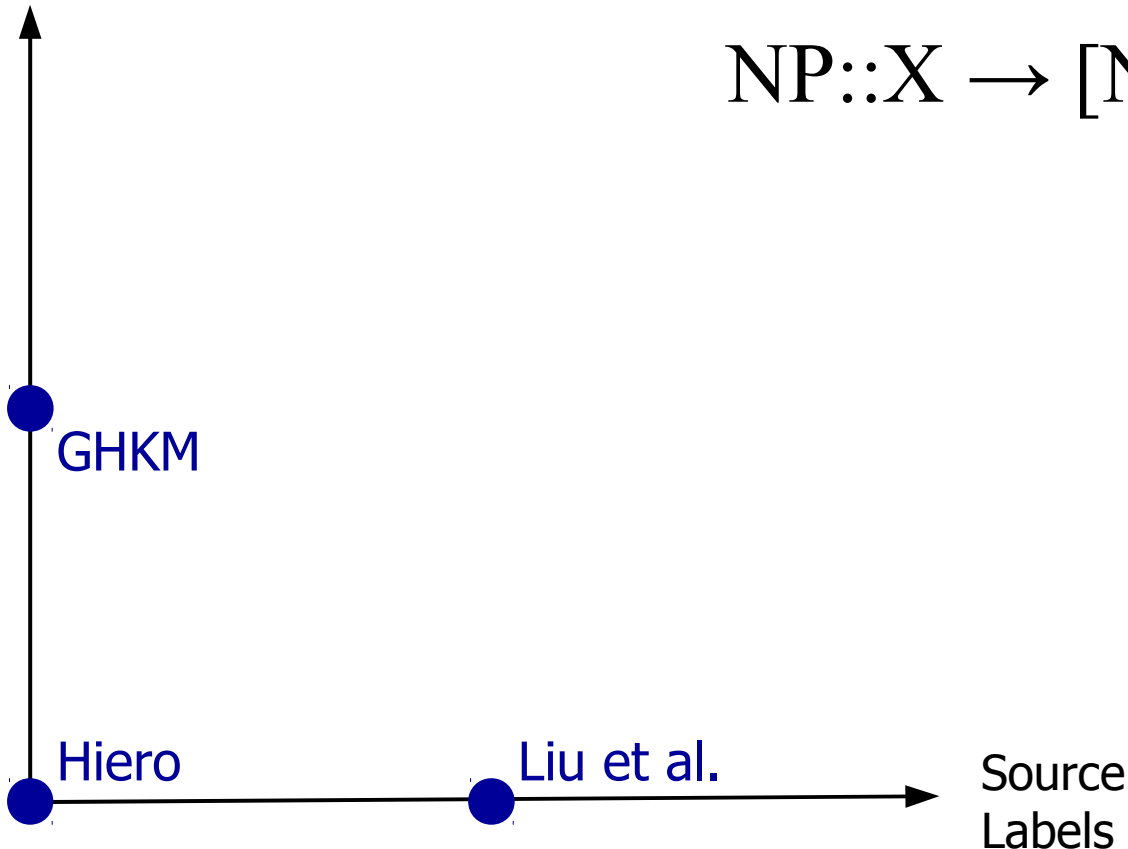
$$X::NP \rightarrow [X^1 \text{ de } X^2]::[N^2 \text{ NP}^1]$$

[Galley et al. 2004]

$$NP::X \rightarrow [NP^1 \text{ de } N^2]::[X^2 \text{ X}^1]$$

[Liu et al. 2006]

Target  
Labels

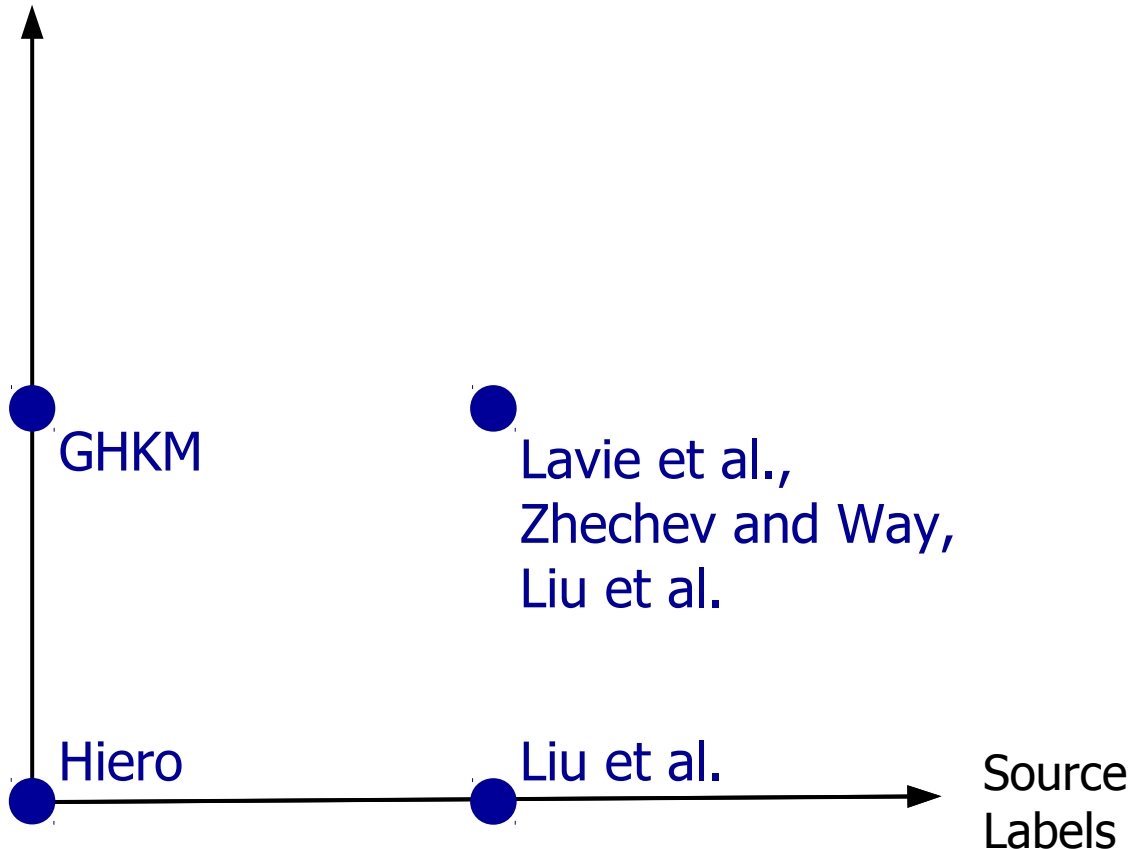


# SCFG Labeling Schemes

$$\text{NP}::\text{NP} \rightarrow [\text{NP}^1 \text{ de } \text{N}^2]::[\text{N}^2 \text{ NP}^1]$$

[Lavie et al. 2008;  
Zhechev and Way 2008;  
Liu et al. 2009]

Target  
Labels



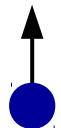
# SCFG Labeling Schemes

$S \rightarrow [\text{il y a } X^1]::[\text{there is } S/\text{NP}^1]$

$\text{NP} \rightarrow [\text{les } X^1 \text{ voitures}]::[\text{the JJ+JJ}^1 \text{ cars}]$

[Zollmann and Venugopal 2006]

Target  
Labels



SAMT



GHKM



Hiero



Lavie et al.,  
Zhechev and Way,  
Liu et al.



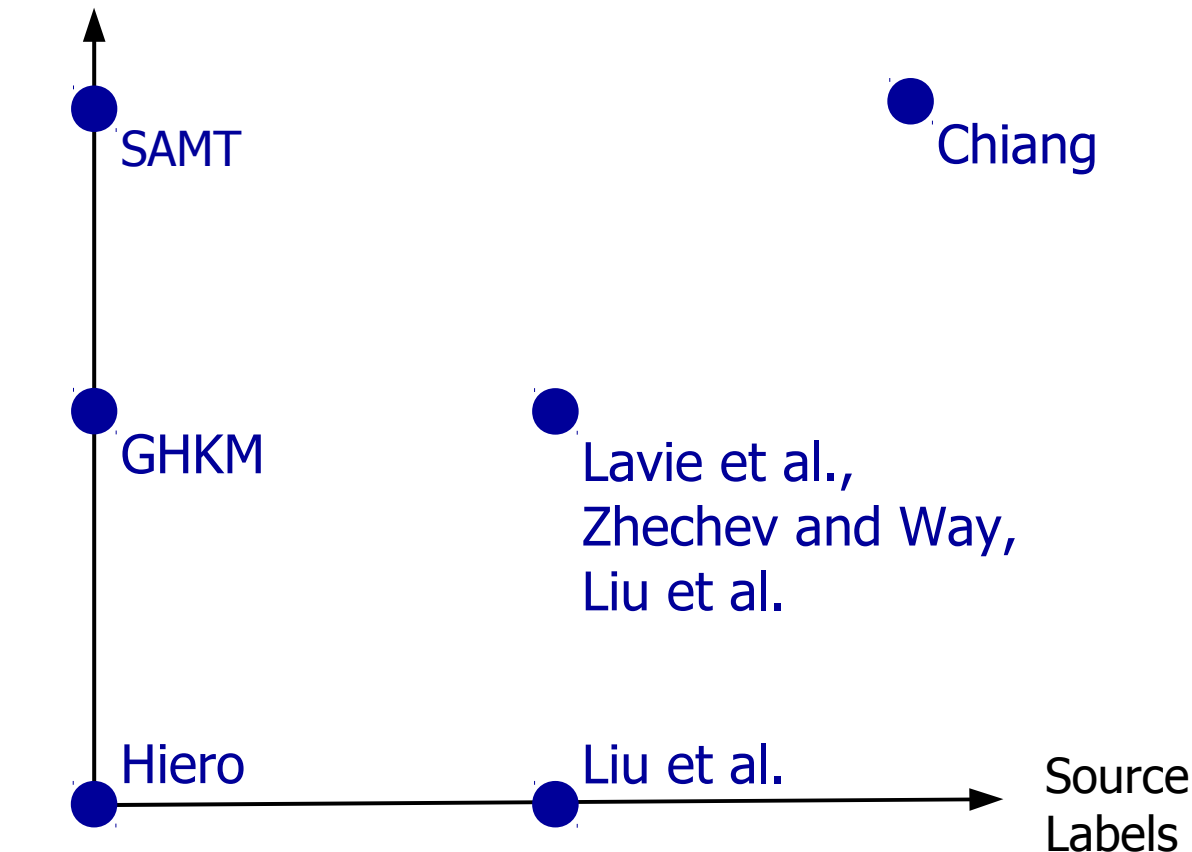
Liu et al.

Source  
Labels

# SCFG Labeling Schemes

$S::S \rightarrow [\text{il y a } S/\text{NP}^1]::[\text{there is } S/\text{NP}^1]$

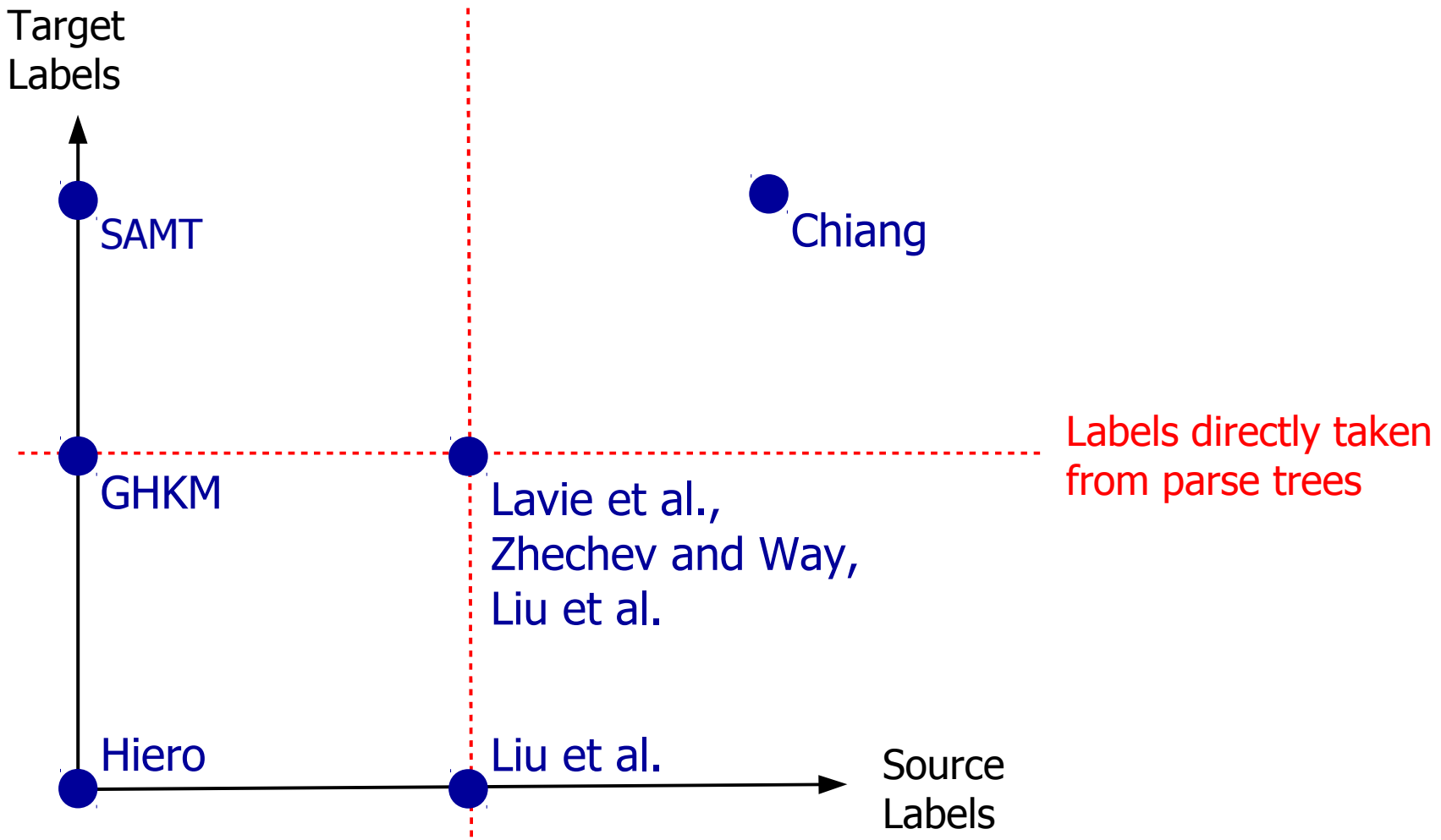
Target Labels  $\text{NP}::\text{NP} \rightarrow [\text{les } D+A^1 \text{ voitures}]::[\text{the } \text{JJ}+\text{JJ}^1 \text{ cars}]$   
[Chiang 2010]



