

Lecture 4:

Machine Translation and IBM Model I (Learning)

Intro to NLP, CS585, Fall 2014

<http://people.cs.umass.edu/~brenocon/inlp2014/>

Brendan O'Connor (<http://brenocon.com>)

- Homework questions?
- Today
 - Quick LM review
 - Machine translation -- learning
- Next week
 - Machine translation -- decoding, broader issues

Interpolation LM [from last time]

Let $\hat{P}$ be the MLE probability distribution

Interpolation estimate is:

$$P(w|w_{-1}, w_{-2}) = \lambda_2 \hat{P}(w|w_{-1}, w_{-2}) + \lambda_1 \hat{P}(w|w_{-1}) + \lambda_0 \hat{P}(w)$$

↑
Sparse

↑
Denser

↑
Dense

Mixing weights: $\lambda_0 + \lambda_1 + \lambda_2 = 1$

$P(w|\text{blacke as})$

Hamlet from NLTK, V=4812

http://people.cs.umass.edu/~brenocon/inlp2014/lectures/04-hamlet_ngrams.txt

$$= \lambda_2 \hat{P}(w|\text{blacke as}) + \lambda_1 \hat{P}(w|\text{as}) + \lambda_0 \hat{P}(w)$$

$$P(w|\text{blacke as})$$

$$= \lambda_2 \hat{P}(w|\text{blacke as}) + \lambda_1 \hat{P}(w|\text{as}) + \lambda_0 \hat{P}(w)$$



3 blacke as

- I blacke as death
- I blacke as hell
- I blacke as his

3 non-zeros

$P(w|\text{blacke as})$

Hamlet from NLTK, V=4812

http://people.cs.umass.edu/~brenocon/inlp2014/lectures/04-hamlet_ngrams.txt

$$= \lambda_2 \hat{P}(w|\text{blacke as}) + \lambda_1 \hat{P}(w|\text{as}) + \lambda_0 \hat{P}(w)$$



3 blacke as
 1 blacke as death
 1 blacke as hell
 1 blacke as his

3 non-zeros



205 as	7 as he	1 as peace
1 as 'twer	1 as	1 as proper
6 as 'twere	healthfull	1 as pure
3 as a	1 as hell	1 as quick-
1 as actiuely	3 as her	siluer
1 as against	1 as his	1 as reason
1 as all	1 as how	1 as sharpe
1 as an	1 as hush	1 as sinewes
2 as any	8 as i	1 as sinnes
1 as are	1 as ice	3 as snow
1 as bad	6 as if	1 as so
2 as before	1 as in	1 as some
3 as by	8 as it	1 as swift
1 as chast	1 as iust	1 as th'art
1 as	2 as kill	13 as the
checking	1 as kinde	1 as therein
1 as	1 as leuell	2 as they
common	1 as liue	3 as this
1 as cunning	1 as loue	3 as thou
1 as damn	1 as low	1 as thus
1 as day	1 as lying	1 as thy
2 as death	1 as mad	5 as to
1 as deepe	1 as made	1 as
1 as dicers	1 as make	vnuallued
1 as doth	1 as many	1 as vulcans
1 as easie	2 as may	1 as
1 as england	1 as	watchmen
1 as ere	meditation	1 as waxe
1 as false	1 as mine	4 as we
1 as farre	1 as mortall	2 as well
1 as fits	1 as most	1 as white
1 as foule	3 as much	2 as will
1 as fresh	2 as my	1 as
2 as from	1 as needfull	womans
1 as gaming	2 as of	1 as would
1 as girdle	1 as oft	12 as you
1 as had	2 as one	2 as your
1 as hamlet	2 as our	1 as yours
1 as hardy	1 as patient	

107 non-zeros

$P(w|blacke as)$

Hamlet from NLTK,V=4812

http://people.cs.umass.edu/~brenocon/inlp2014/lectures/04-hamlet_ngrams.txt

$= \lambda_2 \hat{P}(w|blacke as) + \lambda_1 \hat{P}(w|as) + \lambda_0 \hat{P}(w)$



3	blacke as
1	blacke as death
1	blacke as hell
1	blacke as his

3 non-zeros

205	as	7	as he	1	as peace
1	as 'twer	1	as	1	as proper
6	as 'twere	1	healthfull	1	as pure
3	as a	1	as hell	1	as quick-
1	as actiuely	3	as her	1	siluer
1	as against	1	as his	1	as reason
1	as all	1	as how	1	as sharpe
1	as an	1	as hush	1	as sinewes
2	as any	8	as i	1	as sinnes
1	as are	1	as ice	3	as snow
1	as bad	6	as if	1	as so
2	as before	1	as in	1	as some
3	as by	8	as it	1	as swift
1	as chast	1	as iust	1	as th'art
1	as	2	as kill	13	as the
1	checking	1	as kinde	1	as therein
1	as	1	as leuell	2	as they
1	common	1	as liue	3	as this
1	as cunning	1	as loue	3	as thou
1	as damn	1	as low	1	as thus
1	as day	1	as lying	1	as thy
2	as death	1	as mad	5	as to
1	as deepe	1	as made	1	as
1	as dicers	1	as make	1	vnuallued
1	as doth	1	as many	1	as vulcans
1	as easie	2	as may	1	as
1	as england	1	as	1	watchmen
1	as ere	1	meditation	1	as waxe
1	as false	1	as mine	4	as we
1	as farre	1	as mortall	2	as well
1	as fits	1	as most	1	as white
1	as foule	3	as much	2	as will
1	as fresh	2	as my	1	as
2	as from	1	as needfull	1	womans
1	as gaming	2	as of	1	as would
1	as girdle	1	as oft	12	as you
1	as had	2	as one	2	as your
1	as hamlet	2	as our	1	as yours
1	as hardy	1	as patient		

107 non-zeros

17	&	1	adoption	1	angle	1	attractive
51	'd	2	adores	1	angry	5	audience
200	'em	1	aduanc	1	animals	1	audit
2	'gainst	1	aducement	1	annexment	1	augury
3	'm	1	adue	1	annoint	1	aut
14	're	1	aduantage	1	annual	1	auoid
119	't	5	adue	6	anon	2	auouch
61	'tane	1	aduue	8	another	1	auoyd
1	'tis	1	aduuenturous	1	another	1	auspicious
1	'twas	3	aduice	1	answer	2	author
1	'twere	1	aduise	4	answered	1	authorities
10	'twere	1	aduise	1	answers	1	awake
4	'twixt	1	aduiterate	1	answers	27	away
45	(	1	aduiterate	1	answers	2	awe
44	.	1	aduiterate	1	answers	4	awhile
2892	.	3	aduiterate	1	answers	2	axe
1879	.	1	aduiterate	1	answers	1	aye
3	play	1	aduiterate	1	answers	2	aye
1599	player	1	aduiterate	1	answers	1	ayre
566	.	1	aduiterate	1	answers	12	ayres
298	.	1	aduiterate	1	answers	1	ayrie
459	.	1	aduiterate	1	answers	1	ayry
6	.	1	aduiterate	1	answers	2	babe
497	a	1	aduiterate	1	answers	1	bace
2	a	1	aduiterate	1	answers	6	back
1	a	1	aduiterate	1	answers	1	backe
1	a	1	aduiterate	1	answers	7	backward
1	a	1	aduiterate	1	answers	1	bad
1	a	1	aduiterate	1	answers	2	bait
1	a	1	aduiterate	1	answers	1	bak
1	a	1	aduiterate	1	answers	1	bakers
1	a	1	aduiterate	1	answers	1	bakt-
1	a	1	aduiterate	1	answers	1	meats
1	a	1	aduiterate	1	answers	1	ban
1	a	1	aduiterate	1	answers	1	bands
1	a	1	aduiterate	1	answers	1	banke
1	a	1	aduiterate	1	answers	3	bap
1	a	1	aduiterate	1	answers	1	bapt
1	a	1	aduiterate	1	answers	7	baptista
1	a	1	aduiterate	1	answers	1	bar
1	a	1	aduiterate	1	answers	2	barbars
1	a	1	aduiterate	1	answers	2	barbar
1	a	1	aduiterate	1	answers	1	bare
1	a	1	aduiterate	1	answers	1	bare-foot
1	a	1	aduiterate	1	answers	10	barke
1	a	1	aduiterate	1	answers	8	barn
1	a	1	aduiterate	1	answers	1	barnardo
1	a	1	aduiterate	1	answers	1	barr'd
1	a	1	aduiterate	1	answers	1	barre
1	a	1	aduiterate	1	answers	4	barren
1	a	1	aduiterate	1	answers	1	base
1	a	1	aduiterate	1	answers	1	basenesse
1	a	1	aduiterate	1	answers	2	baser
1	a	1	aduiterate	1	answers	1	basket
1	a	1	aduiterate	1	answers	1	bastard
1	a	1	aduiterate	1	answers	1	bat
1	a	1	aduiterate	1	answers	1	bated
1	a	1	aduiterate	1	answers	1	battalians
1	a	1	aduiterate	1	answers	1	batten
1	a	1	aduiterate	1	answers	1	battery
1	a	1	aduiterate	1	answers	1	battlements
1	a	1	aduiterate	1	answers	1	bawdry
1	a	1	aduiterate	1	answers	191	bawdy
1	a	1	aduiterate	1	answers	1	be
1	a	1	aduiterate	1	answers	1	be-ratted
1	a	1	aduiterate	1	answers	1	beame
1	a	1	aduiterate	1	answers	1	beart
1	a	1	aduiterate	1	answers	6	beard
1	a	1	aduiterate	1	answers	13	beards
1	a	1	aduiterate	1	answers	1	beare
1	a	1	aduiterate	1	answers	2	beares
1	a	1	aduiterate	1	answers	1	beares
1	a	1	aduiterate	1	answers	5	beast
1	a	1	aduiterate	1	answers	2	beasts
1	a	1	aduiterate	1	answers	1	beaten
1	a	1	aduiterate	1	answers	3	beating
1	a	1	aduiterate	1	answers	1	beats
1	a	1	aduiterate	1	answers	1	beauer
1	a	1	aduiterate	1	answers	1	beauteous
1	a	1	aduiterate	1	answers	4	beautie
1	a	1	aduiterate	1	answers	1	beautied
1	a	1	aduiterate	1	answers	1	beauties
1	a	1	aduiterate	1	answers	1	beautified
1	a	1	aduiterate	1	answers	3	beauty
1	a	1	aduiterate	1	answers	1	became
1	a	1	aduiterate	1	answers	1	because
1	a	1	aduiterate	1	answers	1	becke
1	a	1	aduiterate	1	answers	1	beckens
1	a	1	aduiterate	1	answers	1	beckons
1	a	1	aduiterate	1	answers	2	becomes
1	a	1	aduiterate	1	answers	11	bed
1	a	1	aduiterate	1	answers	1	bedded
1	a	1	aduiterate	1	answers	16	bedrid
1	a	1	aduiterate	1	answers	3	bee
1	a	1	aduiterate	1	answers	13	been
1	a	1	aduiterate	1	answers	1	beene
1	a	1	aduiterate	1	answers	1	beer

4812 non-zeros

$$P(w|\text{blacke as})$$

Hamlet from NLTK, V=48 | 2

http://people.cs.umass.edu/~brenocon/inlp2014/lectures/04-hamlet_ngrams.txt

$$= \lambda_2 \hat{P}(w|\text{blacke as}) + \lambda_1 \hat{P}(w|\text{as}) + \lambda_0 \hat{P}(w)$$



3 blacke as

l blacke as death

l blacke as hell

l blacke as his

3 non-zeros

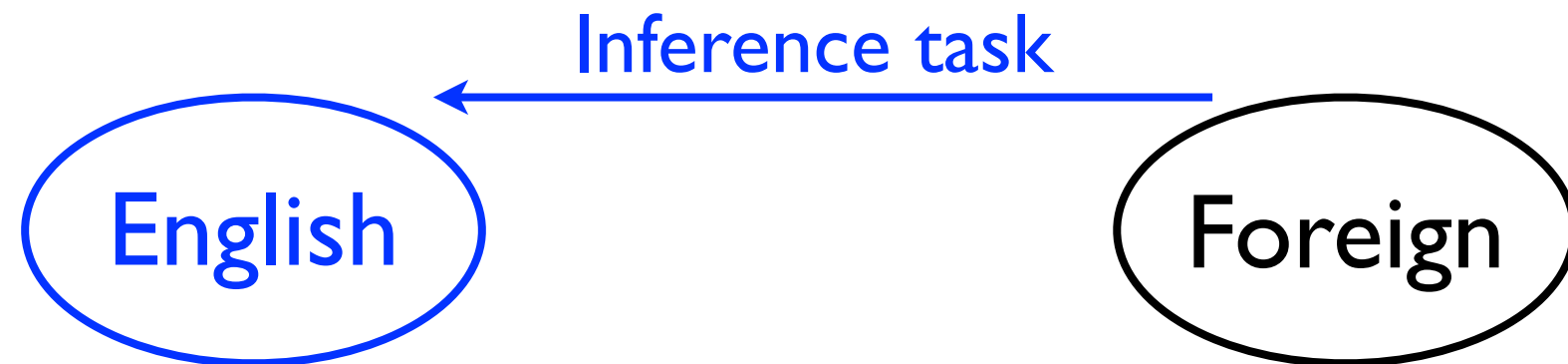
- Generative view:
 - Choose history size by prob ($\lambda_2, \lambda_1, \lambda_0$).
 - Generate w by MLE of that history size.
- Size 1 or 0 are “backoff” choices.
- Prediction time: marginalize (average) over history selection.
- Many other methods have different ways of combining different sized histories.
- Sharing statistical strength: “X as”
- Extension: common prefixes require less smoothing
 “of the _” has 59 non-zeros, so “the _” less in
- To think about: shouldn’t contexts share statistical strength beyond suffix matching?

