

mt1(形態素解析)の大まかな動き「最長一致法」

- 辞書ファイルオープン
- 対訳書き込みファイルオープン
- キーボード(stdin)から文入力
- 初回のループは「次」=「文」
- whileループ「次」が空になるまで
 - 「次」が辞書ファイルの当該行の最初のフィールド(\$dicfield[0])にマッチしたら、
 - 対訳フィールドを画面表示
 - 辞書の当該行をファイル書き込み
 - 「次」に「(マッチした部分の)残り」を代入
 - そうでなければ「次」の末尾1文字をカット(chop)
- ファイルをクローズ

\$bun=

hablo español.

\$tugi=

hablo español.

hablo español

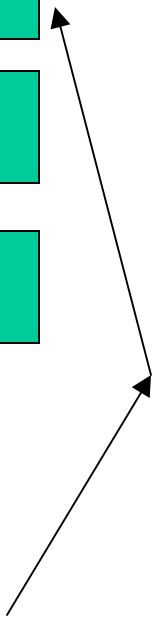
hablo españo

辞書マッチ⇒

hablo

\$nokori=

español.



mt2(構文解析・語順変換)の大まかな流れ

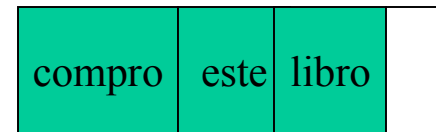
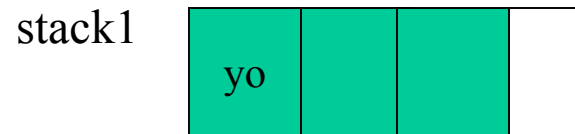
- 対訳ファイル(形態素解析出力)オープン⇒配列にセット(配列番号が形態素番号＝何番目の単語か。0始まり)
- 文頭形態素から順に取り出し、スペイン語品詞を見る。以下を文末形態素まで行う。
 - スタック1開放のトリガーになる品詞＝V(類)または句読点(punto, coma)のとき
 - VならばS1を開放(後ろから前に訳出*)し、自分自身は新たにスタック1に積む。
 - 句読点類ならばS1を開放*し、自分自身も訳出
 - * 途中で「前⇒後ろの順で訳出する箇所」は別途スタック2に
 - トリガー品詞でない場合は、スタック1に積む

基本は(①←V②←)yo compro este libro.

形態素配列 0 1 2 3 4

* スタック(先入れ後出しのための仕組み)

トリガー品詞が出るまではスタック1(後⇒前訳出用スタック)に積む(push)



前⇒後訳出部分はS1開放(pop)時にスタック2に積む