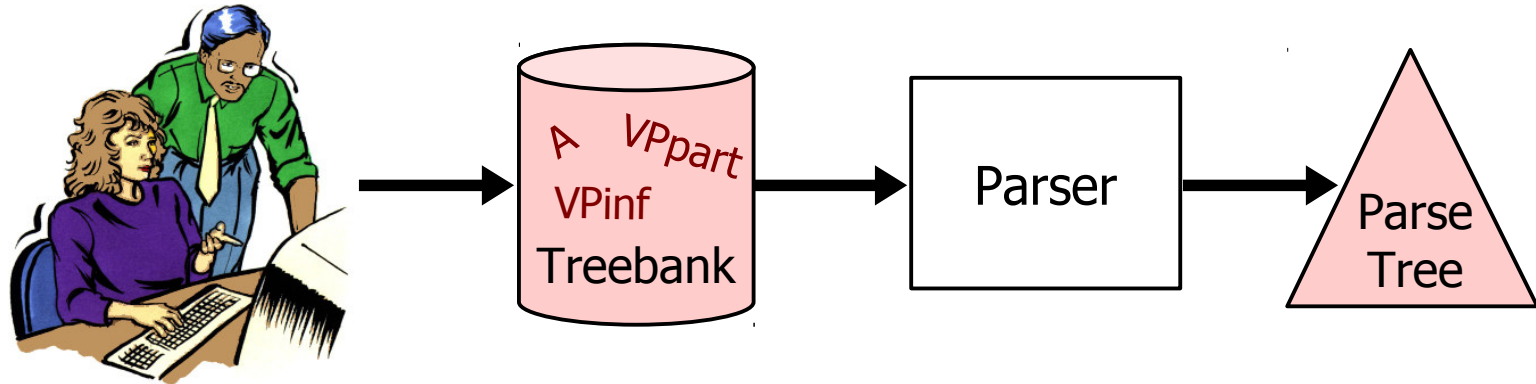
# SCFG Labeling Schemes

Labels directly taken  
from parse trees

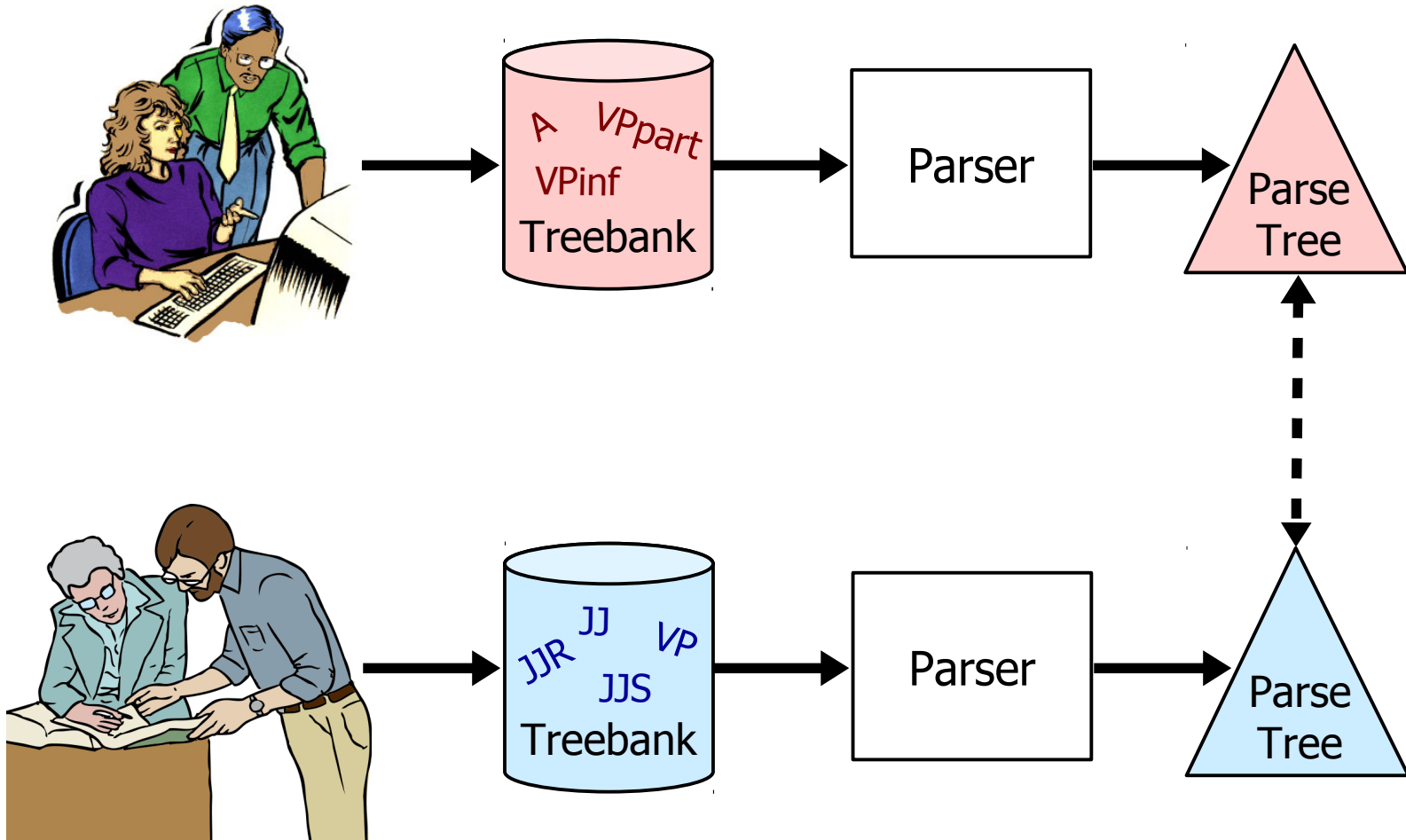
Target  
Labels



# Parse Tree Labels Assume a Lot



# Parse Tree Labels Assume a Lot



# Thesis Statement

The nonterminal set used in SCFG-based MT is important.

Using category labels from treebanks is suboptimal for MT, but we can make labels better automatically.

# Thesis Proposal

- Define and measure the effect labels have
  - Spurious ambiguity, rule sparsity, and reordering precision
- Explore the space of labeling schemes

– Collapsing labels

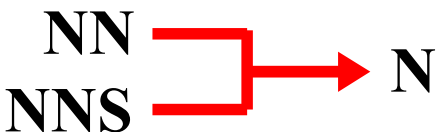


Diagram illustrating collapsing labels: **NN** and **NNS** are grouped by a red bracket and mapped to **N**.

– Refining labels

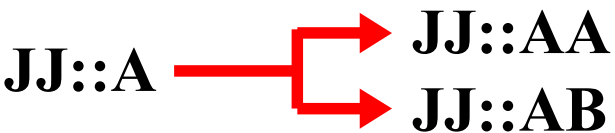


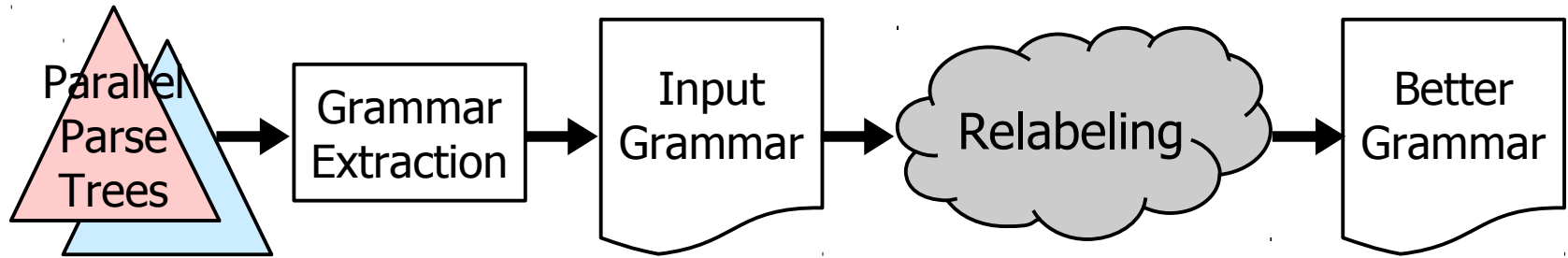
Diagram illustrating refining labels: **JJ::A** is mapped to a red bracket, which then splits into two red arrows pointing to **JJ::AA** and **JJ::AB**.

– Correcting local labeling errors



Diagram illustrating correcting local labeling errors: **PRO** is mapped to **N**.

# Thesis Claims

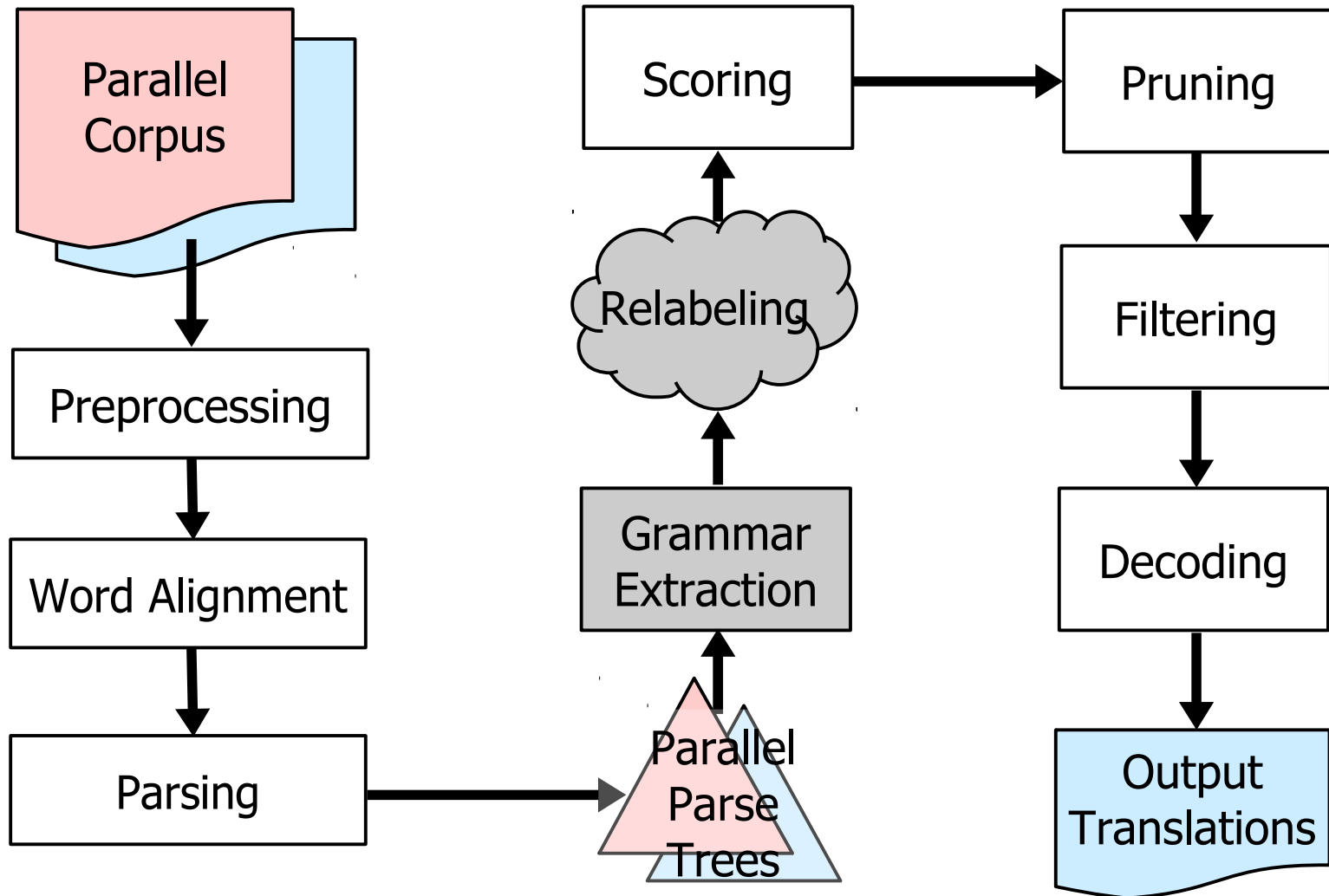


- “Better” will be defined by
  - Significant improvement on automatic metrics
  - Reduced spurious ambiguity and rule sparsity
  - Maintained or improved reordering precision
- Improvements will be demonstrated on at least two language pairs

# Outline

- Introduction
- Baseline system and experimental setup
- Label-based system properties
- Extraction and relabeling techniques
- Putting it all together

# Experimental Setup



# Baseline System

[Hanneman et al. 2010]

- French–English, based on WMT 2010 data
  - 8.6 million sentence pairs of Europarl, news commentary, and UN document text
- Stanford and Berkeley parsers
- Extracted 14 million unique hierarchical rules and 41 million unique phrase pairs
- Pruned to 10,000 most frequent rules and all phrase pairs matching test set
- Joshua SCFG-based decoder

# Language Pairs and Data

- French–English
  - Millions of sentence pairs of Europarl, news commentary, UN document, and Web text
  - Annual WMT evaluation
- Arabic–English and Chinese–English
  - Millions of sentence pairs released from LDC
  - Periodic NIST evaluation
- Chinese–English proving ground
  - FBIS corpus of 300,000 sentence pairs

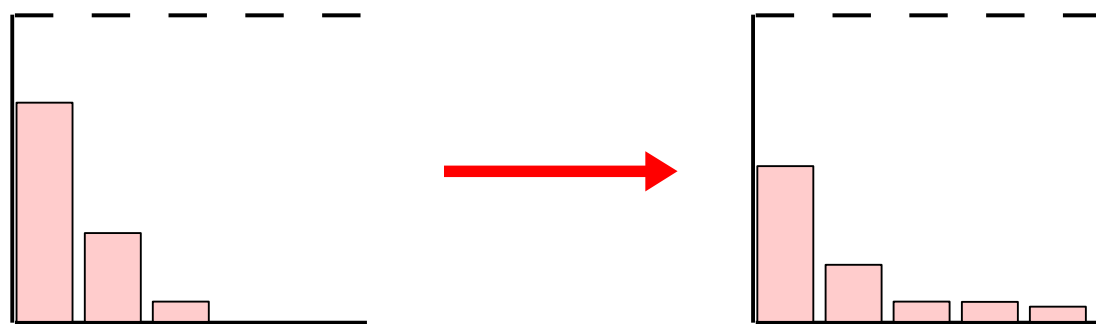
# Outline

- Introduction
- Baseline system and experimental setup
- Label-based system properties
  - Spurious ambiguity
  - Rule sparsity
  - Reordering precision
- Extraction and relabeling techniques
- Putting it all together

# (1) Spurious Ambiguity

- “Many derivations that are distinct yet have the same model feature vectors and give the same translation” [Chiang 2005]
- Weakens conditional probabilities