107 non-zeros

48 | 2 non-zeros

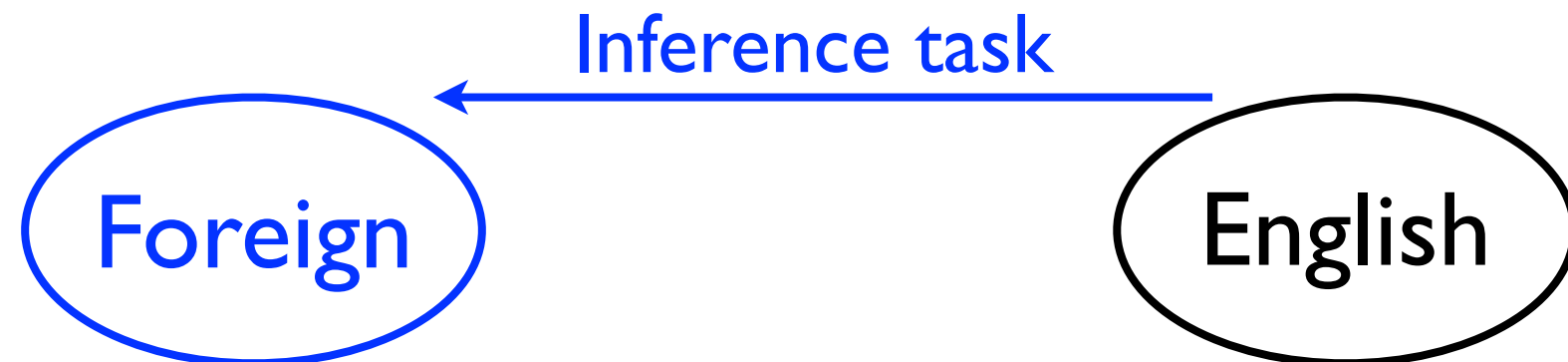
MT as Noisy Channel

Convention in Collins/Knight: translating from ***f*** into ***e***



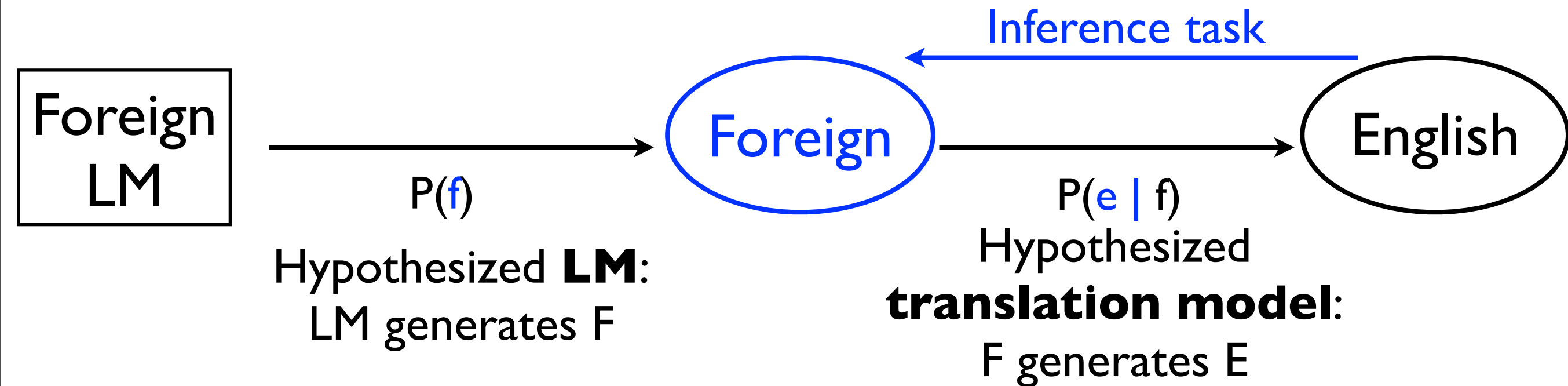
MT as Noisy Channel

Convention today (sorry!): translating from **e** into **f**



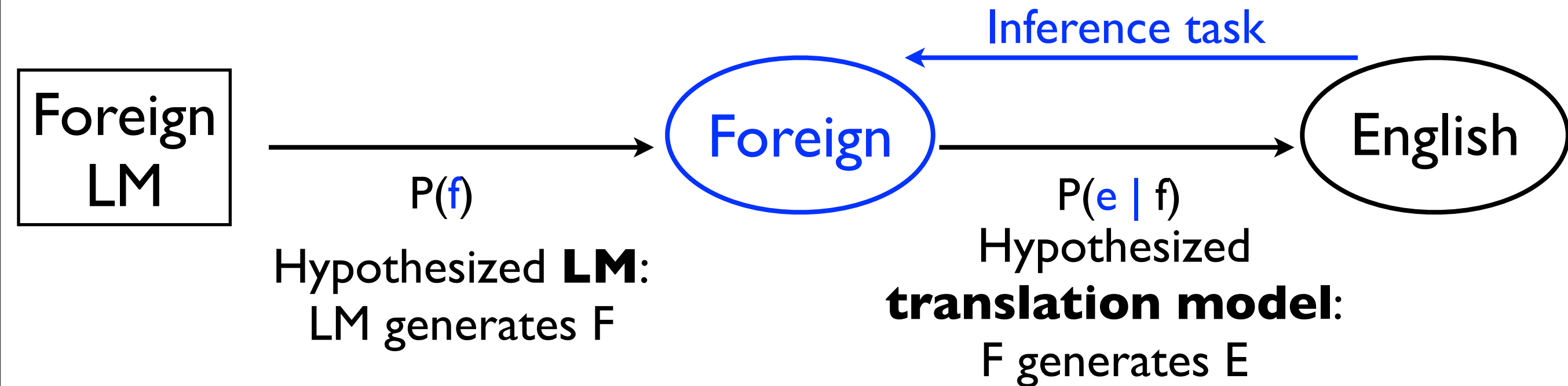
MT as Noisy Channel

Convention today (sorry!): translating from **e** into **f**



MT as Noisy Channel

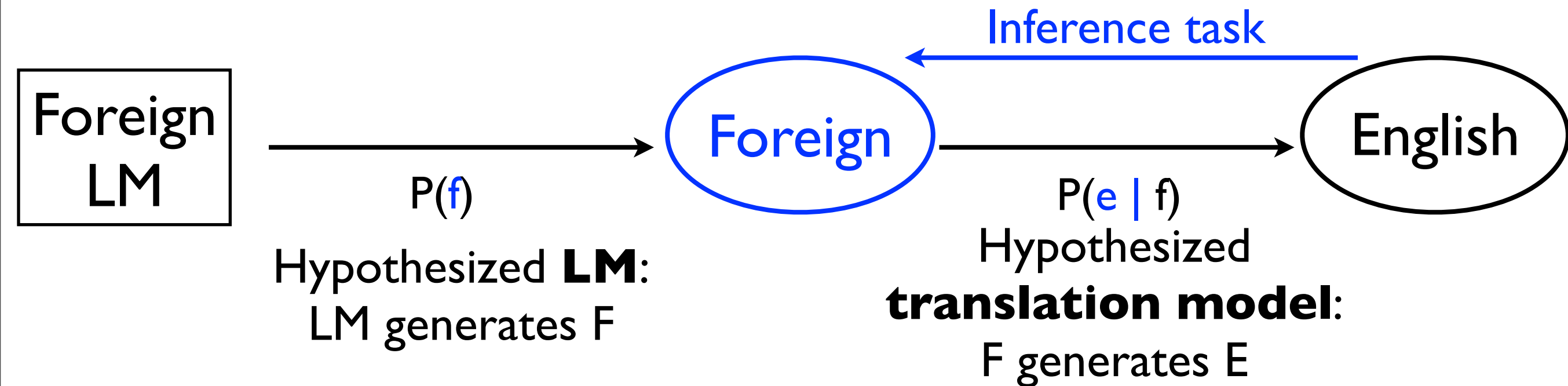
Convention today (sorry!): translating from **e** into **f**



$$P(\mathbf{f} | \mathbf{e}) = P(\mathbf{e} | \mathbf{f}) P(\mathbf{f}) / P(\mathbf{e})$$

MT as Noisy Channel

Convention today (sorry!): translating from **e** into **f**



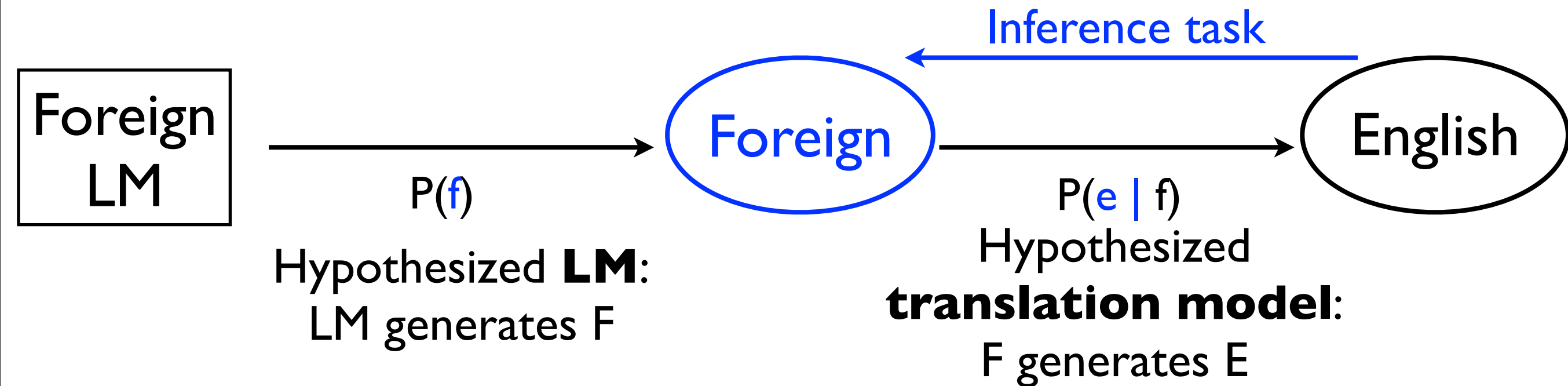
$$P(f | e) = P(e | f) P(f) / P(e)$$

Inference task ("Decoding")

$$\operatorname{argmax}_f P(f | e) = \operatorname{argmax}_f P(e | f) P(f)$$

MT as Noisy Channel

Convention today (sorry!): translating from **e** into **f**



$$P(f | e) = P(e | f) P(f) / P(e)$$

Inference task
("Decoding")

$$\operatorname{argmax}_f P(f | e) = \operatorname{argmax}_f P(e | f) P(f)$$

Learning task

$P(f)$: from large
monolingual corpus

$P(e | f)$ today:
lexical translation models

Data to learn statistical MT

- Statistical machine translation is almost entirely based on *parallel corpora*, a.k.a. *bitexts*.
- Training data: human-translated sentence pairs (**e**, **f**)

Lexical Translation

- How do we translate a word? Look it up in the dictionary

Haus : house, home, shell, household

- Multiple translations
 - Different word senses, different registers, different inflections (?)
 - *house, home* are common
 - *shell* is specialized (the Haus of a snail is a shell)

How common is each?