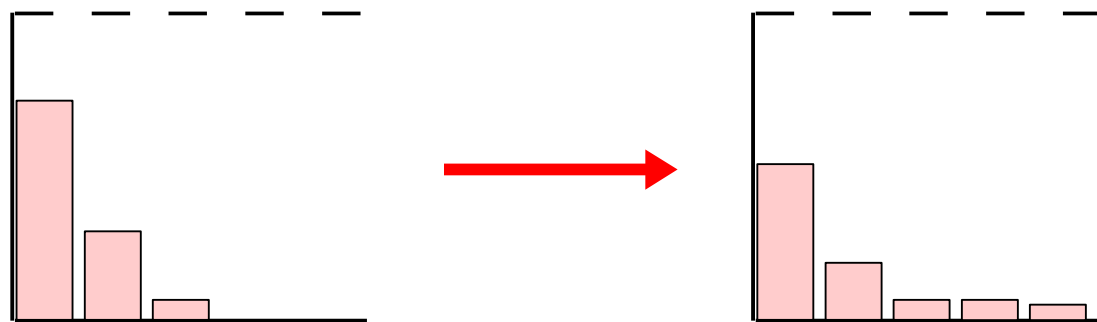
$$\begin{array}{ccc} P(\ell_s, \ell_t, r_s \mid r_t) & P(\ell_s, \ell_t \mid r_s) & P(\ell_s, \ell_t \mid r_s, r_t) \\ P(\ell_s, \ell_t, r_t \mid r_s) & P(\ell_s, \ell_t \mid r_t) & P(r_s, r_t \mid \ell) \end{array}$$



# (1) Spurious Ambiguity

- “Many derivations that are distinct yet ~~have the same model feature vectors and~~ give the same translation” [Chiang 2005]
- Weakens conditional probabilities

$$\begin{array}{lll} P(\ell_s, \ell_t, r_s \mid r_t) & P(\ell_s, \ell_t \mid r_s) & P(\ell_s, \ell_t \mid r_s, r_t) \\ P(\ell_s, \ell_t, r_t \mid r_s) & P(\ell_s, \ell_t \mid r_t) & P(r_s, r_t \mid \ell) \end{array}$$



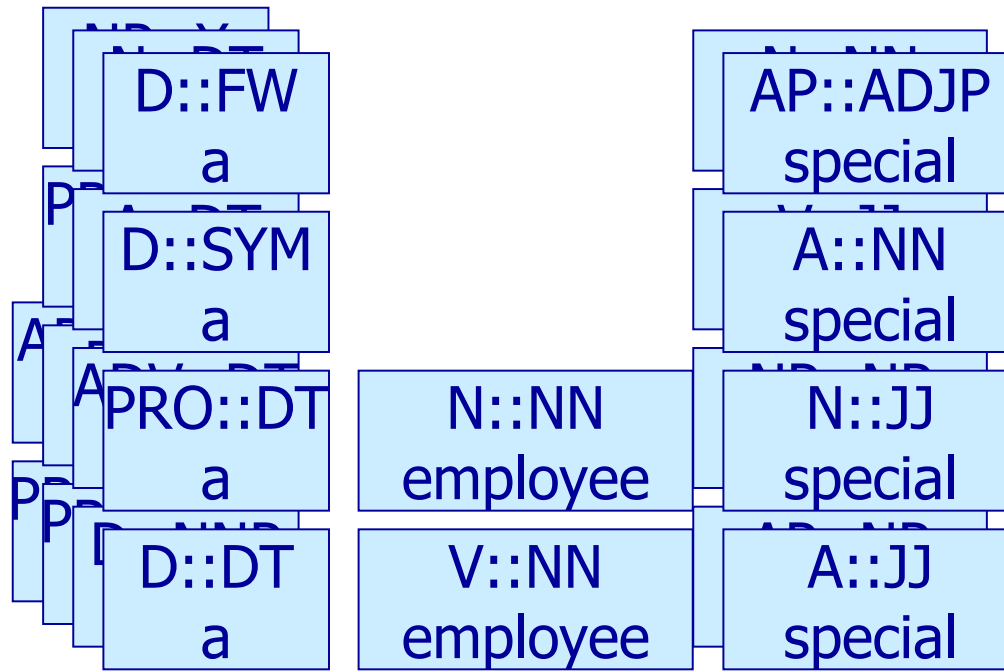
# **(1) Spurious Ambiguity**

- Generates duplicate chart entries with varying left-hand-side labels

une employée spéciale

# (1) Spurious Ambiguity

- Generates duplicate chart entries with varying left-hand-side labels



une employée spéciale

## **(2) Rule Sparsity**

- Necessary rule not found in grammar
- Prevents building of translation fragment

l' office monétaire de Hongkong

## (2) Rule Sparsity

- Necessary rule not found in grammar
- Prevents building of translation fragment

$NP::NP \rightarrow [NP^1 \text{ de } NP^2]::[NP^2 NP^1]$

$NP::NP \rightarrow [NP^1 \text{ de } N^2]::[NN^2 NP^1]$

NP::NP  
monetary office

N::NN  
office

A::JJ  
monetary

N::NNP  
Hong Kong

l' office monétaire de Hongkong

## (2) Rule Sparsity

- Necessary rule not found in grammar
- Prevents building of translation fragment

$NP::NP \rightarrow [NP^1 \text{ de } NP^2]::[NP^2 NP^1]$

$NP::NP \rightarrow [NP^1 \text{ de } N^2]::[NN^2 NP^1]$

Need an  
 $NP::NP$  or  
an  $N::NN$



$NP::NP$   
monetary office

$N::NN$   
office

$A::JJ$   
monetary

$N::NNP$   
Hong Kong

l' office monétaire de Hongkong

# (3) Reordering Precision

- Restricts rules to only apply in contexts specified by their labels

$$X::X \rightarrow [X^1 \text{ de } X^2]::[X^1 X^2]$$

$$X::X \rightarrow [X^1 \text{ de } X^2]::[X^2 X^1]$$

$$PP::X \rightarrow [P^1 \text{ de } NP^2]::[X^1 X^2]$$

$$NP::X \rightarrow [N^1 \text{ de } N^2]::[X^2 X^1]$$

$$PP::PP \rightarrow [P^1 \text{ de } NP^2]::[IN^1 NP^2]$$

$$NP::NP \rightarrow [N^1 \text{ de } N^2]::[NN^2 NNS^1]$$

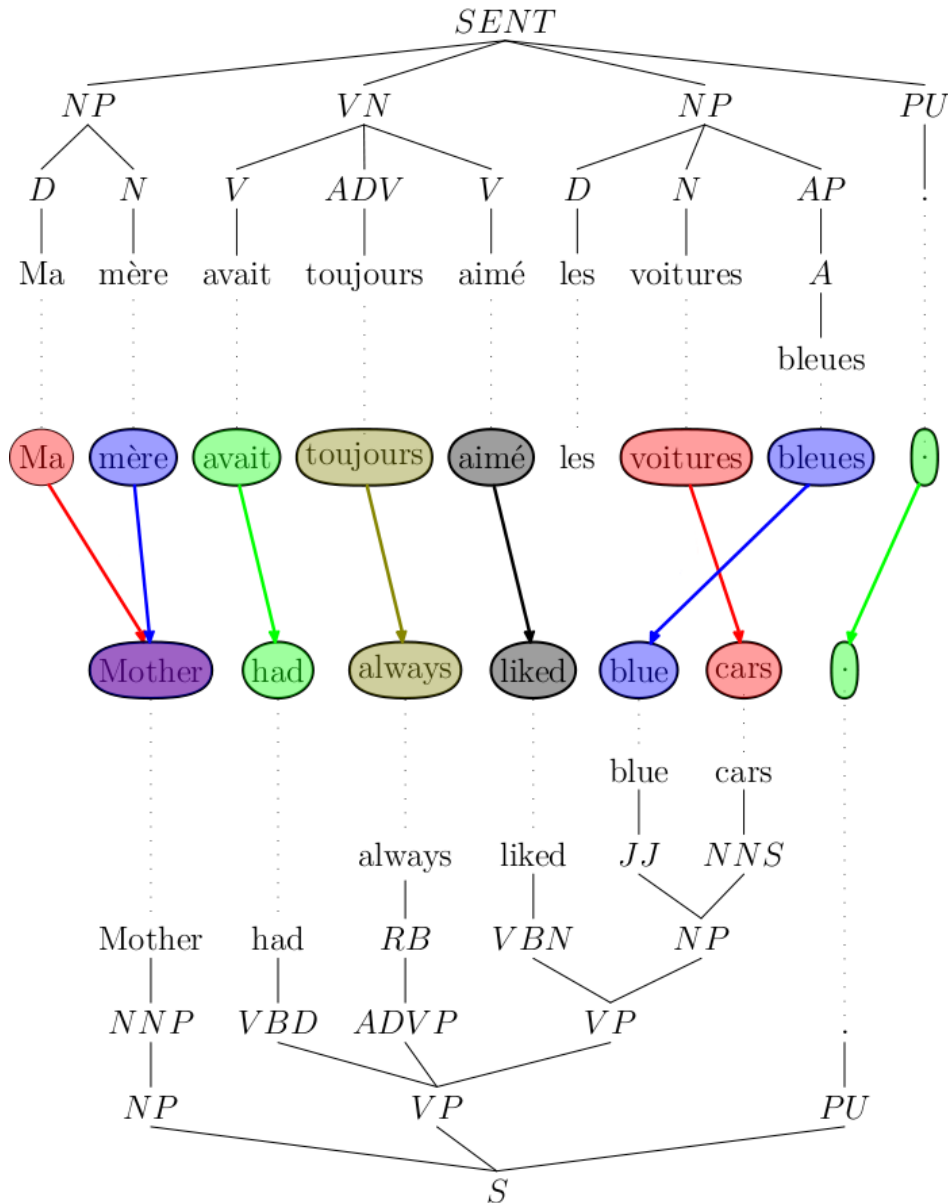
# Measuring These Properties

- Simple counting in the grammar
  - Left-hand-side labels per right-hand side
  - Fraction of possible labelings that instantiate a reordering pattern
- Counting within runtime system
  - Cell contents in completed decoder charts
  - Words out of order on a test set
- Constrained decoding [Auli et al. 2009]
  - Reference reachability

# Outline

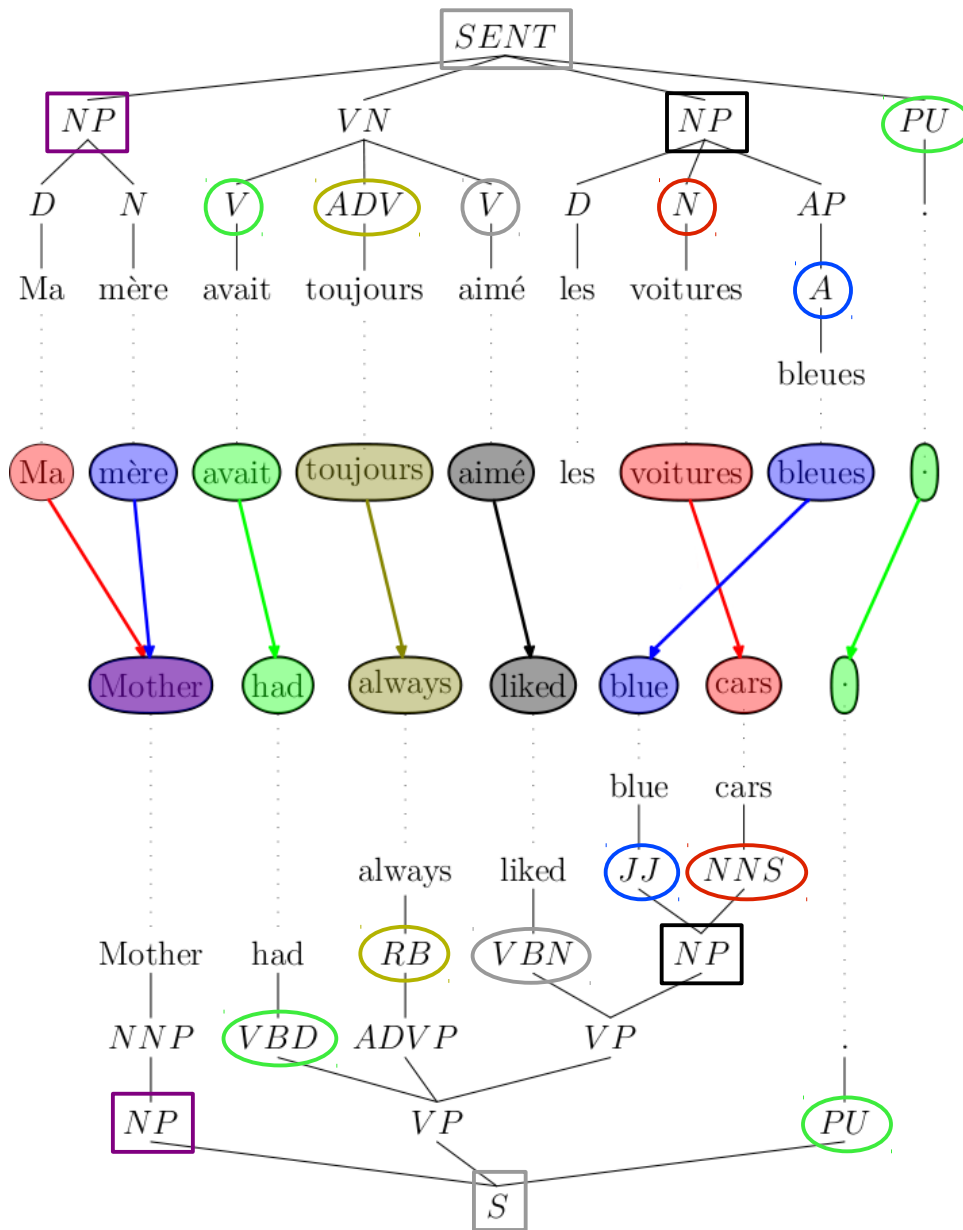
- Introduction
- Baseline system and experimental setup
- Label-based system properties
- **Extraction and relabeling techniques**
  - General-purpose rule extractor
  - Label collapsing
  - Label refining
  - Correcting local labeling errors
- Putting it all together