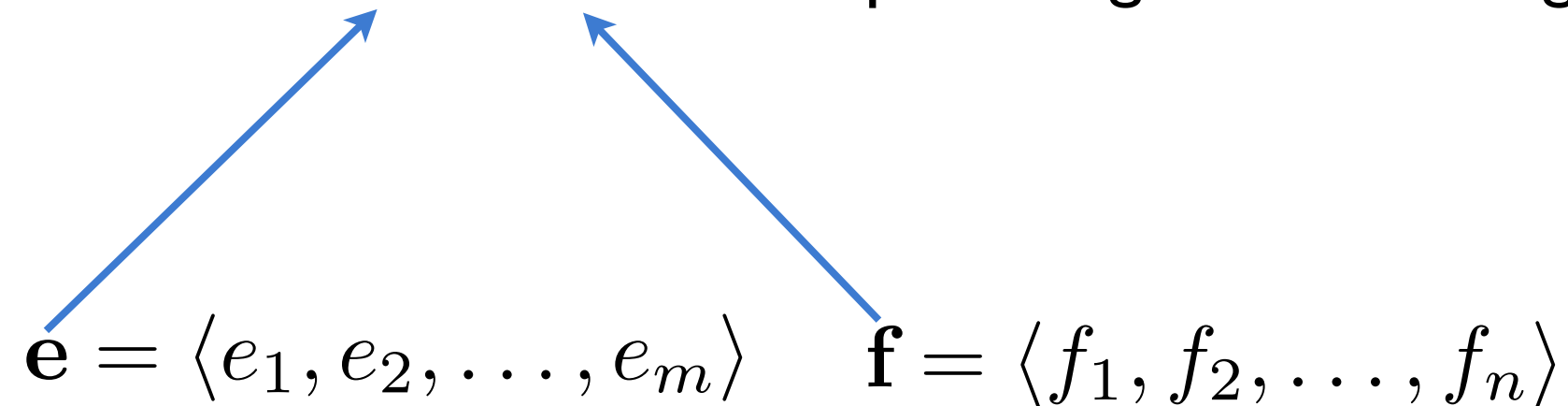
Translation	Count
house	5000
home	2000
shell	100
household	80

MLE

$$\hat{p}_{\text{MLE}}(e \mid \text{Haus}) = \begin{cases} 0.696 & \text{if } e = \text{house} \\ 0.279 & \text{if } e = \text{home} \\ 0.014 & \text{if } e = \text{shell} \\ 0.011 & \text{if } e = \text{household} \\ 0 & \text{otherwise} \end{cases}$$

Lexical Translation

- Goal: a model $p(\mathbf{e} \mid \mathbf{f}, m)$
- where $\mathbf{e}$ and $\mathbf{f}$ are complete English and Foreign sentences



The diagram illustrates the relationship between the goal model and the sentence vectors. Two blue arrows point from the vectors $\mathbf{e}$ and $\mathbf{f}$ to the model $p(\mathbf{e} \mid \mathbf{f}, m)$ in the list above. The vector $\mathbf{e}$ is defined as $\langle e_1, e_2, \dots, e_m \rangle$ and the vector $\mathbf{f}$ is defined as $\langle f_1, f_2, \dots, f_n \rangle$.

$$\mathbf{e} = \langle e_1, e_2, \dots, e_m \rangle \quad \mathbf{f} = \langle f_1, f_2, \dots, f_n \rangle$$

Lexical Translation

- Goal: a model $p(\mathbf{e} \mid \mathbf{f}, m)$
- where $\mathbf{e}$ and $\mathbf{f}$ are complete English and Foreign sentences
- Lexical translation makes the following **assumptions**:
 - Each word in e_i in $\mathbf{e}$ is generated from exactly one word in $\mathbf{f}$
 - Thus, we have an *alignment* a_i that indicates which word e_i “came from”, specifically it came from f_{a_i} .
 - Given the alignments $\mathbf{a}$, translation decisions are conditionally independent of each other and depend *only* on the aligned source word f_{a_i} .

Lexical Translation

$$\mathbf{e} = \langle e_1, e_2, \dots, e_m \rangle \quad \mathbf{f} = \langle f_1, f_2, \dots, f_n \rangle$$
$$\mathbf{a} = \langle a_1, a_2, \dots, a_m \rangle \quad \text{each } a_i \in \{0, 1, \dots, n\}$$

Chain rule

$$p(\mathbf{e} \mid \mathbf{f}, m) = \sum_{\mathbf{a} \in \{0, 1, \dots, n\}^m} p(\mathbf{a} \mid \mathbf{f}, m) \times p(\mathbf{e} \mid \mathbf{a}, \mathbf{f}, m)$$

Modeling assumptions

$$p(\mathbf{e} \mid \mathbf{f}, m) = \sum_{\mathbf{a} \in \{0, 1, \dots, n\}^m} p(\mathbf{a} \mid \mathbf{f}, m) \times \prod_{i=1}^m p(e_i \mid f_{a_i})$$

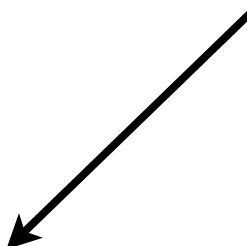
[Alignment] × [Translation | Alignment]

Lexical Translation

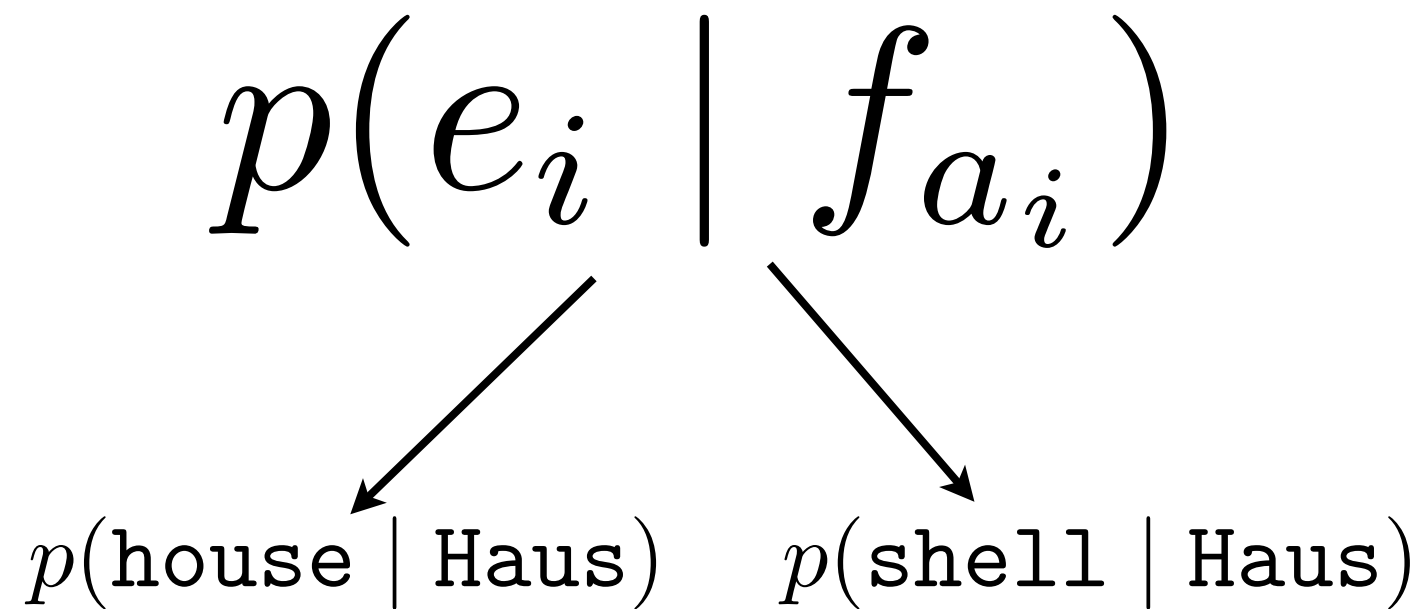
$$p(e_i \mid f_{a_i})$$

Lexical Translation

$$p(e_i | f a_i)$$


$$p(\text{house} | \text{Haus})$$

Lexical Translation



Lexical Translation

$$\mathbf{e} = \langle e_1, e_2, \dots, e_m \rangle \quad \mathbf{f} = \langle f_1, f_2, \dots, f_n \rangle$$
$$\mathbf{a} = \langle a_1, a_2, \dots, a_m \rangle \quad \text{each } a_i \in \{0, 1, \dots, n\}$$

Modeling assumptions

$$p(\mathbf{e} \mid \mathbf{f}, m) = \sum_{\mathbf{a} \in \{0, 1, \dots, n\}^m} p(\mathbf{a} \mid \mathbf{f}, m) \times \prod_{i=1}^m p(e_i \mid f_{a_i})$$

[Alignment] × [Translation | Alignment]

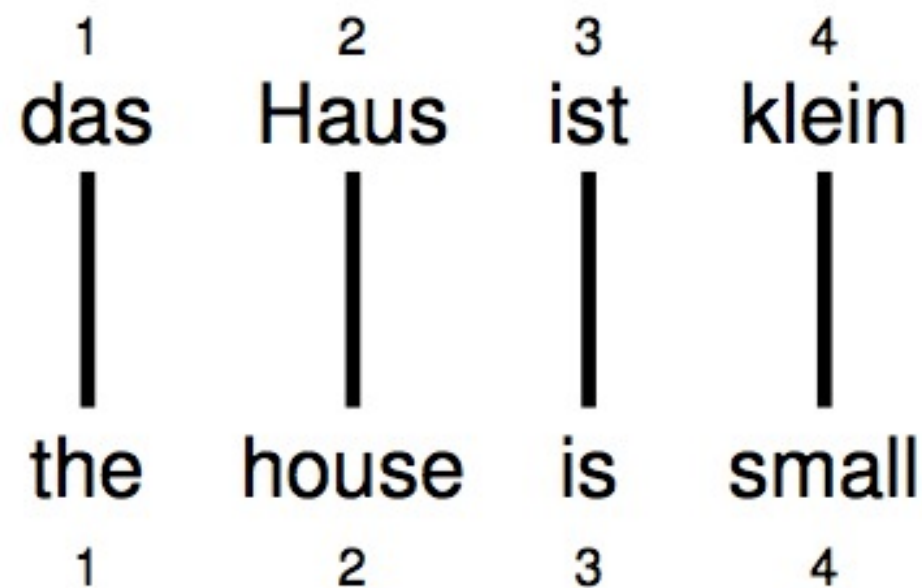
Alignment

$$p(\mathbf{a} \mid \mathbf{f}, m)$$

Most of the action for the first 10 years of MT was here. Words weren't the problem, word *order* was hard.

Alignment

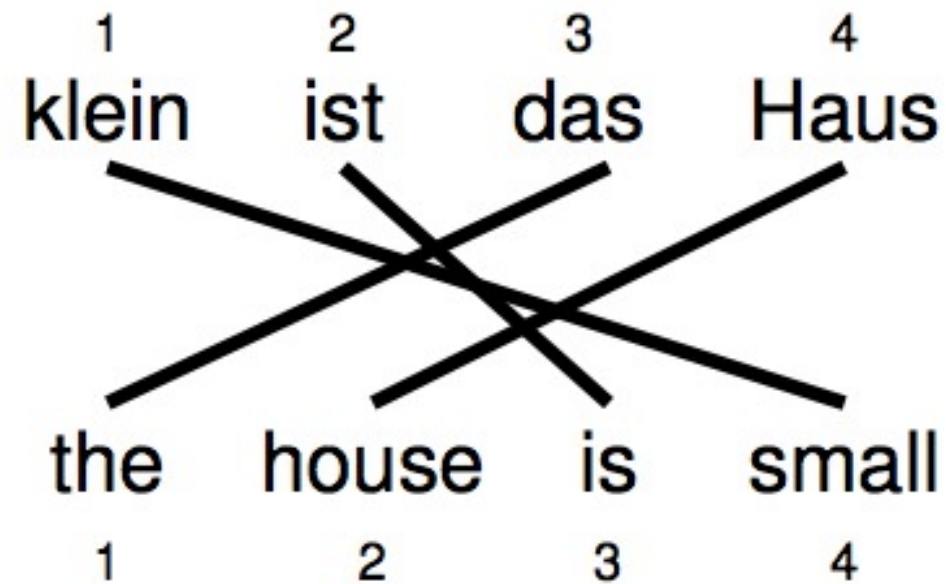
- Alignments can be visualized in by drawing links between two sentences, and they are represented as vectors of positions:



$$\mathbf{a} = (1, 2, 3, 4)^{\top}$$

Reordering

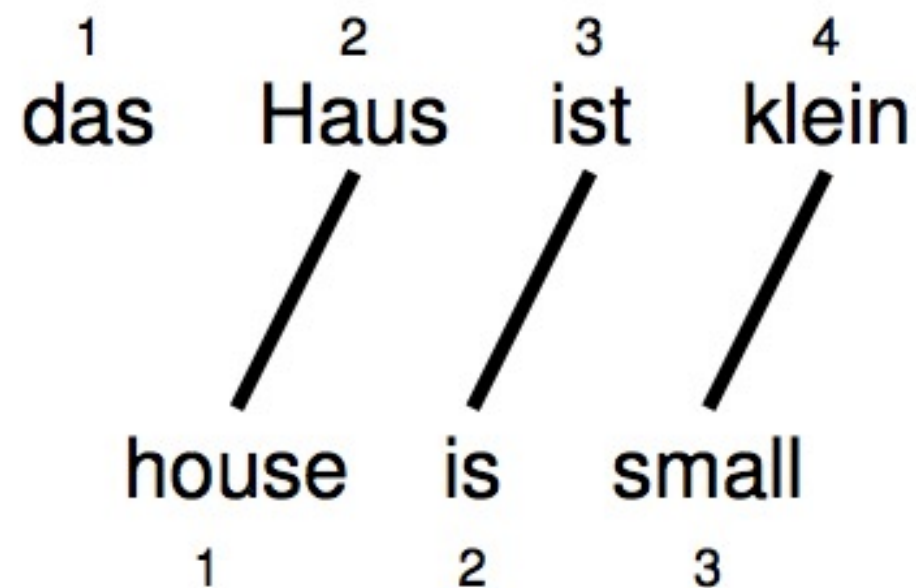
- Words may be reordered during translation.