# (1) General Rule Extractor



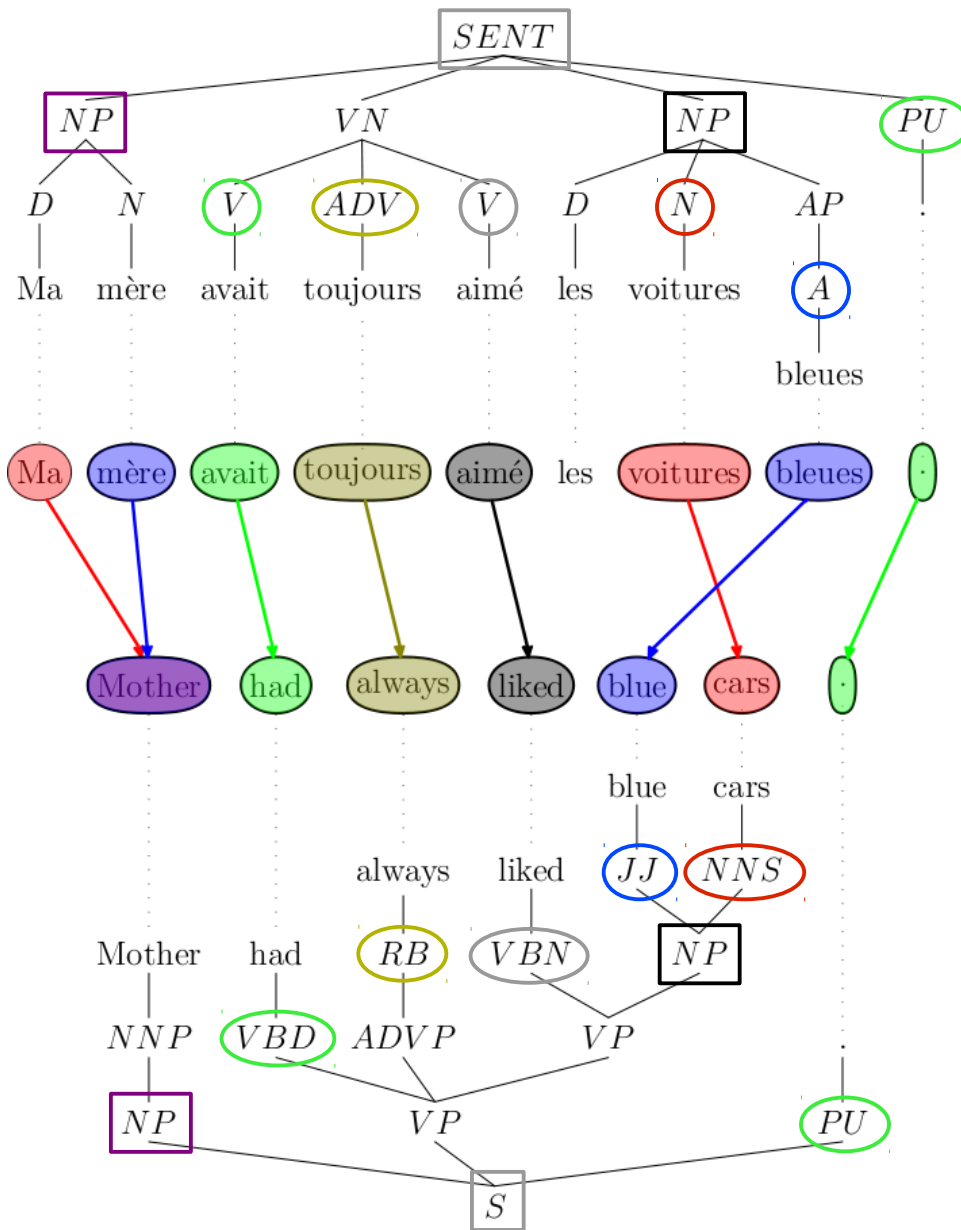
# Baseline Extraction Method

[Lavie et al. 2008]



# Baseline Extraction Method

[Lavie et al. 2008]



NP::NP → [les voitures bleues]::  
[blue cars]

NP::NP → [les N<sup>1</sup> A<sup>2</sup>]::  
[JJ<sup>2</sup> NNS<sup>1</sup>]

# Proposed Extensions

- Node alignment
  - Remove constraints on ambiguous alignments
  - Align single node to sequence of sibling nodes
- Rule extraction
  - Produce multiple tree decompositions

$\text{NP}::\text{NP} \rightarrow [\text{les voitures bleues}]::[\text{blue cars}]$

$\text{NP}::\text{NP} \rightarrow [\text{les voitures } A^2]::[\text{JJ}^2 \text{ cars}]$

$\text{NP}::\text{NP} \rightarrow [\text{les } N^1 \text{ bleues}]::[\text{blue NNS}^1]$

$\text{NP}::\text{NP} \rightarrow [\text{les } N^1 A^2]::[\text{JJ}^2 \text{ NNS}^1]$

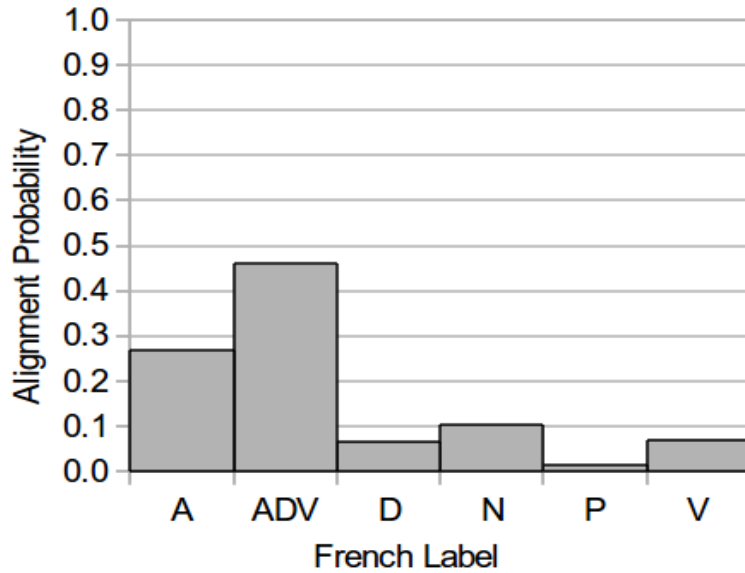
## (2) Label Collapsing

- Spurious ambiguity and rule sparsity caused by large label sets
- Idea: Cluster and collapse the label set
- Intuition: Categories that translate differently still “deserve” different labels
  - English JJ vs. JJR and JJS

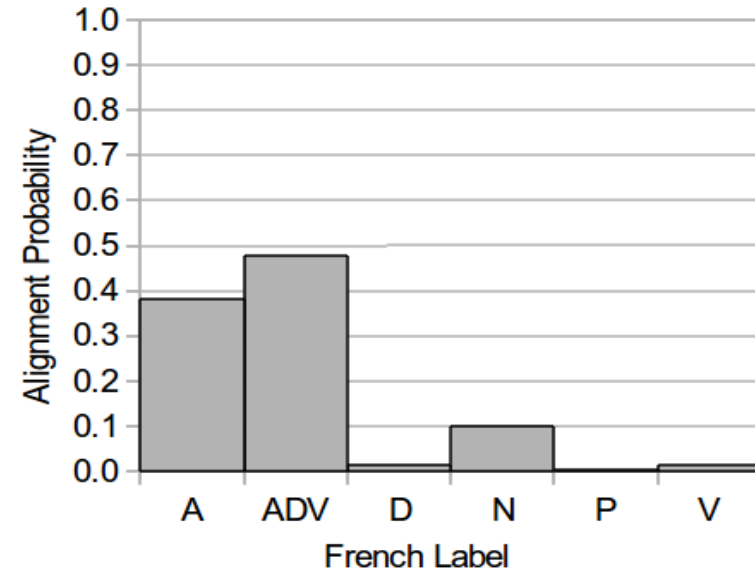
|                        |                                  |
|------------------------|----------------------------------|
| the <u>large</u> car   | la <u>grande</u> voiture         |
| the <u>larger</u> car  | la <u>plus grande</u> voiture    |
| the <u>largest</u> car | la voiture <u>la plus grande</u> |

# Alignment Distribution Difference

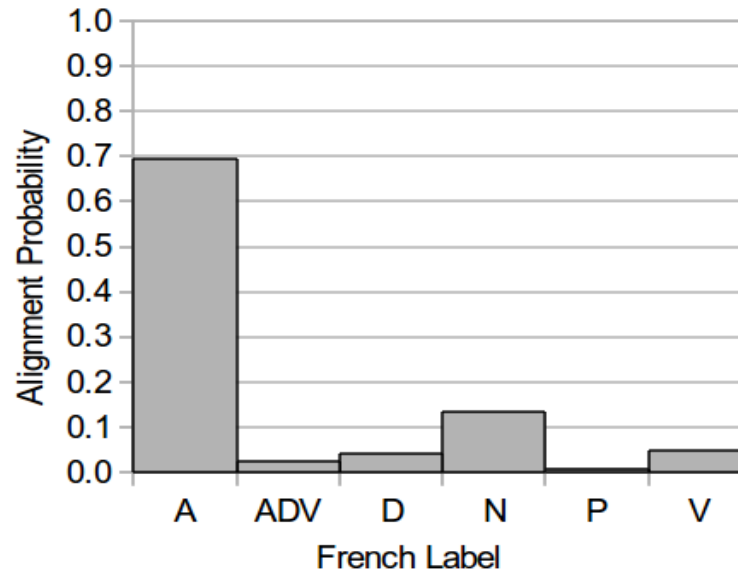
JJR



JJS



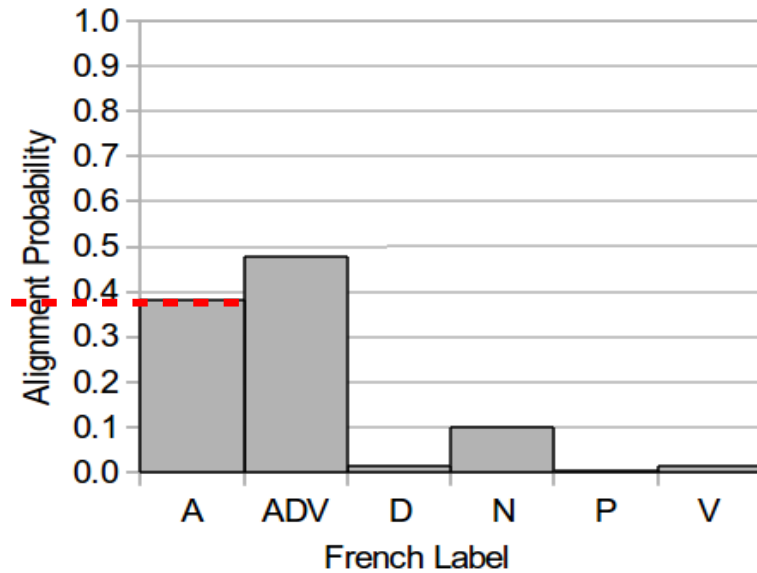
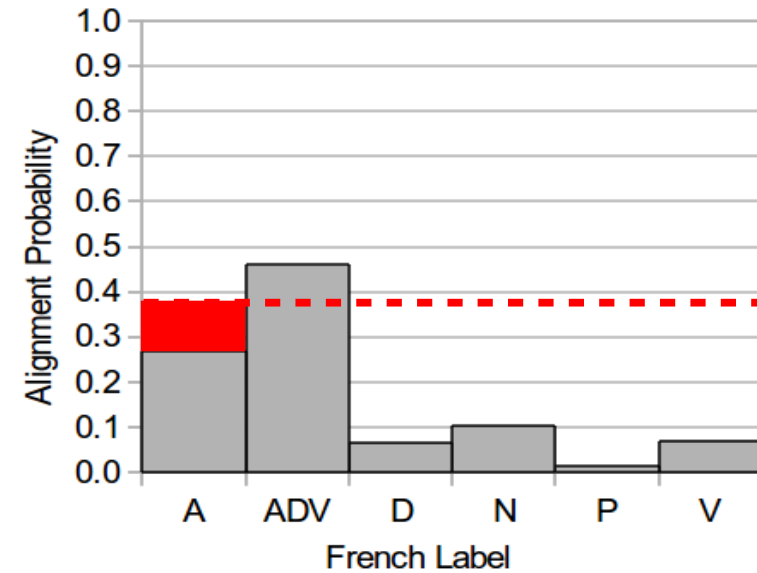
JJ



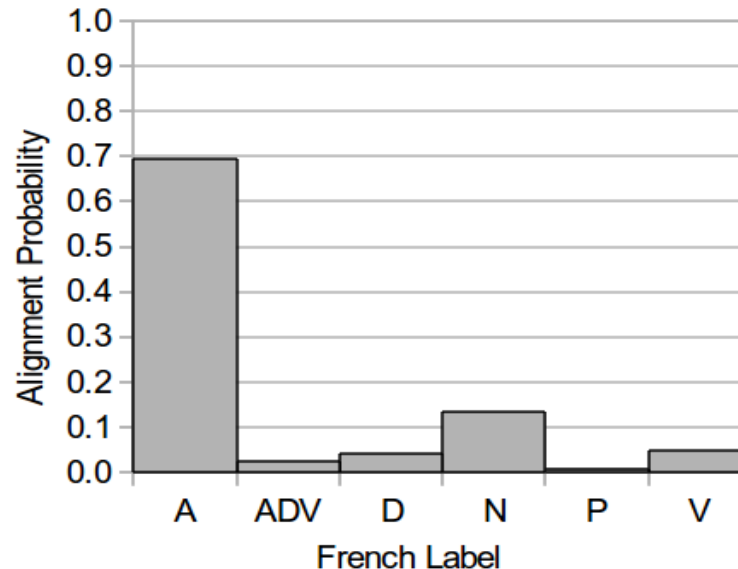
# Alignment Distribution Difference

JJR

JJS



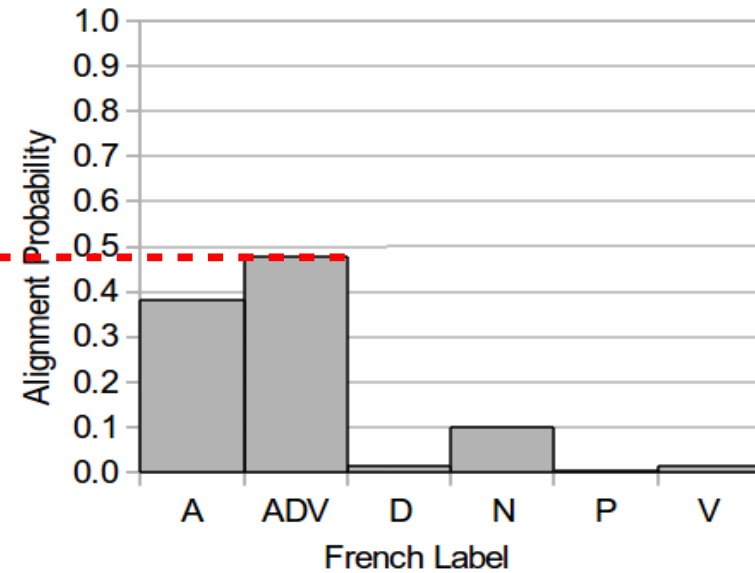
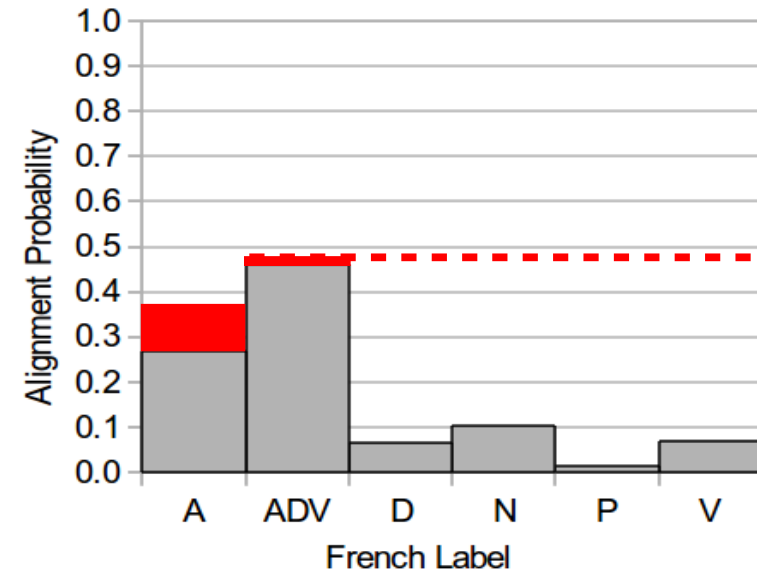
JJ



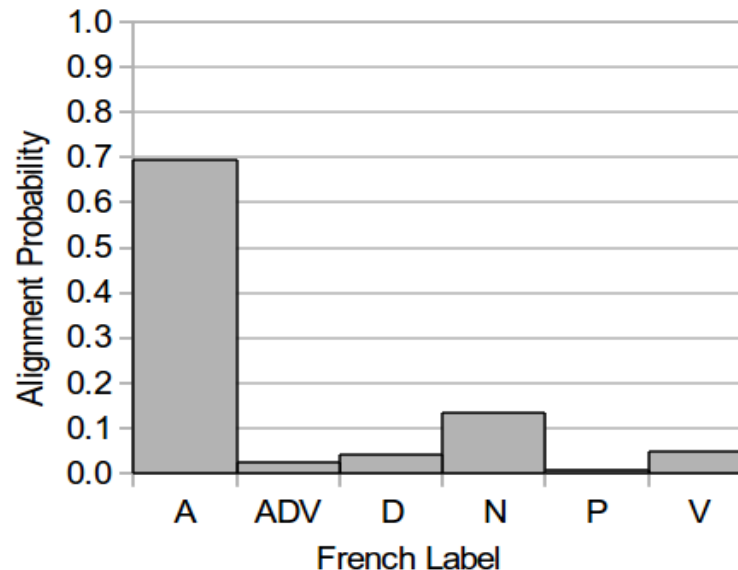
# Alignment Distribution Difference

JJR

JJS



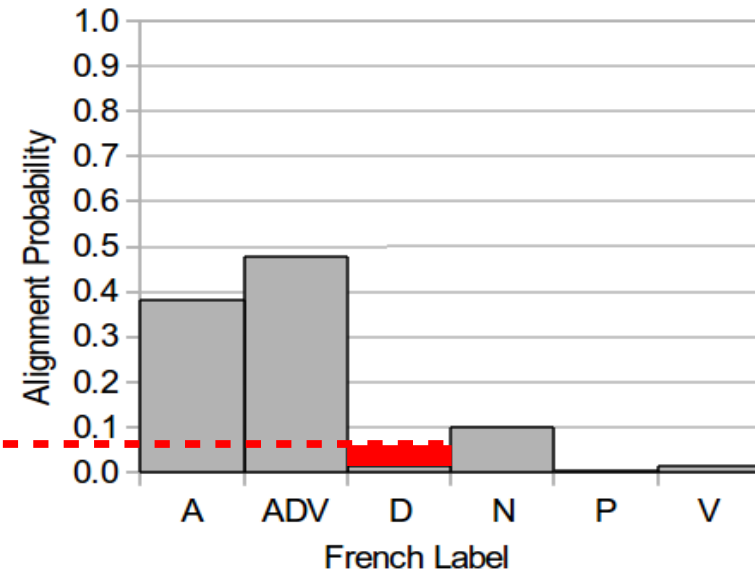
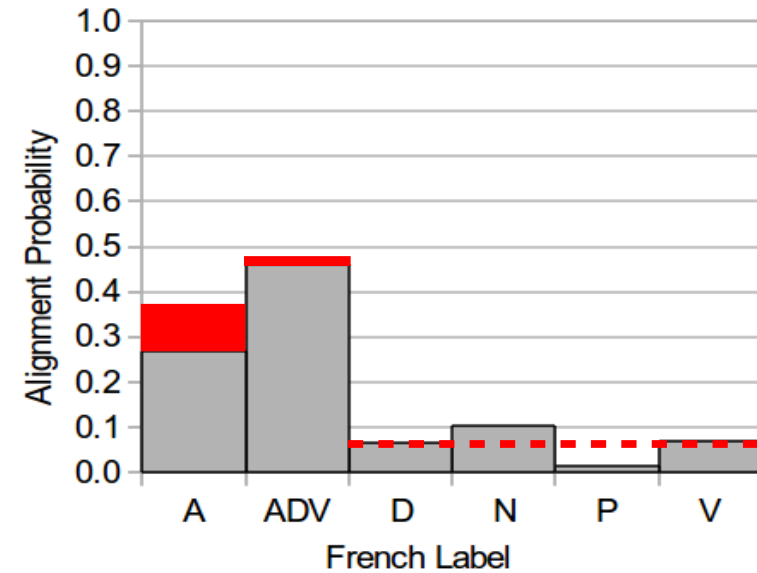
JJ



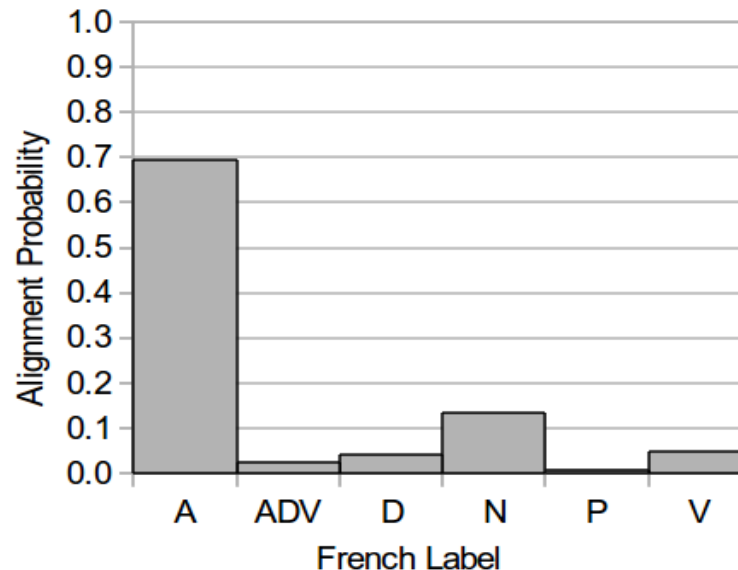
# Alignment Distribution Difference

JJR

JJS



JJ

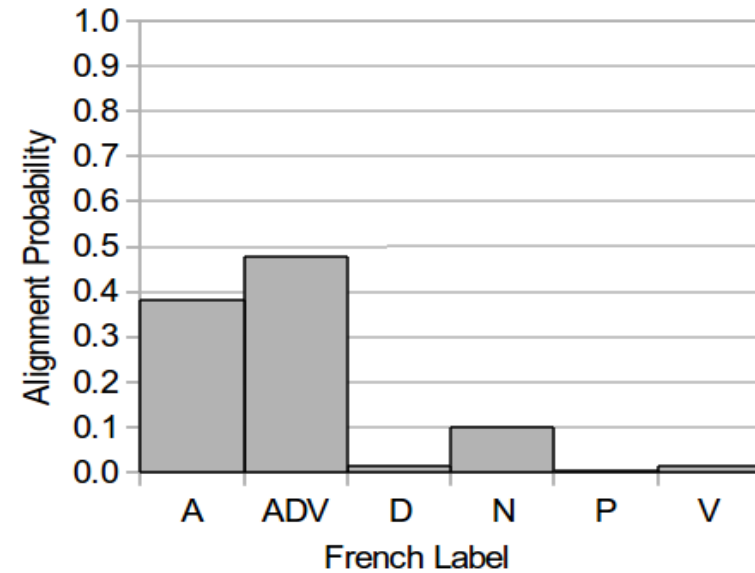
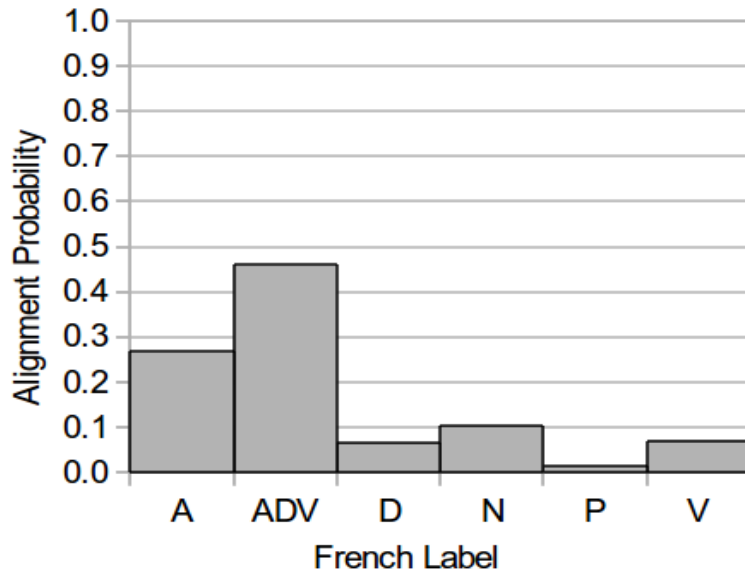