$$\mathbf{a} = (3, 4, 2, 1)^{\top}$$

Word Dropping

- A source word may not be translated at all



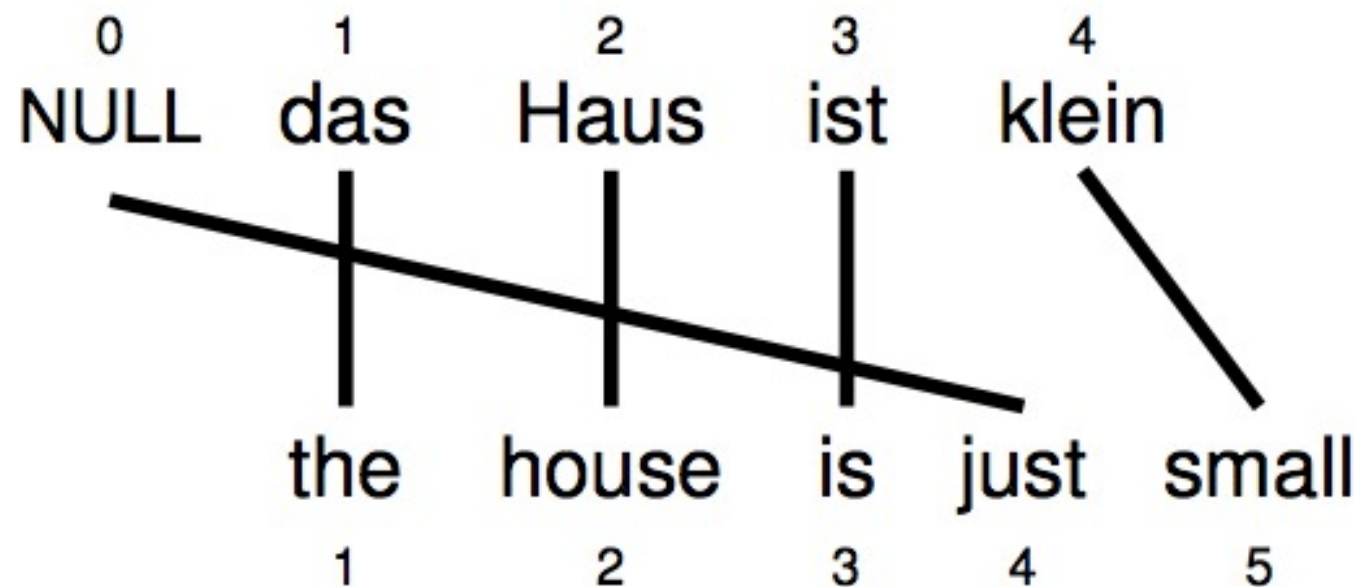
$$\mathbf{a} = (2, 3, 4)^{\top}$$

Word Insertion

- Words may be inserted during translation

English *just* does not have an equivalent

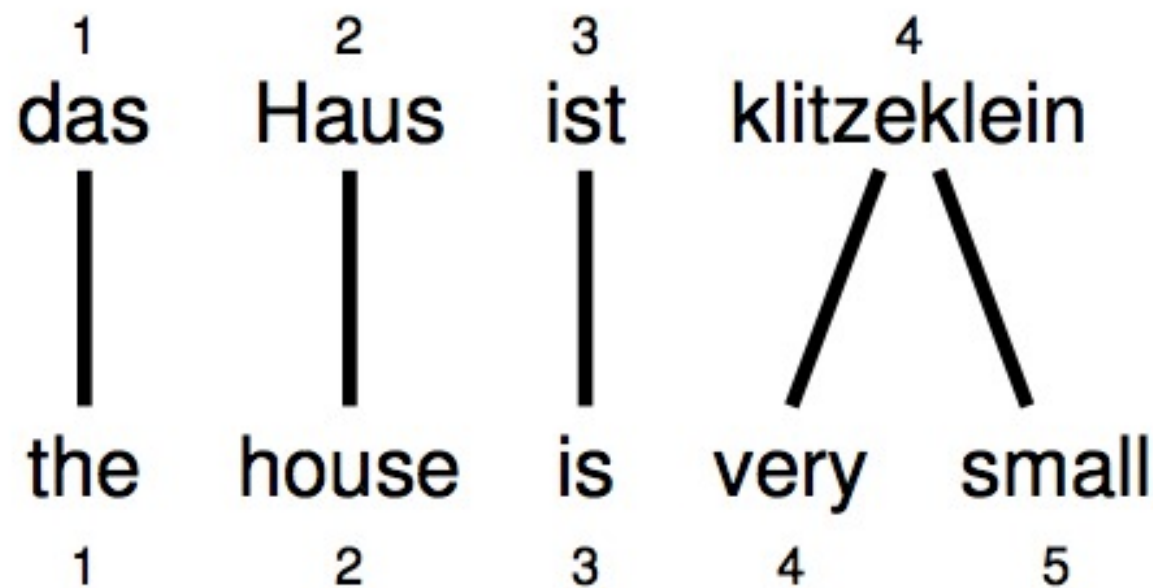
But it must be explained - we typically assume every source sentence contains a NULL token



$$\mathbf{a} = (1, 2, 3, 0, 4)^{\top}$$

One-to-many Translation

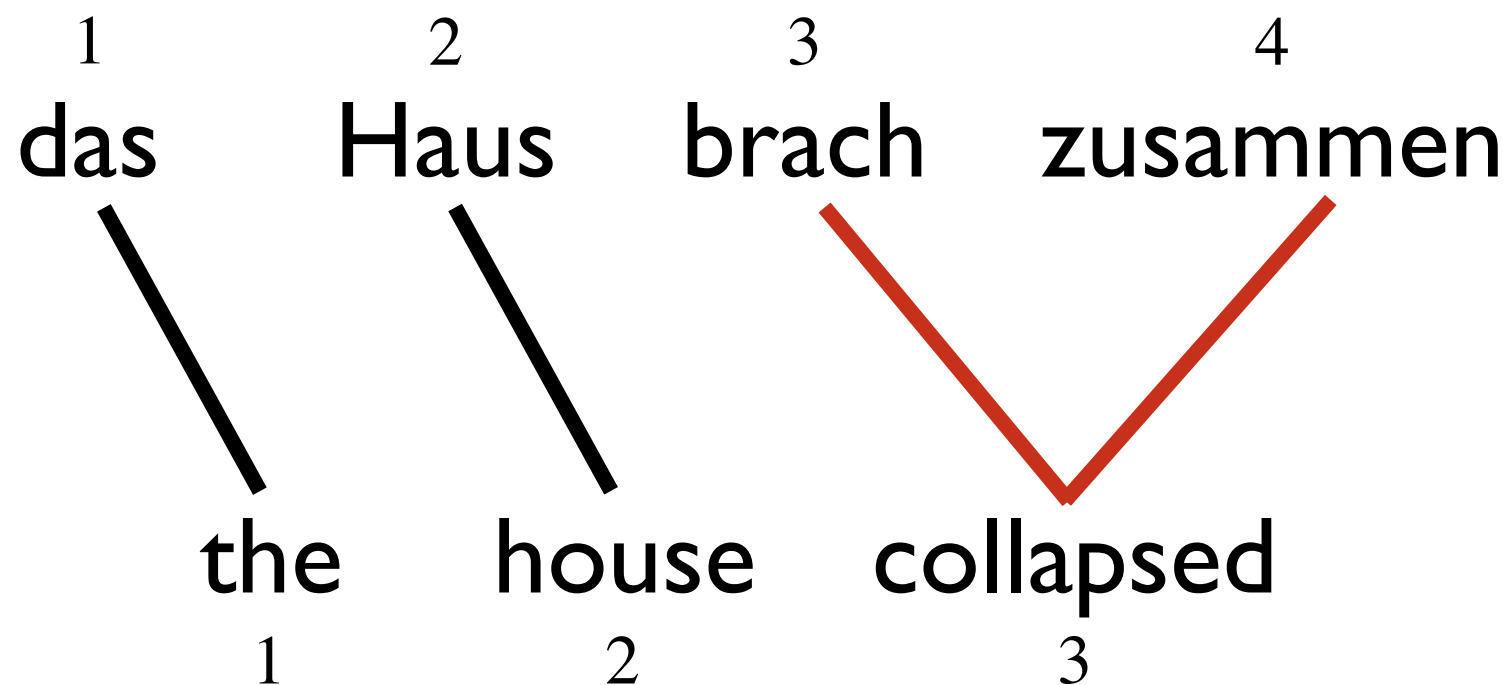
- A source word may translate into **more than one** target word



$$\mathbf{a} = (1, 2, 3, 4, 4)^{\top}$$

Many-to-one Translation

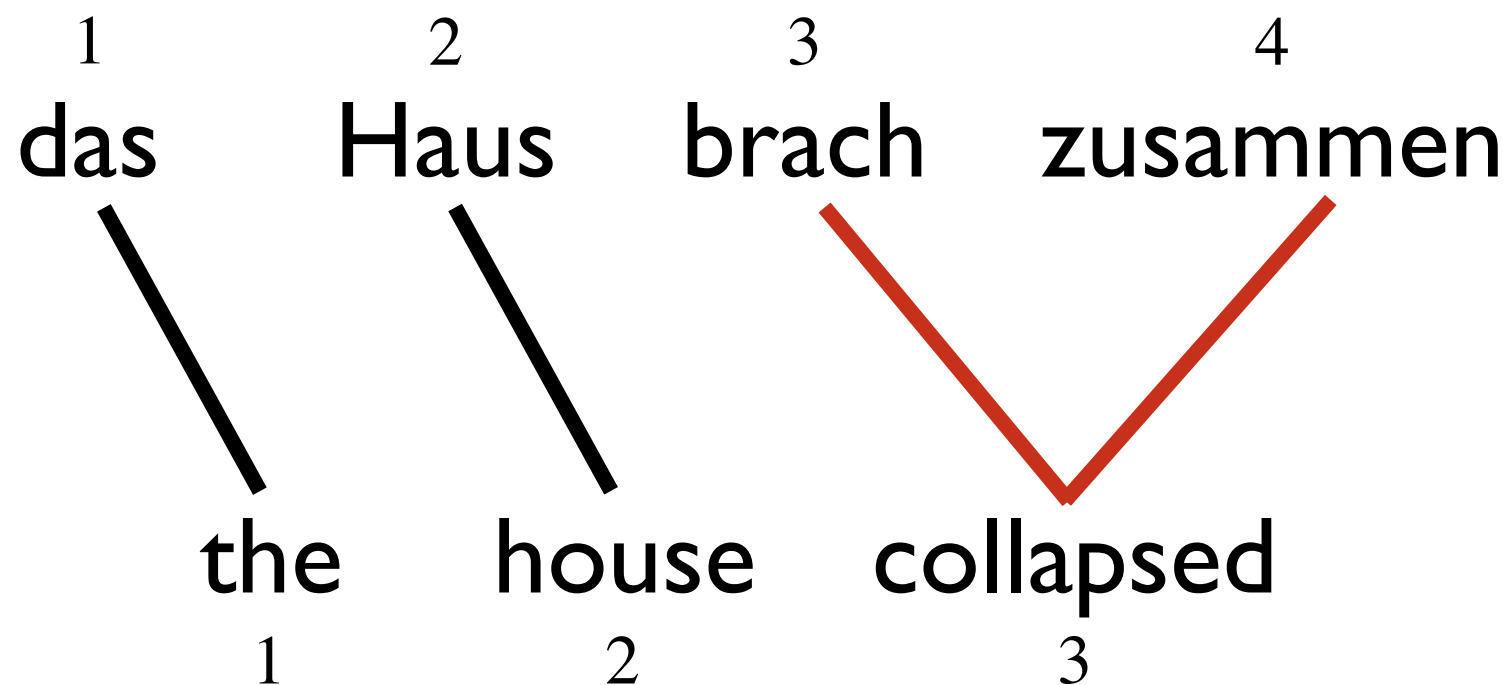
- **More than one source word** may **not** translate as a unit in lexical translation



a = ???