# Alignment Distribution Difference

JJR

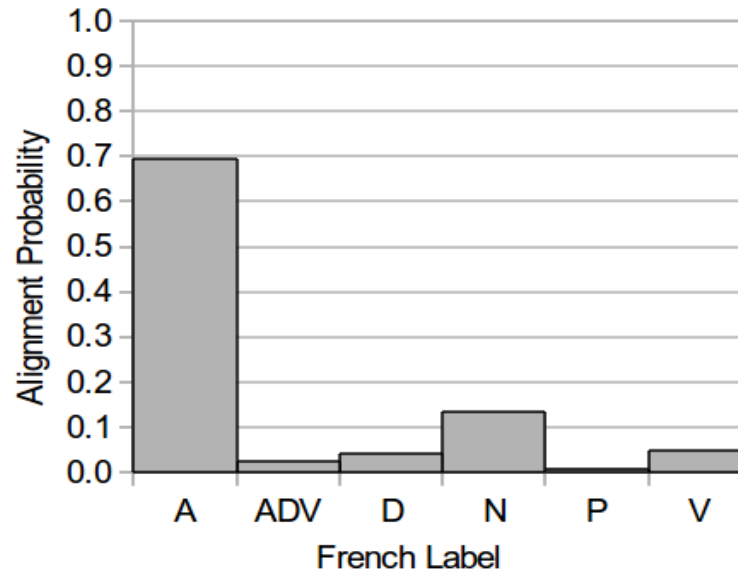
JJS

0.2688



JJ

0.9952

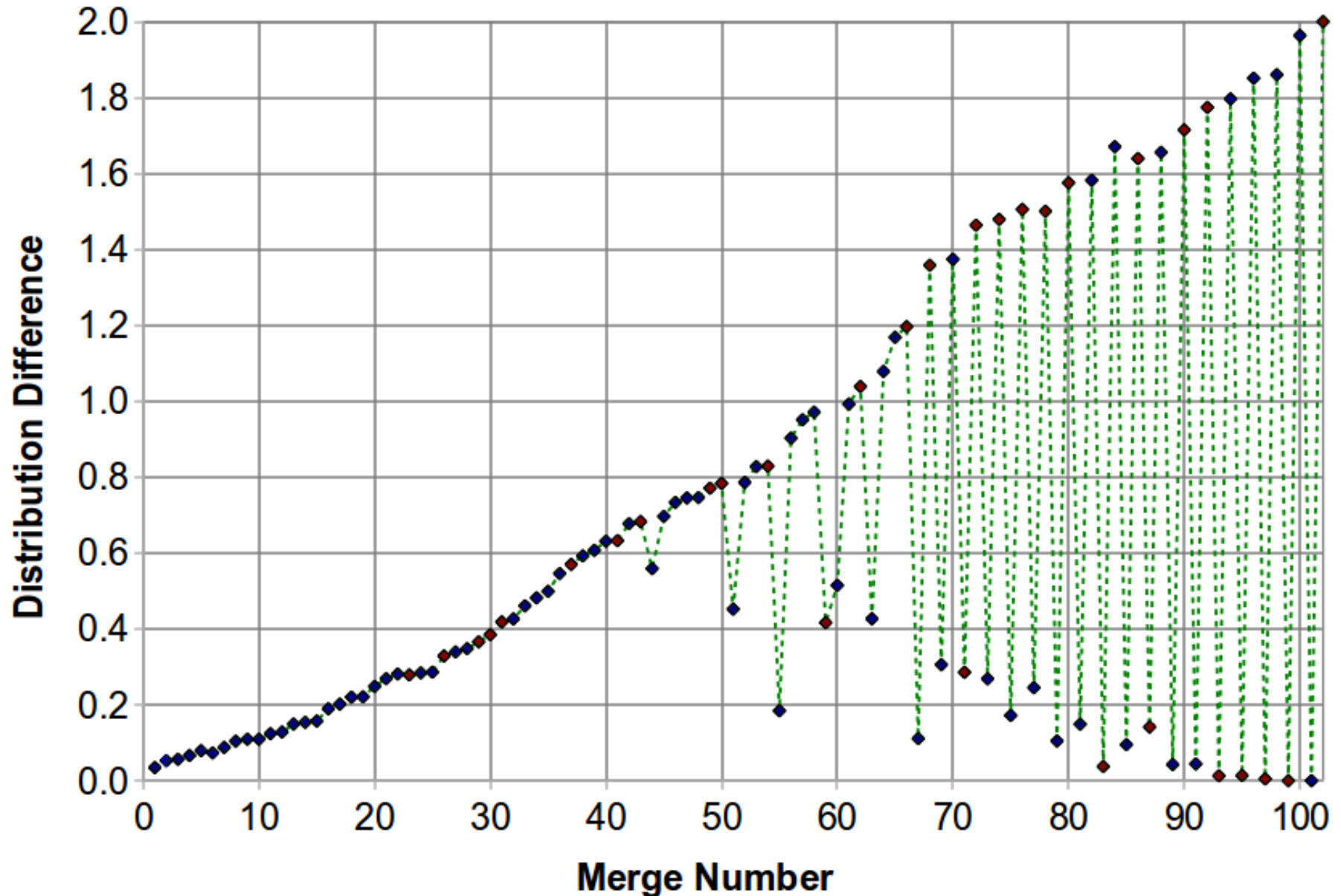


0.9114

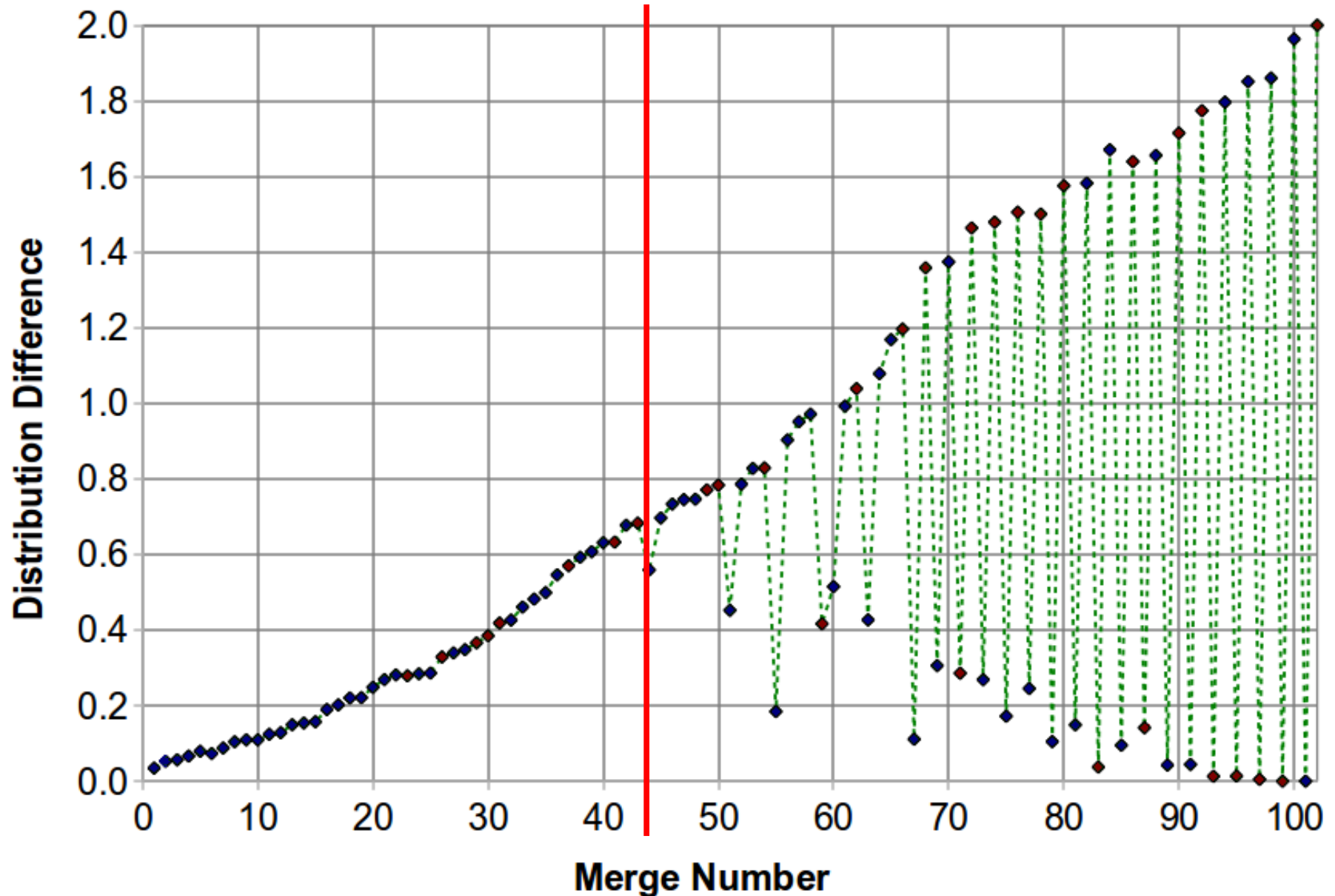
# Label Collapsing Algorithm

- Greedy procedure
  - Compute distribution difference between all pairs of source and target labels
  - Merge the label pair with smallest value
  - Update alignment statistics
  - Repeat
- Stopping criterion...?

# Label Collapsing Algorithm



# Label Collapsing Algorithm



# Results: Grammar

- Label set
  - Before: 33 French  $\times$  72 English  $\Rightarrow$  1134 joint
  - After: 25 French  $\times$  37 English  $\Rightarrow$  502 joint
- Spurious ambiguity
  - 1000 most frequent phrase pairs have 21% fewer left-hand-side labels
- Rule sparsity
  - Label sequences 7 times more likely to fit 10 frequent reordering patterns
  - More rule applications in MT test sets

# Results: MT Output

- Collapsing the label set improves scores
  - news-test2009

|                  | <b>BLEU</b>  | <b>METR</b>  | <b>TER</b>   |
|------------------|--------------|--------------|--------------|
| <b>Original</b>  | 21.75        | 51.96        | 59.44        |
| <b>Collapsed</b> | <u>22.78</u> | <u>53.00</u> | <u>59.16</u> |

– news-test2010

|                  | <b>BLEU</b>  | <b>METR</b>  | <b>TER</b>   |
|------------------|--------------|--------------|--------------|
| <b>Original</b>  | 22.03        | 53.13        | 58.00        |
| <b>Collapsed</b> | <u>23.14</u> | <u>54.03</u> | <u>57.86</u> |

# Proposed Work

- Investigate choice of
  - Stopping criterion
  - Distance metric
  - Greedy collapsing vs. all at once
- Run collapsing from different initial label granularities [Petrov et al. 2006]
  - Default: S, NP, NNS, A, VPinf
  - Refined: S-1, S-4, NP-2, NP-17, NNS-8, A-11, A-13, VPinf-3, etc.

# (3) Label Refining

- Bilingual SCFG rules encode knowledge not used in monolingual parsers
- Idea: Split existing labels using reordering information
- Intuition: Rules that differ only in target-side order indicate hidden categories

$$\text{NP}::\text{NP} \rightarrow [\text{DT}^1 \text{ JJ}^2 \text{ NN}^3]::[\text{D}^1 \text{ N}^3 \text{ A}^2]$$

$$\text{NP}::\text{NP} \rightarrow [\text{DT}^1 \text{ JJ}^2 \text{ NN}^3]::[\text{D}^1 \text{ A}^2 \text{ N}^3]$$

# Proposed Algorithm

- Identify rule set of interest

$$\text{NP}::\text{NP} \rightarrow [\text{DT}^1 \text{ JJ}^2 \text{ NN}^3]::[\text{D}^1 \text{ N}^3 \text{ A}^2]$$
$$\text{NP}::\text{NP} \rightarrow [\text{DT}^1 \text{ JJ}^2 \text{ NN}^3]::[\text{D}^1 \text{ A}^2 \text{ N}^3]$$

- From parsed corpus, tabulate what plugs into the rules at each extracted instance
- Items much more likely to plug into one rule than other form a category subtype

# Refining Example

$\text{NP}::\text{NP} \rightarrow [\text{DT}^1 \text{JJ}^2 \text{NN}^3]::[\text{_____}]$

**D<sup>1</sup>** N<sup>3</sup> A<sup>2</sup>

|       |        |
|-------|--------|
| un    | 0.3165 |
| une   | 0.2654 |
| la    | 0.1288 |
| le    | 0.0815 |
| cette | 0.0428 |
| l'    | 0.0371 |
| ce    | 0.0368 |
| ...   | ...    |

**D<sup>1</sup>** A<sup>2</sup> N<sup>3</sup>

|       |        |
|-------|--------|
| une   | 0.2545 |
| un    | 0.2067 |
| la    | 0.1688 |
| le    | 0.1054 |
| cette | 0.0563 |
| Le    | 0.0515 |
| La    | 0.0435 |
| ...   | ...    |

# Refining Example

$\text{NP}::\text{NP} \rightarrow [\text{DT}^1 \text{JJ}^2 \text{NN}^3]::[\text{_____}]$

**D**<sup>1</sup> N<sup>3</sup> A<sup>2</sup>

|       |        |
|-------|--------|
| un    | 0.3165 |
| une   | 0.2654 |
| la    | 0.1288 |
| le    | 0.0815 |
| cette | 0.0428 |
| l'    | 0.0371 |
| ce    | 0.0368 |
| ...   | ...    |

0.4154



**D**<sup>1</sup> A<sup>2</sup> N<sup>3</sup>

|       |        |
|-------|--------|
| une   | 0.2545 |
| un    | 0.2067 |
| la    | 0.1688 |
| le    | 0.1054 |
| cette | 0.0563 |
| Le    | 0.0515 |
| La    | 0.0435 |
| ...   | ...    |

# Refining Example

$\text{NP}::\text{NP} \rightarrow [\text{DT}^1 \text{JJ}^2 \text{NN}^3]::[\text{_____}]$

$\text{D}^1 \text{N}^3 \text{A}^2$

|            |        |
|------------|--------|
| rôle       | 0.0349 |
| position   | 0.0241 |
| question   | 0.0213 |
| communauté | 0.0212 |
| solution   | 0.0172 |
| situation  | 0.0171 |
| base       | 0.0154 |
| ...        | ...    |

0.8981



$\text{D}^1 \text{A}^2 \text{N}^3$

|          |        |
|----------|--------|
| fois     | 0.0397 |
| temps    | 0.0355 |
| point    | 0.0258 |
| majorité | 0.0212 |
| question | 0.0199 |
| chose    | 0.0193 |
| problème | 0.0174 |
| ...      | ...    |

# Refining Example

$NP::NP \rightarrow [DT^1 JJ^2 NN^3]::[ \text{_____} ]$

$D^1 N^3 A^2$

|                |        |
|----------------|--------|
| commune        | 0.0338 |
| important      | 0.0286 |
| politique      | 0.0285 |
| internationale | 0.0267 |
| européenne     | 0.0218 |
| européen       | 0.0186 |
| juridique      | 0.0144 |
| ...            | ...    |

1.8703



$D^1 A^2 N^3$

|          |        |
|----------|--------|
| même     | 0.0957 |
| nouvelle | 0.0837 |
| première | 0.0676 |
| bonne    | 0.0593 |
| nouveau  | 0.0521 |
| deuxième | 0.0509 |
| bon      | 0.0406 |
| ...      | ...    |

# Refining Example

$\text{NP}::\text{NP} \rightarrow [\text{DT}^1 \text{JJ}^2 \text{NN}^3]::[\text{D}^1 \text{N}^3 \text{AA}^2]$

$\text{NP}::\text{NP} \rightarrow [\text{DT}^1 \text{JJ}^2 \text{NN}^3]::[\text{D}^1 \text{AB}^2 \text{N}^3]$

**AA:** commune, politique, internationale,  
européenne, européen, juridique, ...

**AA & AB:** important, importante, ...

**AB:** même, nouvelle, première, bonne,  
nouveau, deuxième, bon, autre, ...

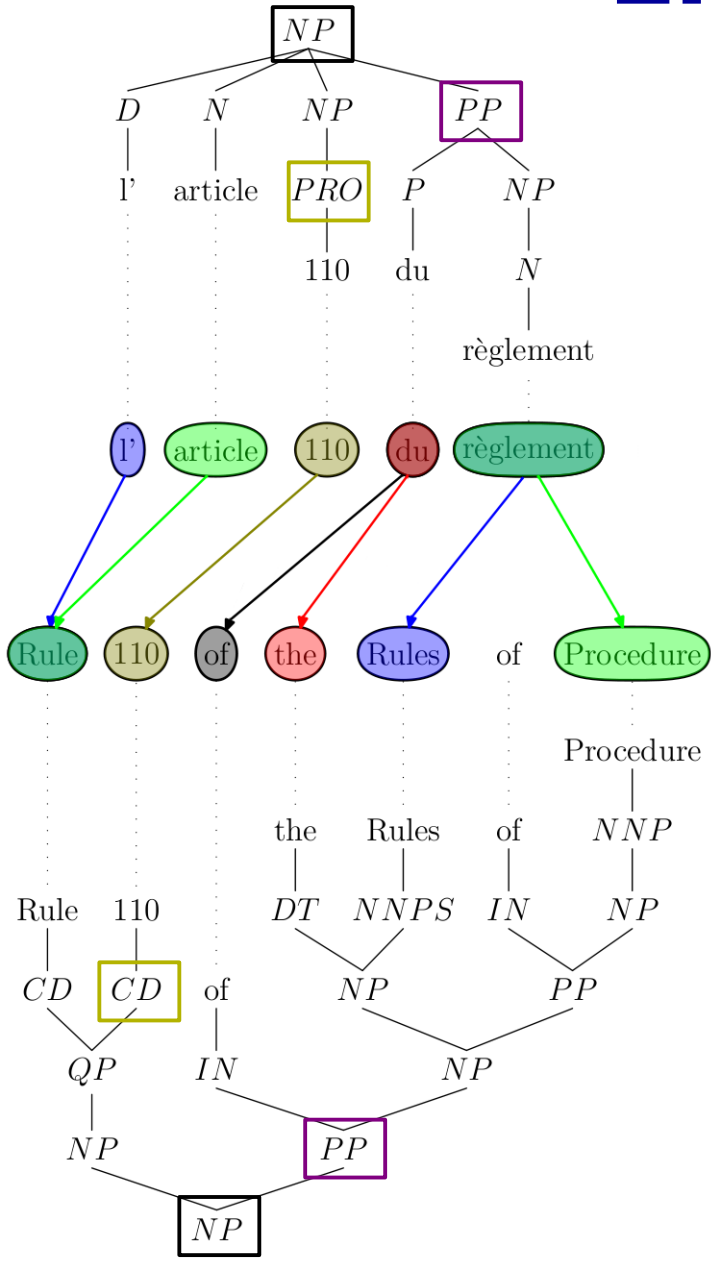
## **(4) Correcting Local Label Errors**

- Sentences parsed independently may contain local labeling errors
- Idea: Correct them using rules extracted from the entire parsed parallel corpus
- Intuition: Rule probabilities from whole corpus can fix label errors; fixing local label errors improves rule probabilities

# Proposed Algorithm

- EM algorithm
  - Initial grammar extraction:  
compute  $P(\text{LHS} \mid \text{RHS})$  for each rule
  - E step: Re-derive most likely labeled version of each aligned tree pair using rule scores
  - M step: Re-extract grammar from modified trees and re-compute rule scores
  - Repeat E and M steps until grammar doesn't change

# EM Example

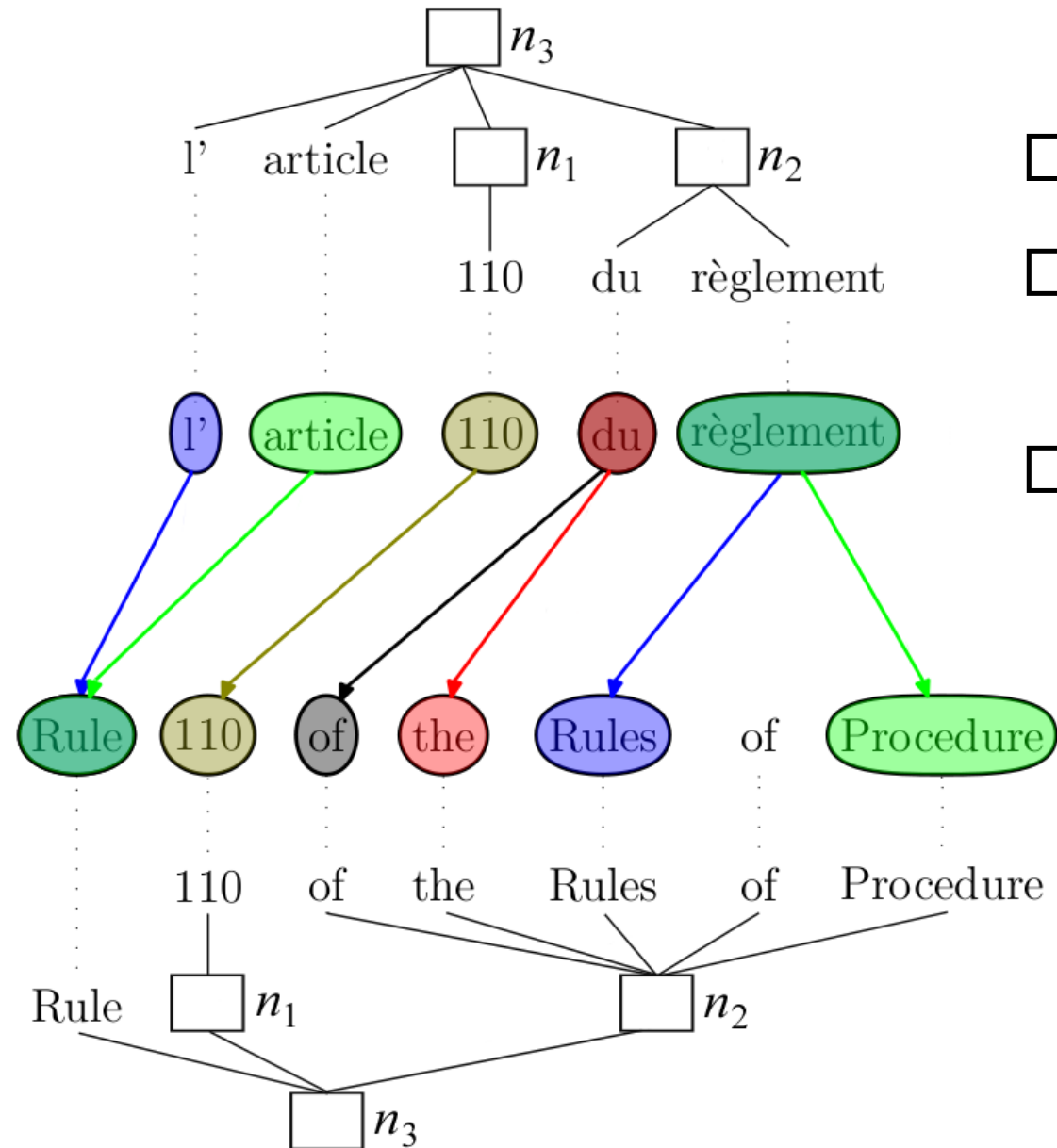


**PRO**::CD  $\rightarrow$  [110]::[110]

PP::PP  $\rightarrow$  [du règlement]::  
[of the Rules of Procedure]

NP::NP  $\rightarrow$  [l' article **PRO**<sup>1</sup> PP<sup>2</sup>]::  
[Rule CD<sup>1</sup> PP<sup>2</sup>]

# EM Example



$\square :: \square \rightarrow [110] :: [110]$

$\square :: \square \rightarrow [\text{du règlement}] :: [\text{of the R. of P.}]$

$\square :: \square \rightarrow [l' \text{ article } \square^1 \square^2] :: [\text{Rule } \square^1 \square^2]$

# EM Example



$n_1$ :  $\square::\square \rightarrow [110]::[110]$

D::CD      0.5081

PRO::CD    0.2789

N::CD      0.1004

...

...

$n_2$ :  $\square::\square \rightarrow [\text{du règlement}]::[\text{of the Rules of Procedure}]$

PP::PP      0.7736

NP::PP      0.2264

# EM Example

$$P(\text{LHS} \mid \text{RHS}) \cdot P(\text{RHS})$$

$n_3$ :  $\square :: \square \rightarrow [I' \text{ article } \square^1 \square^2] :: [\text{Rule } \square^1 \square^2]$

$\text{NP} :: \text{NP} \rightarrow [I' \text{ article } \text{PRO}^1 \text{PP}^2] :: [\text{Rule } \text{CD}^1 \text{PP}^2]$

$$1 \cdot 0.2789 \cdot 0.7736 = 0.2158$$

|         |        |
|---------|--------|
| D::CD   | 0.5081 |
| PRO::CD | 0.2789 |
| N::CD   | 0.1004 |
| ...     | ...    |

|        |        |
|--------|--------|
| PP::PP | 0.7736 |
| NP::PP | 0.2264 |

# EM Example

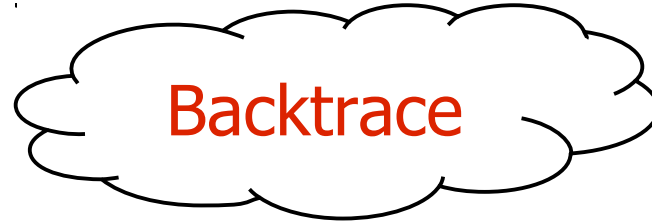

$$P(\text{LHS} \mid \text{RHS}) \cdot P(\text{RHS})$$

$n_3$ :  $\square :: \square \rightarrow [l' \text{ article } \square^1 \square^2] :: [\text{Rule } \square^1 \square^2]$

NP::NP      0.7018

NP::VP      0.0023

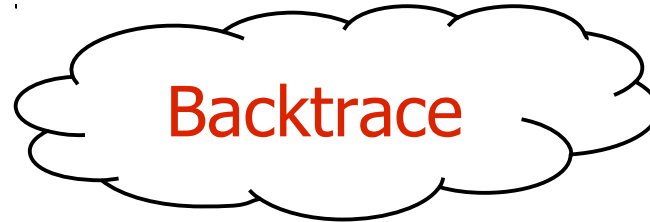
# EM Example



$n_3$ : NP::NP  $\rightarrow$  [l' article  $\square^1 \square^2$ ]::[Rule  $\square^1 \square^2$ ]

|        |        |
|--------|--------|
| NP::NP | 0.7018 |
| NP::VP | 0.0023 |

# EM Example



$n_3$ : NP::NP  $\rightarrow$  [I' article D<sup>1</sup> PP<sup>2</sup>]::[Rule CD<sup>1</sup> PP<sup>2</sup>]

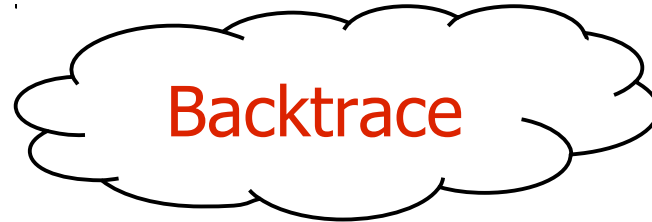
NP::NP      0.7018

NP::VP      0.0023

NP::NP  $\rightarrow$  [I' article D<sup>1</sup> PP<sup>2</sup>]::[Rule CD<sup>1</sup> PP<sup>2</sup>]

1 · 0.5081 · 0.7736 = 0.3931

# EM Example



$n_3$ : NP::NP  $\rightarrow$  [l' article D<sup>1</sup> PP<sup>2</sup>]::[Rule CD<sup>1</sup> PP<sup>2</sup>]

NP::NP      0.7018

NP::VP      0.0023

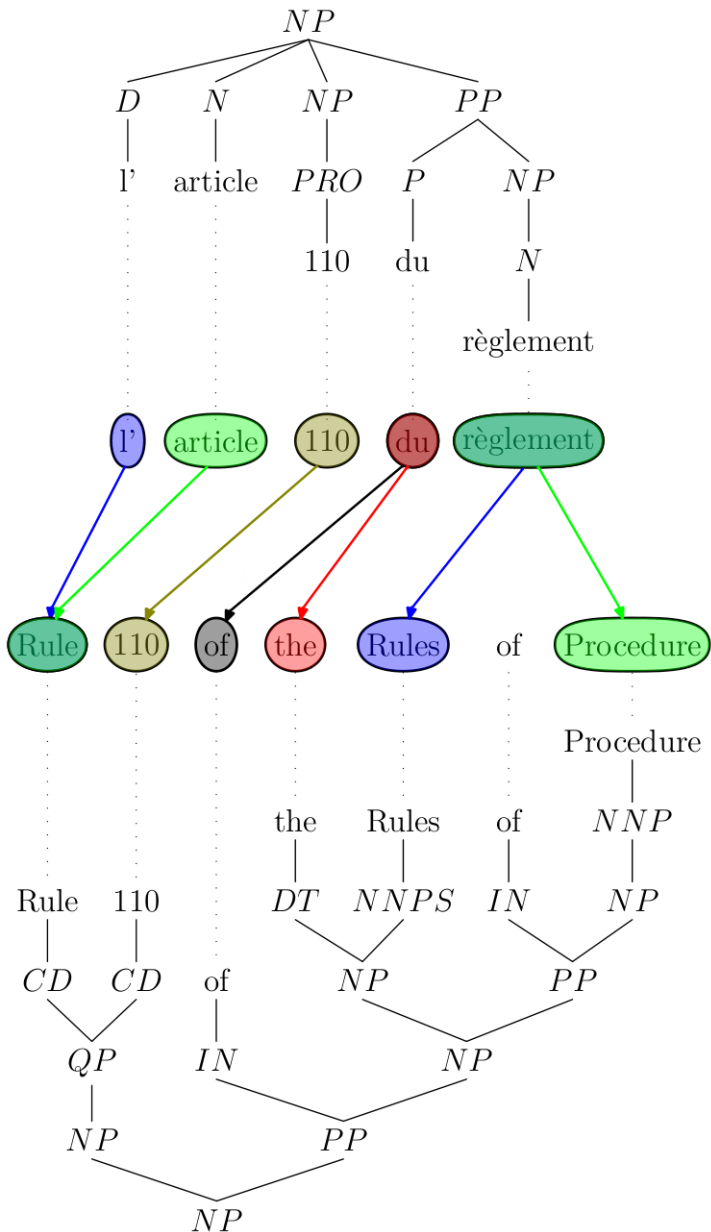
NP::NP  $\rightarrow$  [l' article D<sup>1</sup> PP<sup>2</sup>]::[Rule CD<sup>1</sup> PP<sup>2</sup>]

$1 \cdot 0.5081 \cdot 0.7736 = 0.3931$

$n_1$ : D::CD  $\rightarrow$  [1 10]::[1 10]

$n_2$ : PP::PP  $\rightarrow$  [du règlement]::[of the Rules of Procedure]

# EM Example


$$\mathbf{D}::\mathbf{CD} \rightarrow [110]::[110]$$

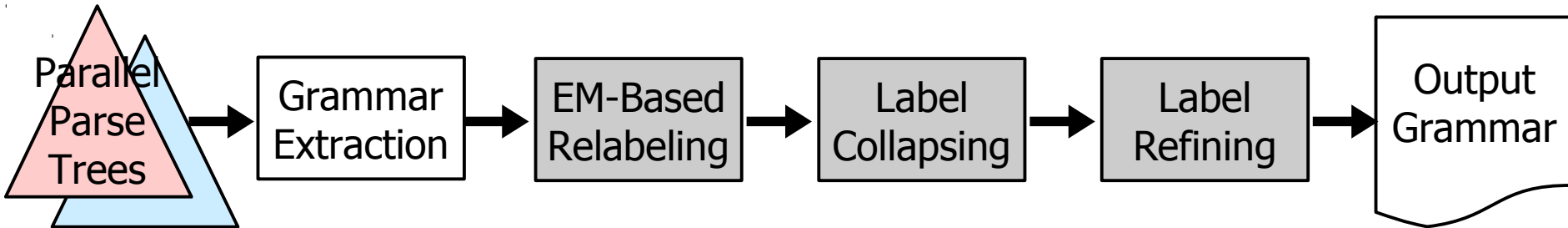
PP::PP → [du règlement]::  
[of the Rules of Procedure]

$$\text{NP}::\text{NP} \rightarrow [1' \text{ article } \mathbf{D}^1 \text{ PP}^2]:: [\text{Rule CD}^1 \text{ PP}^2]$$

# Outline

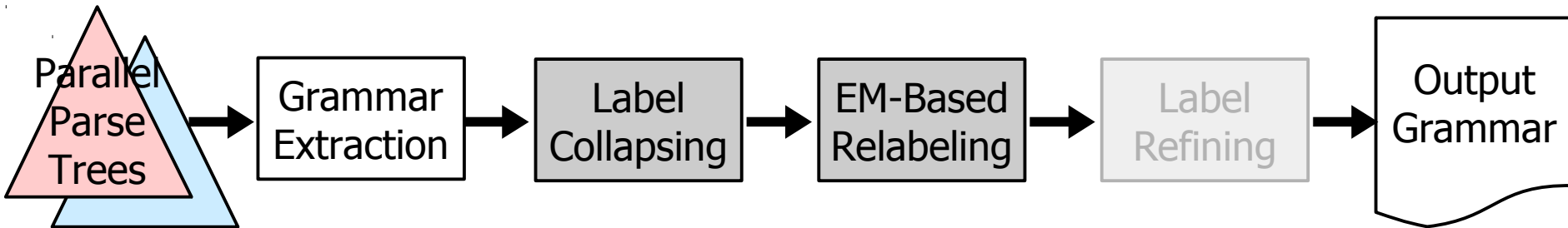
- Introduction
- Baseline system and experimental setup
- Label-based system properties
- Extraction and relabeling techniques
- Putting it all together

# Proposed Combination: Default



- Large-scale system with standard labels
  - Remove easy-to-fix occasional labeling errors first (–SA)
  - Remove redundant labels next (–SA, –RS)
  - Add necessary latent categories (+RP)

# Proposed Combination: Variants



- Fine-grained input categories
  - Move label collapsing to first (–fragmentation)
  - Remove label refining (focus on RS)
- Small-data systems
  - Remove label refining (focus on RS)