Many-to-one Translation

- **More than one source word** may **not** translate as a unit in lexical translation



$$\mathbf{a} = ???$$

$$\mathbf{a} = (1, 2, (3, 4)^\top)^\top \quad ?$$

IBM Model I

- Simplest possible lexical translation model
- Additional assumptions
 - The m alignment decisions are independent
 - The alignment distribution for each a_i is uniform over all source words and NULL

for each $i \in [1, 2, \dots, m]$

$$a_i \sim \text{Uniform}(0, 1, 2, \dots, n)$$

$$e_i \sim \text{Categorical}(\boldsymbol{\theta}_{f_{a_i}})$$

IBM Model I

for each $i \in [1, 2, \dots, m]$

$$a_i \sim \text{Uniform}(0, 1, 2, \dots, n)$$

$$e_i \sim \text{Categorical}(\boldsymbol{\theta}_{f_{a_i}})$$

$$p(\mathbf{e}, \mathbf{a} \mid \mathbf{f}, m) = \prod_{i=1}^m$$

IBM Model I

for each $i \in [1, 2, \dots, m]$

$$a_i \sim \text{Uniform}(0, 1, 2, \dots, n)$$

$$e_i \sim \text{Categorical}(\boldsymbol{\theta}_{f_{a_i}})$$

$$p(\mathbf{e}, \mathbf{a} \mid \mathbf{f}, m) = \prod_{i=1}^m \frac{1}{1+n}$$

IBM Model I

for each $i \in [1, 2, \dots, m]$

$a_i \sim \text{Uniform}(0, 1, 2, \dots, n)$

$e_i \sim \text{Categorical}(\boldsymbol{\theta}_{f_{a_i}})$

$$p(\mathbf{e}, \mathbf{a} \mid \mathbf{f}, m) = \prod_{i=1}^m \frac{1}{1+n} p(e_i \mid f_{a_i})$$

IBM Model I

for each $i \in [1, 2, \dots, m]$

$a_i \sim \text{Uniform}(0, 1, 2, \dots, n)$

$e_i \sim \text{Categorical}(\boldsymbol{\theta}_{f_{a_i}})$

$$p(\mathbf{e}, \mathbf{a} \mid \mathbf{f}, m) = \prod_{i=1}^m \frac{1}{1+n} p(e_i \mid f_{a_i})$$

IBM Model I

for each $i \in [1, 2, \dots, m]$

$a_i \sim \text{Uniform}(0, 1, 2, \dots, n)$

$e_i \sim \text{Categorical}(\boldsymbol{\theta}_{f_{a_i}})$

$$p(\mathbf{e}, \mathbf{a} \mid \mathbf{f}, m) = \prod_{i=1}^m \frac{1}{1+n} p(e_i \mid f_{a_i})$$

$$p(e_i, a_i \mid \mathbf{f}, m) = \frac{1}{1+n} p(e_i \mid f_{a_i})$$

$$p(\mathbf{e}, \mathbf{a} \mid \mathbf{f}, m) = \prod_{i=1}^m p(e_i, a_i \mid \mathbf{f}, m)$$

Marginal probability

$$p(e_i, a_i \mid \mathbf{f}, m) = \frac{1}{1+n} p(e_i \mid f_{a_i})$$
$$p(e_i \mid \mathbf{f}, m) = \sum_{a_i=0}^n \frac{1}{1+n} p(e_i \mid f_{a_i})$$

Recall our independence assumption: all alignment decisions are independent of each other, and given alignments all translation decisions are independent of each other, so **all translation decisions are independent of each other.**

Marginal probability

$$p(e_i, a_i \mid \mathbf{f}, m) = \frac{1}{1+n} p(e_i \mid f_{a_i})$$

$$p(e_i \mid \mathbf{f}, m) = \sum_{a_i=0}^n \frac{1}{1+n} p(e_i \mid f_{a_i})$$

Recall our independence assumption: all alignment decisions are independent of each other, and given alignments all translation decisions are independent of each other, so **all translation decisions are independent of each other**.

$$p(a, b, c, d) = p(a)p(b)p(c)p(d)$$

Marginal probability

$$p(e_i, a_i \mid \mathbf{f}, m) = \frac{1}{1+n} p(e_i \mid f_{a_i})$$

$$p(e_i \mid \mathbf{f}, m) = \sum_{a_i=0}^n \frac{1}{1+n} p(e_i \mid f_{a_i})$$

Recall our independence assumption: all alignment decisions are independent of each other, and given alignments all translation decisions are independent of each other, so **all translation decisions are independent of each other.**

Marginal probability

$$p(e_i, a_i \mid \mathbf{f}, m) = \frac{1}{1+n} p(e_i \mid f_{a_i})$$

$$p(e_i \mid \mathbf{f}, m) = \sum_{a_i=0}^n \frac{1}{1+n} p(e_i \mid f_{a_i})$$

Recall our independence assumption: all alignment decisions are independent of each other, and given alignments all translation decisions are independent of each other, so **all translation decisions are independent of each other**.

$$p(\mathbf{e} \mid \mathbf{f}, m) = \prod_{i=1}^m p(e_i \mid \mathbf{f}, m)$$

Marginal probability

$$p(e_i, a_i \mid \mathbf{f}, m) = \frac{1}{1+n} p(e_i \mid f_{a_i})$$

$$p(e_i \mid \mathbf{f}, m) = \sum_{a_i=0}^n \frac{1}{1+n} p(e_i \mid f_{a_i})$$

$$p(\mathbf{e} \mid \mathbf{f}, m) = \prod_{i=1}^m p(e_i \mid \mathbf{f}, m)$$

Marginal probability

$$p(e_i, a_i \mid \mathbf{f}, m) = \frac{1}{1+n} p(e_i \mid f_{a_i})$$

$$p(e_i \mid \mathbf{f}, m) = \sum_{a_i=0}^n \frac{1}{1+n} p(e_i \mid f_{a_i})$$

$$p(\mathbf{e} \mid \mathbf{f}, m) = \prod_{i=1}^m p(e_i \mid \mathbf{f}, m)$$

$$= \prod_{i=1}^m \sum_{a_i=0}^n \frac{1}{1+n} p(e_i \mid f_{a_i})$$

Marginal probability

$$p(e_i, a_i \mid \mathbf{f}, m) = \frac{1}{1+n} p(e_i \mid f_{a_i})$$

$$p(e_i \mid \mathbf{f}, m) = \sum_{a_i=0}^n \frac{1}{1+n} p(e_i \mid f_{a_i})$$

$$p(\mathbf{e} \mid \mathbf{f}, m) = \prod_{i=1}^m p(e_i \mid \mathbf{f}, m)$$

$$= \prod_{i=1}^m \sum_{a_i=0}^n \frac{1}{1+n} p(e_i \mid f_{a_i})$$

$$= \frac{1}{(1+n)^m} \prod_{i=1}^m \sum_{a_i=0}^n p(e_i \mid f_{a_i})$$

Example

0	1	2	3	4
NULL	das	Haus	ist	klein

<u>1</u>	<u>2</u>	<u>3</u>	<u>4</u>
----------	----------	----------	----------

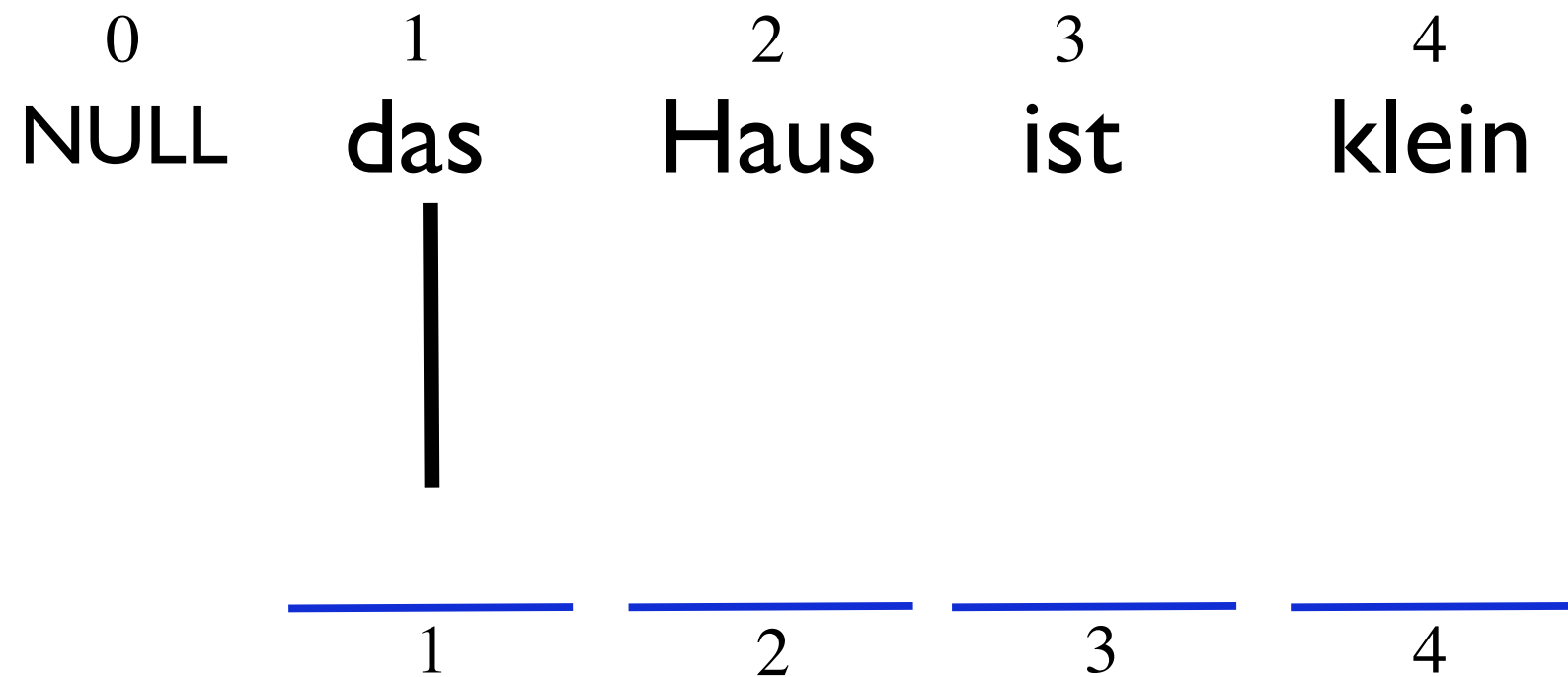
Start with a foreign sentence and a target length.

Example

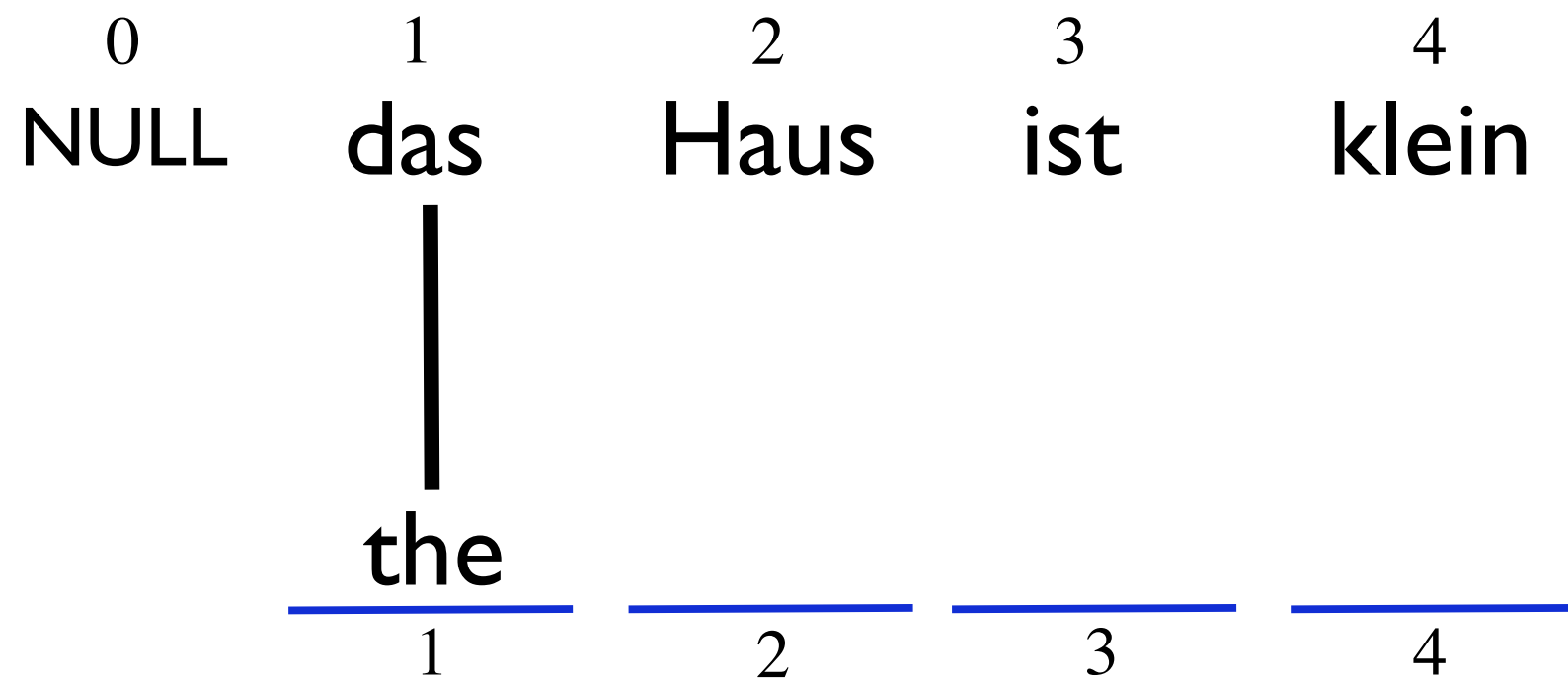
0	1	2	3	4
NULL	das	Haus	ist	klein

<u>1</u>	<u>2</u>	<u>3</u>	<u>4</u>
----------	----------	----------	----------

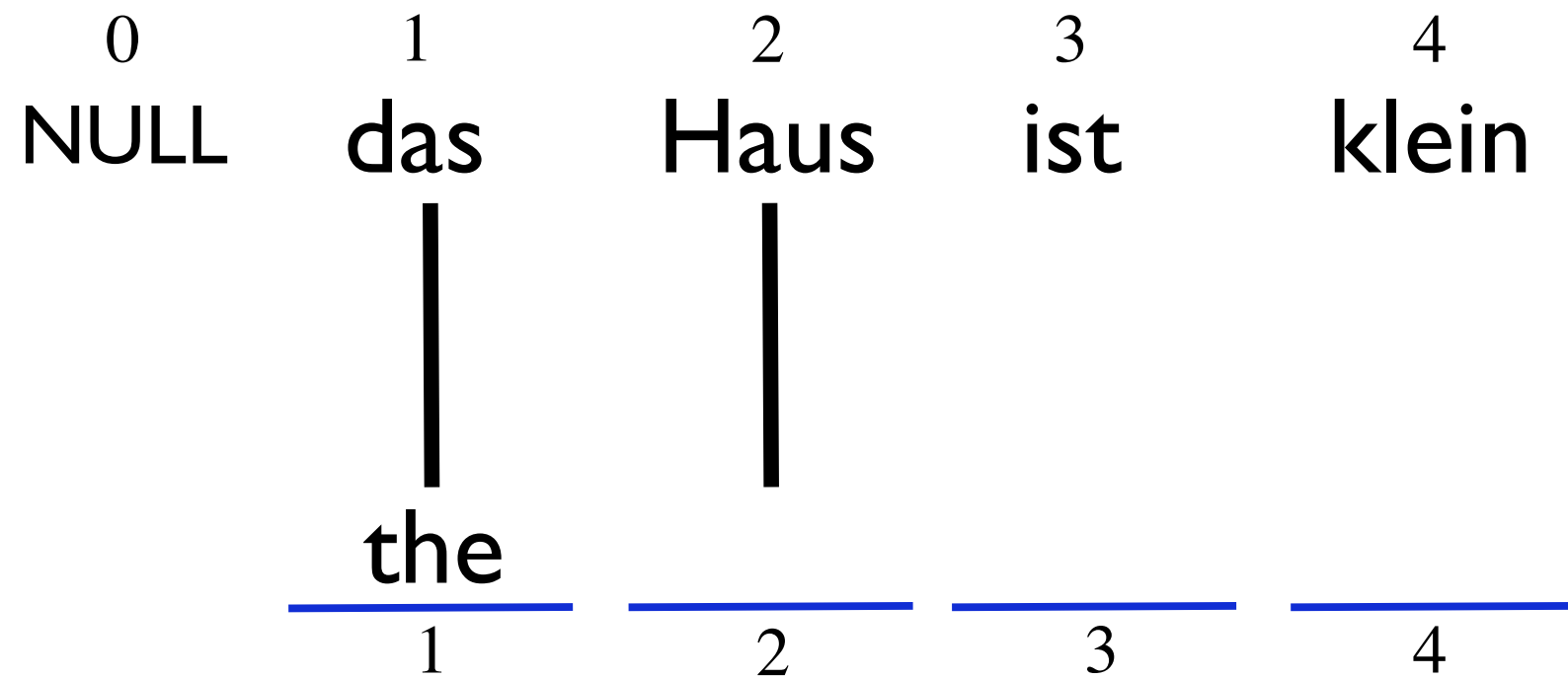
Example



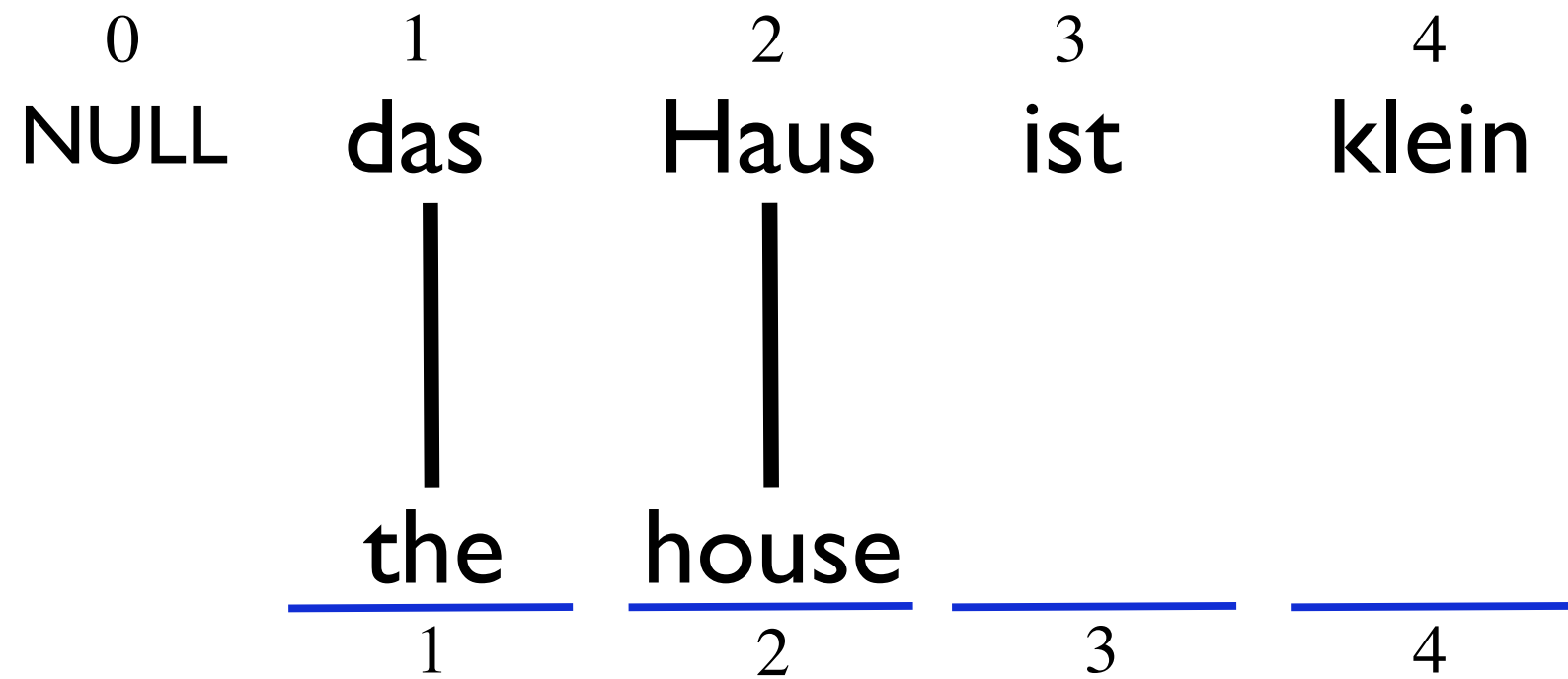
Example