# Work Summary

|                                                                | Completed Work                          | Proposed Work                                                 |
|----------------------------------------------------------------|-----------------------------------------|---------------------------------------------------------------|
| <b>Spurious ambiguity, rule sparsity, reordering precision</b> | Defined, measured by simple counting    | Measured within system, constrained decoding                  |
| <b>General-purpose grammar extractor</b>                       | Node alignment stage                    | Rule extraction stage                                         |
| <b>Label collapsing</b>                                        | Basic algorithm, end-to-end experiments | Algorithm extensions                                          |
| <b>Label refining</b>                                          | Motivational example, basic approach    | Scoping, end-to-end experiments                               |
| <b>EM-based relabeling</b>                                     | Motivational example, basic approach    | Implementation with grammar extractor, end-to-end experiments |
| <b>Combined relabeling scheme</b>                              | Default proposals for different systems | Experimentation                                               |

# Timeline

| Time Period            | Work                                                                                                                                                                  |
|------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Jan.-March 2011</b> | Implementation of grammar extractor<br>Full label collapsing experiments<br>Tests of extensions to label collapsing algorithm<br>Submission of EMNLP paper (March 23) |
| <b>April-July 2011</b> | Development of constrained decoding setup<br>Development of label refining techniques<br>Full label refining experiments<br>French-English submission to WMT 2011     |
| <b>Aug.-Oct. 2011</b>  | Development of EM-based relabeling framework<br>Full EM-based relabeling experiments<br>Arabic or Chinese submission to NIST 2011                                     |
| <b>Nov.-Dec. 2011</b>  | Finalization of all individual-technique experiments                                                                                                                  |
| <b>Jan.-March 2012</b> | Development of combined-technique pipeline<br>Finalization of systems and experiments                                                                                 |
| <b>April-May 2012</b>  | Thesis writing<br>Thesis defense                                                                                                                                      |



# Grammar Pruning

- Hiero
  - Phrase pairs limited to length 10
  - Maximum **two non-adjacent** nonterminals in rules
  - No fully abstract rules
- GHKM
  - Phrase pairs filtered to **sentence** level
  - Save  $n$  most frequent translations **per source RHS**, or all rules extracted at least  $k$  times ( $n = 40, k = 50$ )
- Our systems
  - Phrase pairs filtered to **test set** level
  - Save  $n$  most **globally** frequent rules ( $n = 10,000$ )

# The State of the Art

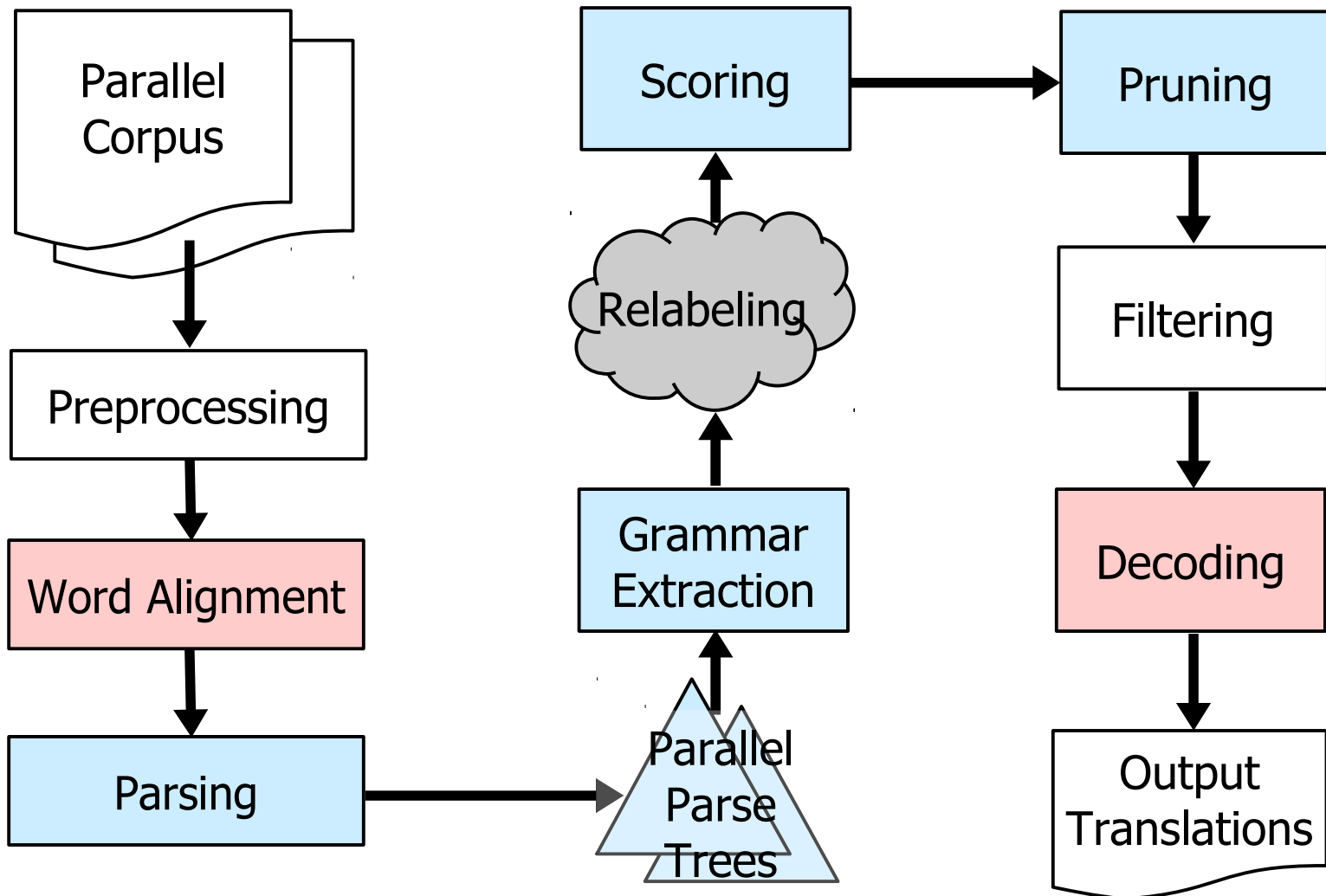
- WMT news-test2010 test set

|                             | BLEU        |      |
|-----------------------------|-------------|------|
| Best Syntax-Only Baseline   | 24.1        | +2.1 |
| Best Overall Baseline       | 26.4        | +1.2 |
| <b>Best WMT 2010 System</b> | <b>29.6</b> |      |

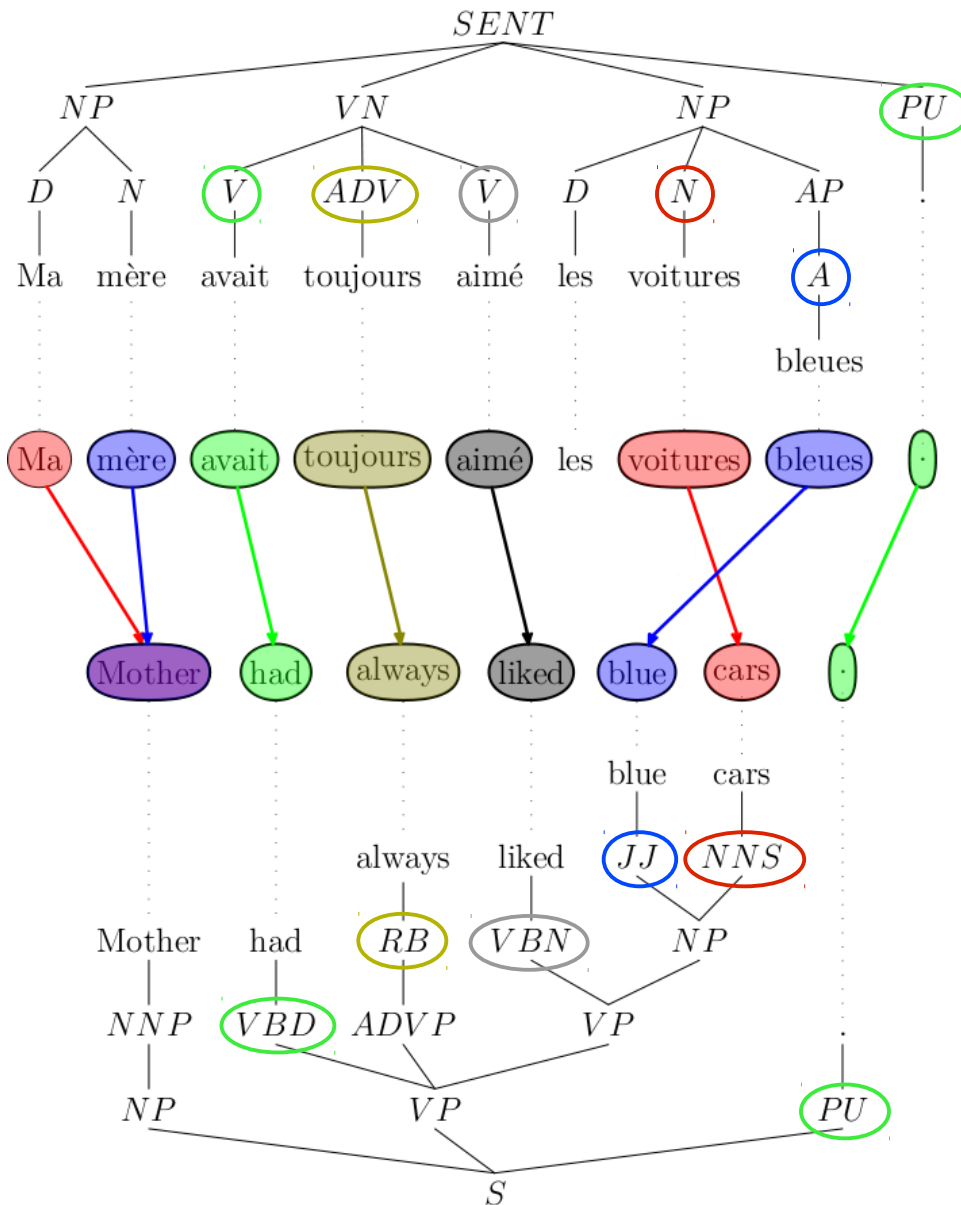
- Syntax-only better for grammar analysis
- Expected improvements from improved grammar pruning and rule extraction

# Computing Demands

 = Local server     = Parallel cluster     = MapReduce cluster



# Baseline Extraction Method



V::VBD → [avait]::[had]

ADV::RB → [toujours]::  
[always]

V::VBN → [aimé]::[liked]

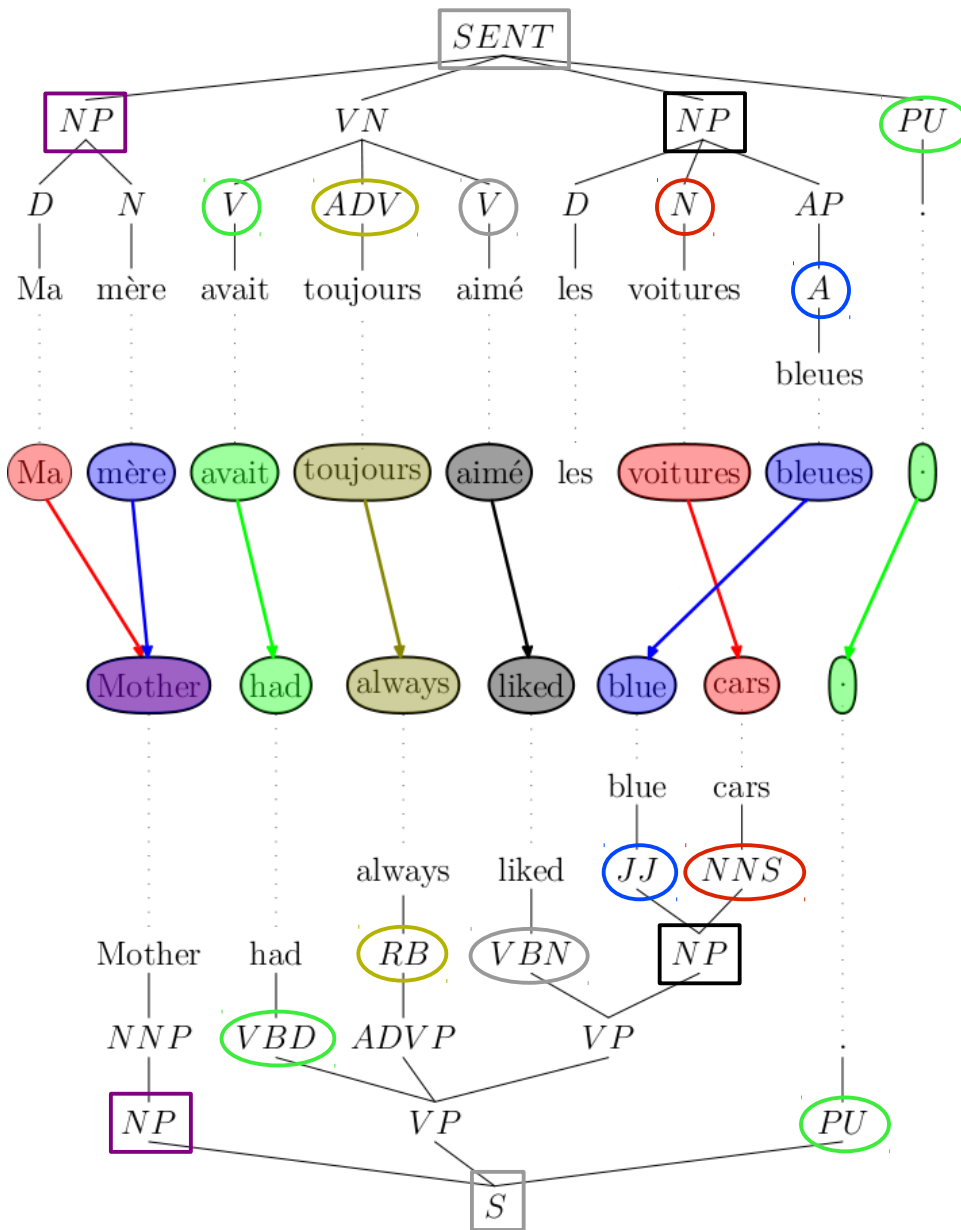
N::NNS → [voitures]::[cars]

A::JJ → [bleues]::[blue]

PU::PU → [.]::[.]

[Lavie et al. 2008]

# Baseline Extraction Method



NP::NP → [Ma mère]::  
[Mother]

NP::NP → [les voitures bleues]::  
[blue cars]

NP::NP → [les N<sup>1</sup> A<sup>2</sup>]::  
[JJ<sup>2</sup> NNS<sup>1</sup>]

SENT::S →  
[NP<sup>1</sup> V<sup>2</sup> ADV<sup>3</sup> V<sup>4</sup> NP<sup>5</sup> PU<sup>6</sup>]::  
[NP<sup>1</sup> VBD<sup>2</sup> RB<sup>3</sup> VBN<sup>4</sup> NP<sup>5</sup> PU<sup>6</sup>]

[Lavie et al. 2008]