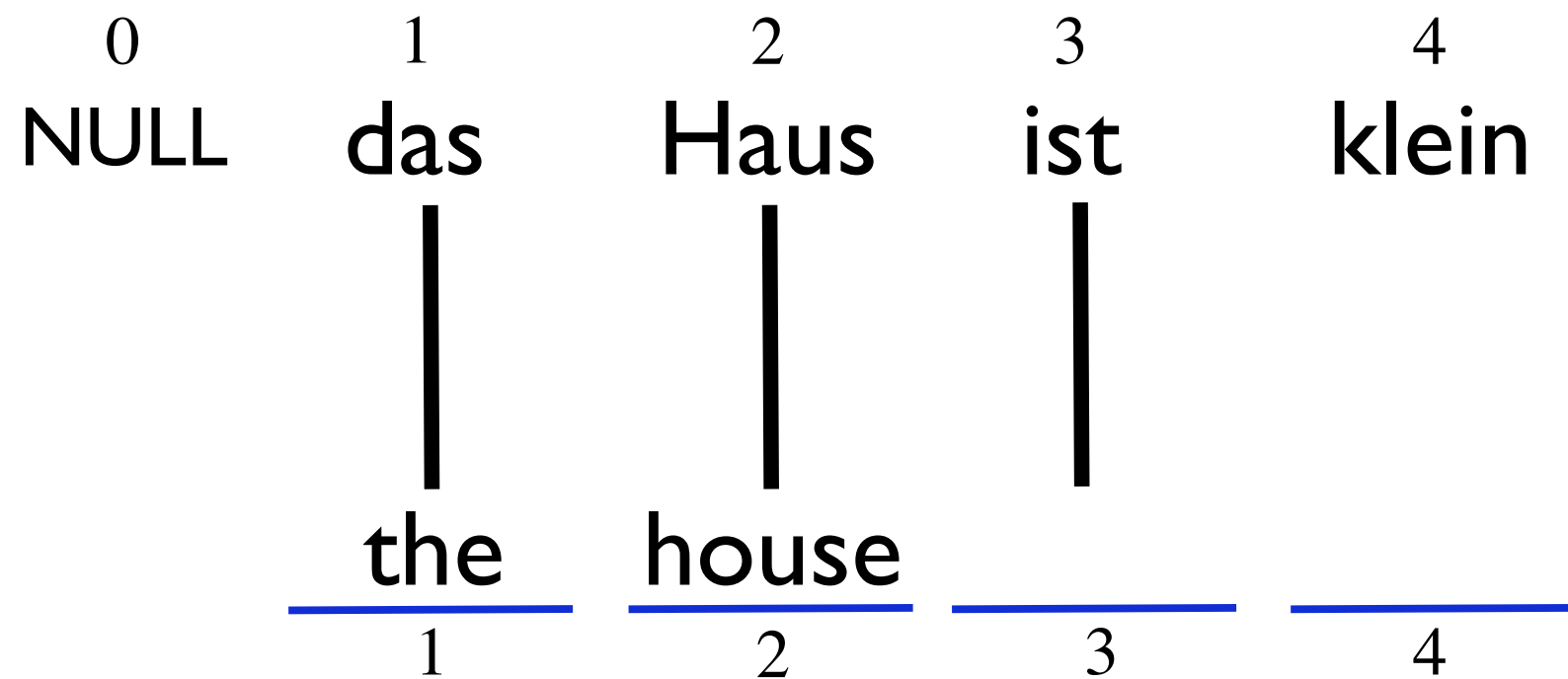
Example



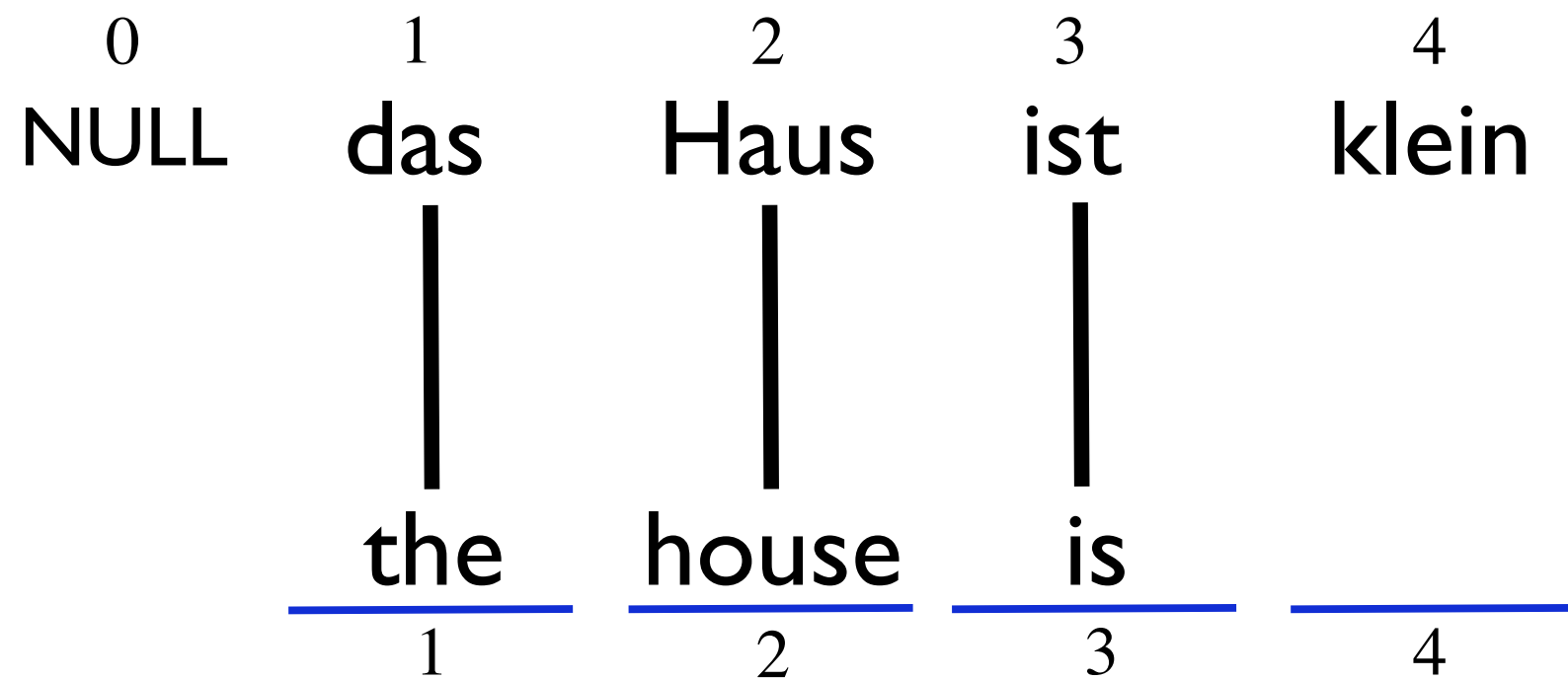
Example



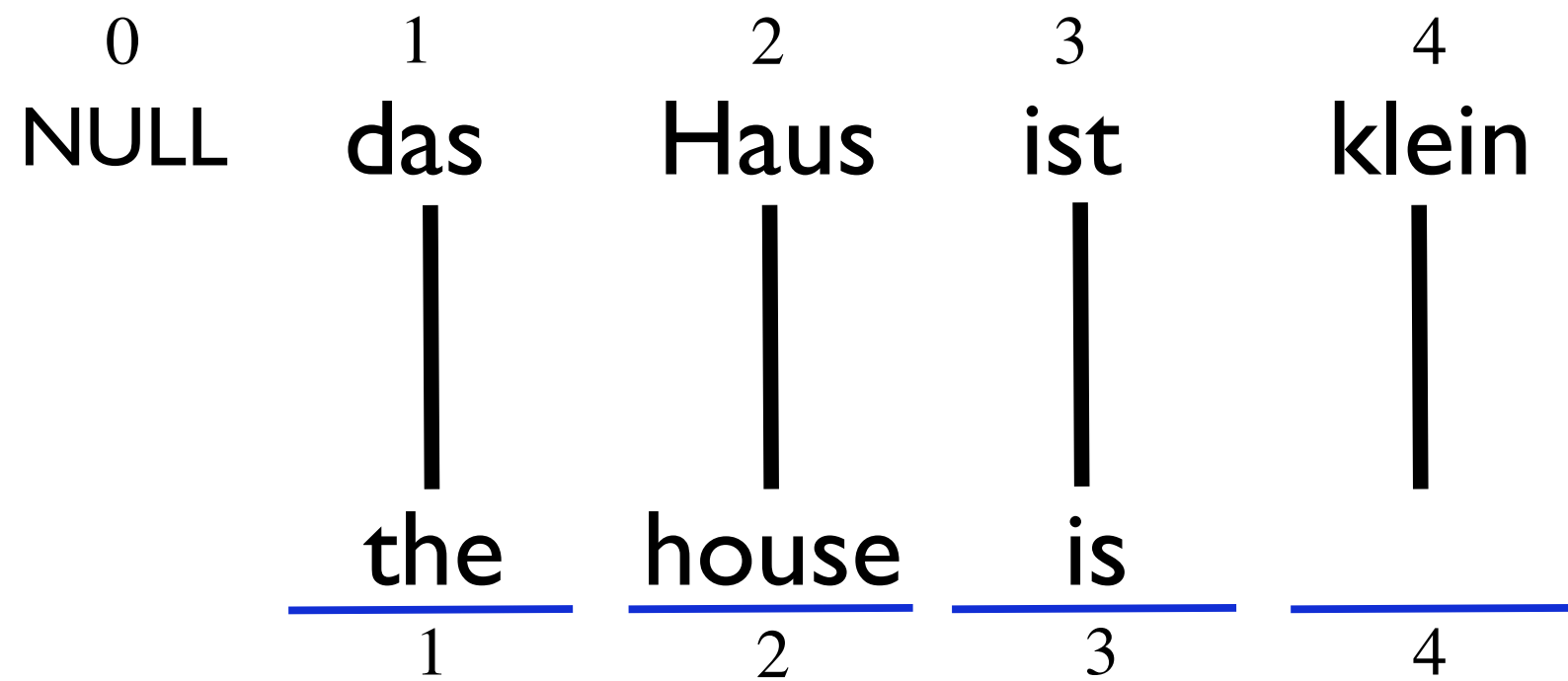
Example



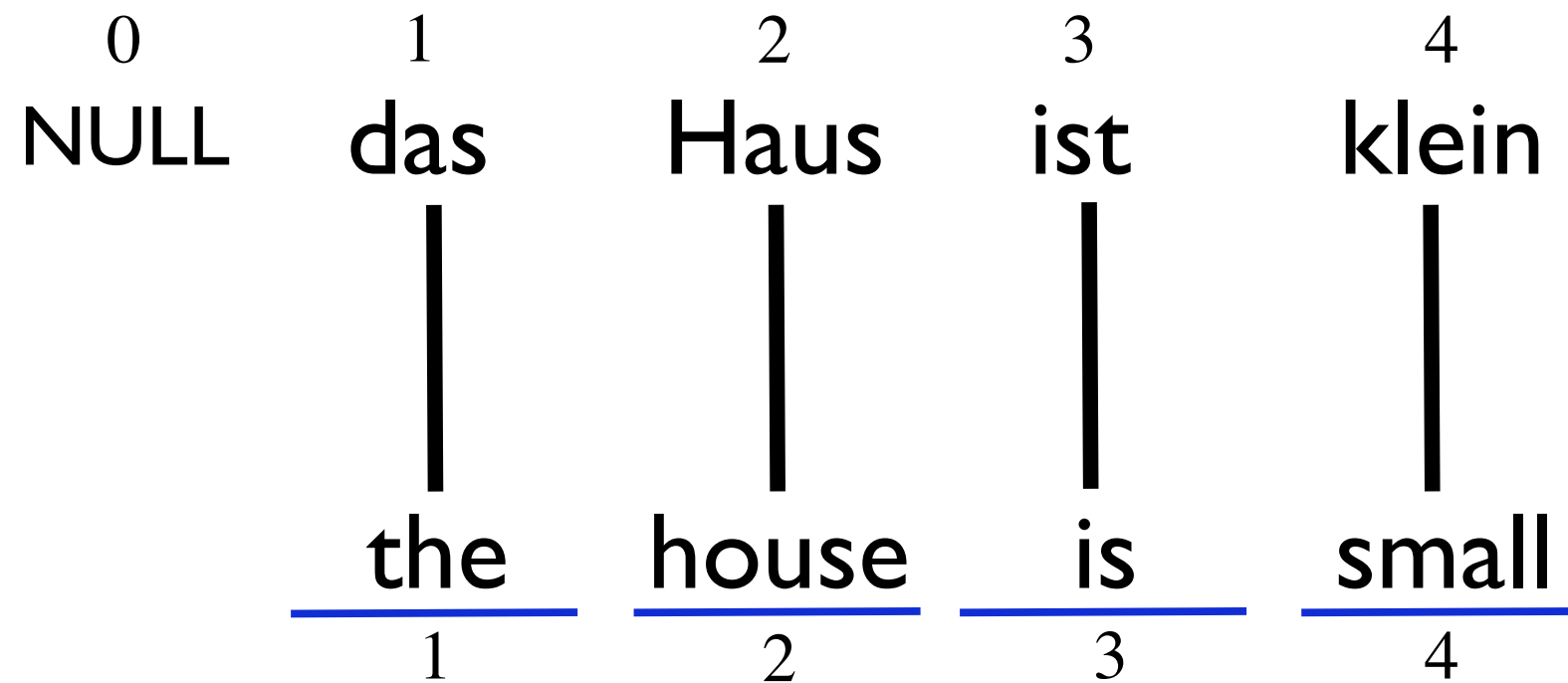
Example



Example



Example

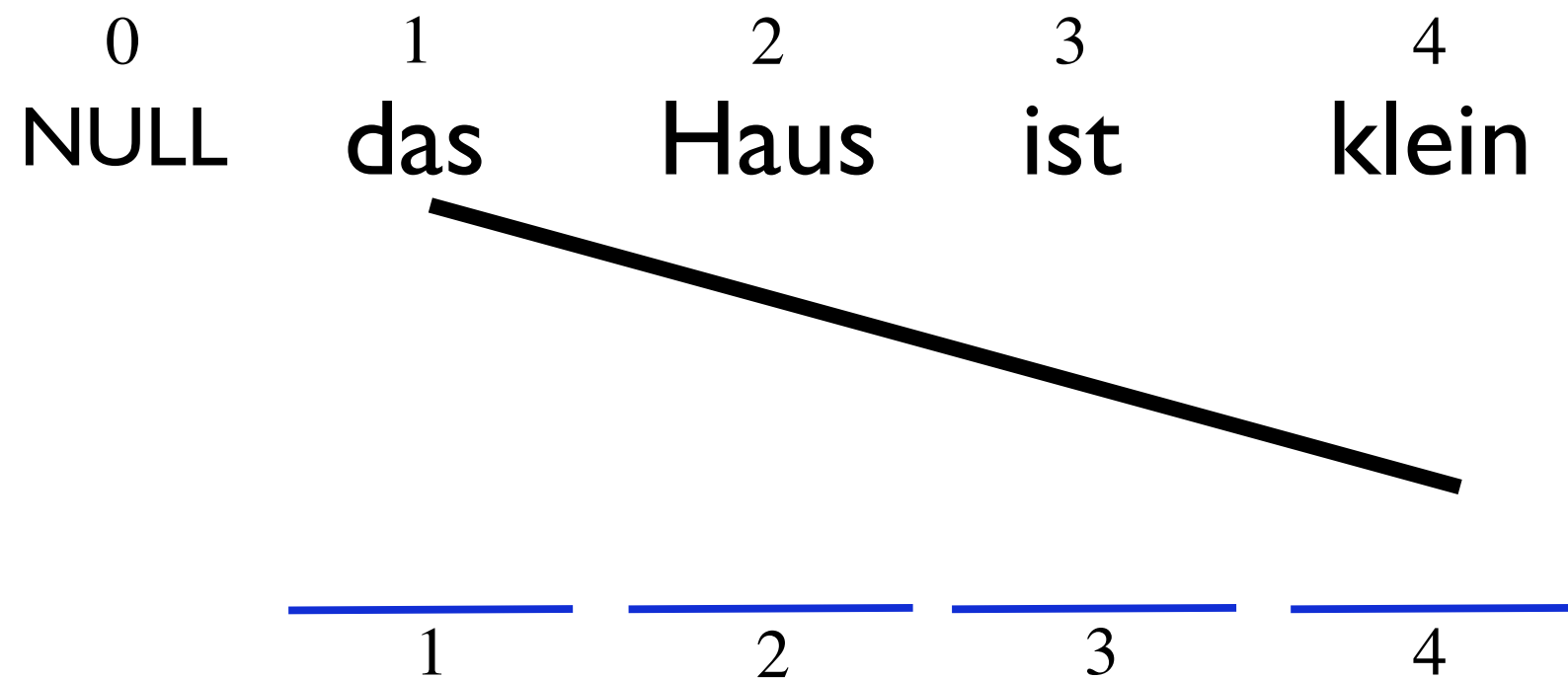


Example

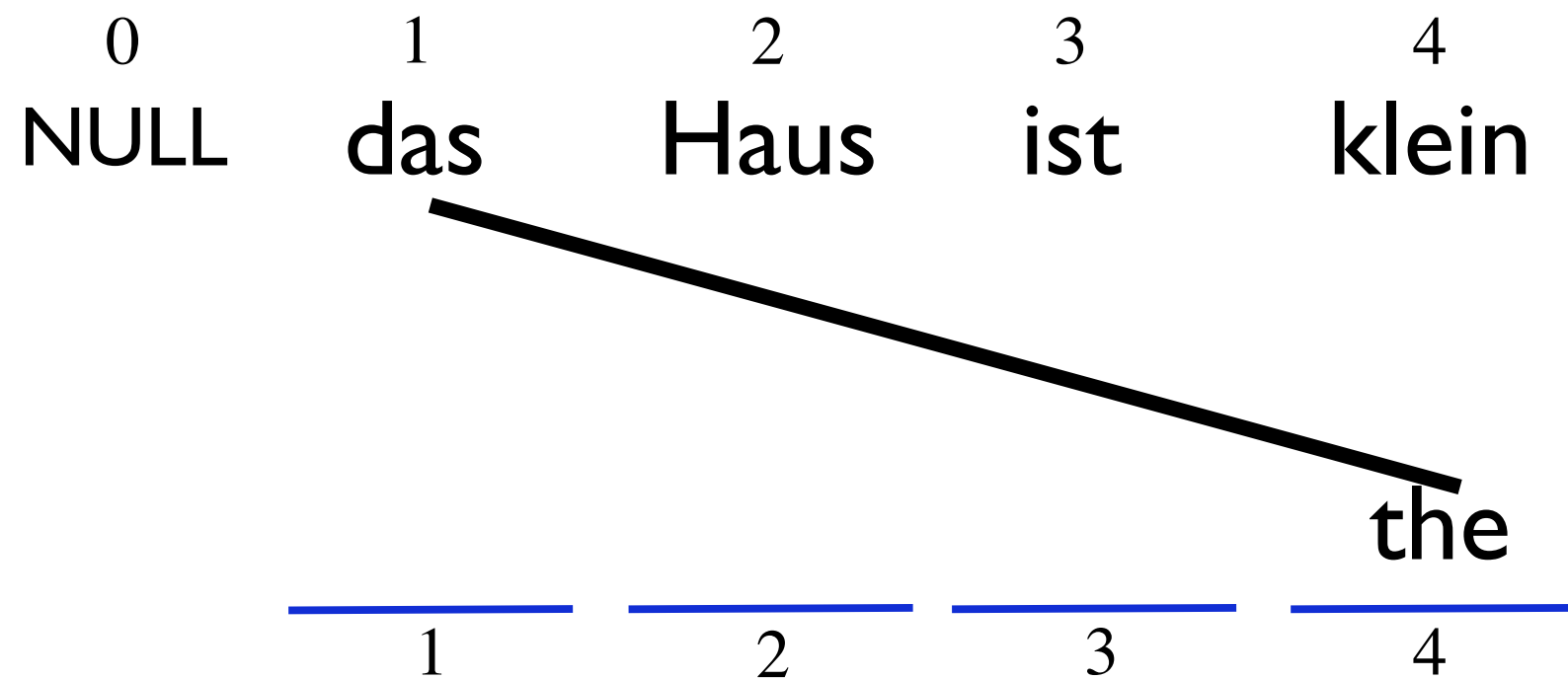
0	1	2	3	4
NULL	das	Haus	ist	klein

<u>1</u>	<u>2</u>	<u>3</u>	<u>4</u>
----------	----------	----------	----------

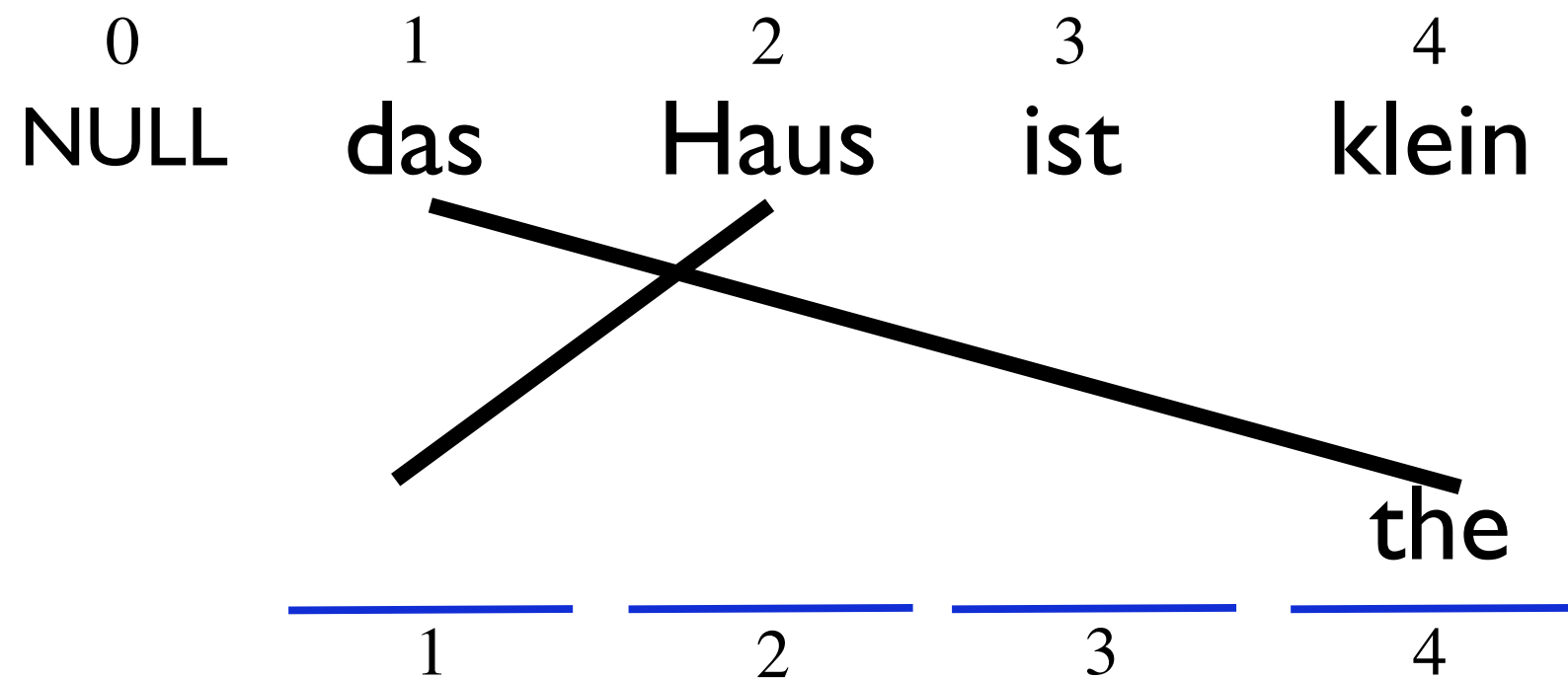
Example



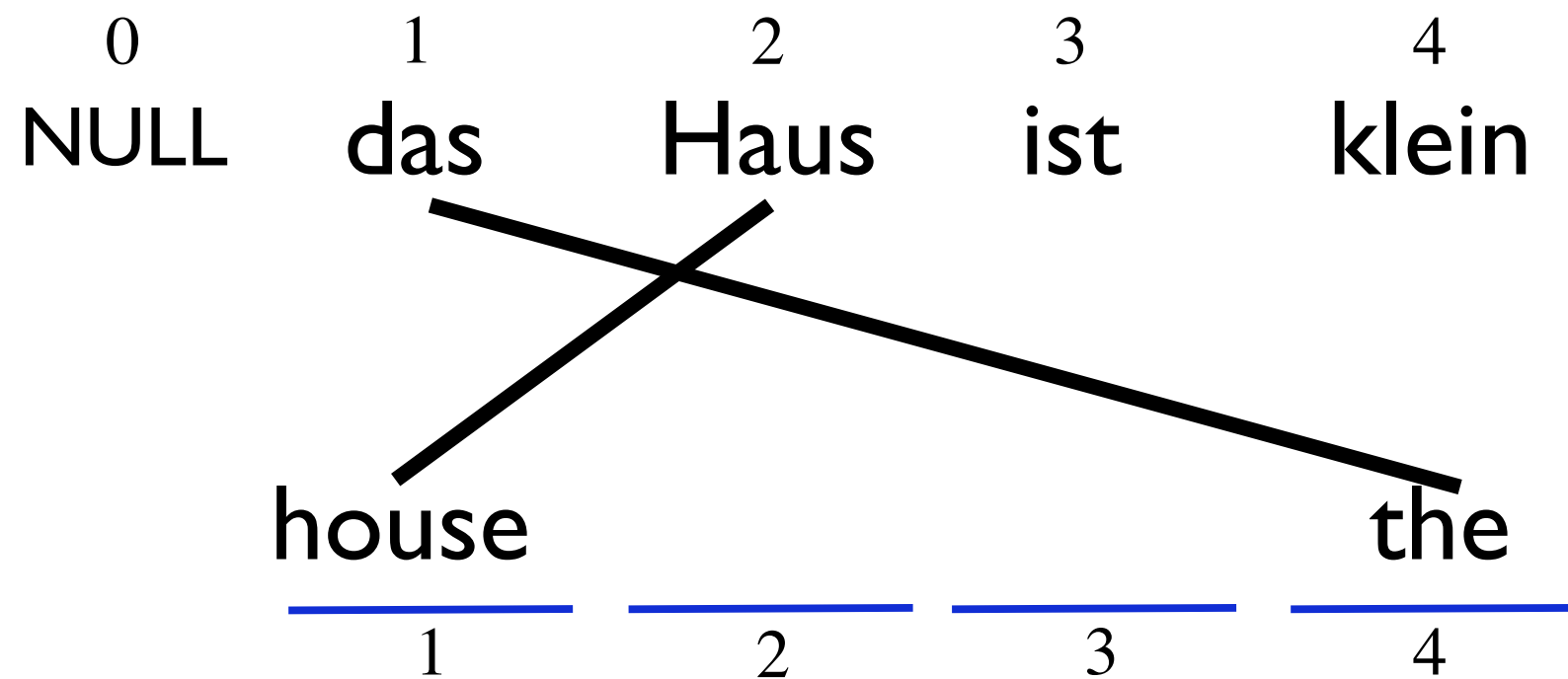
Example



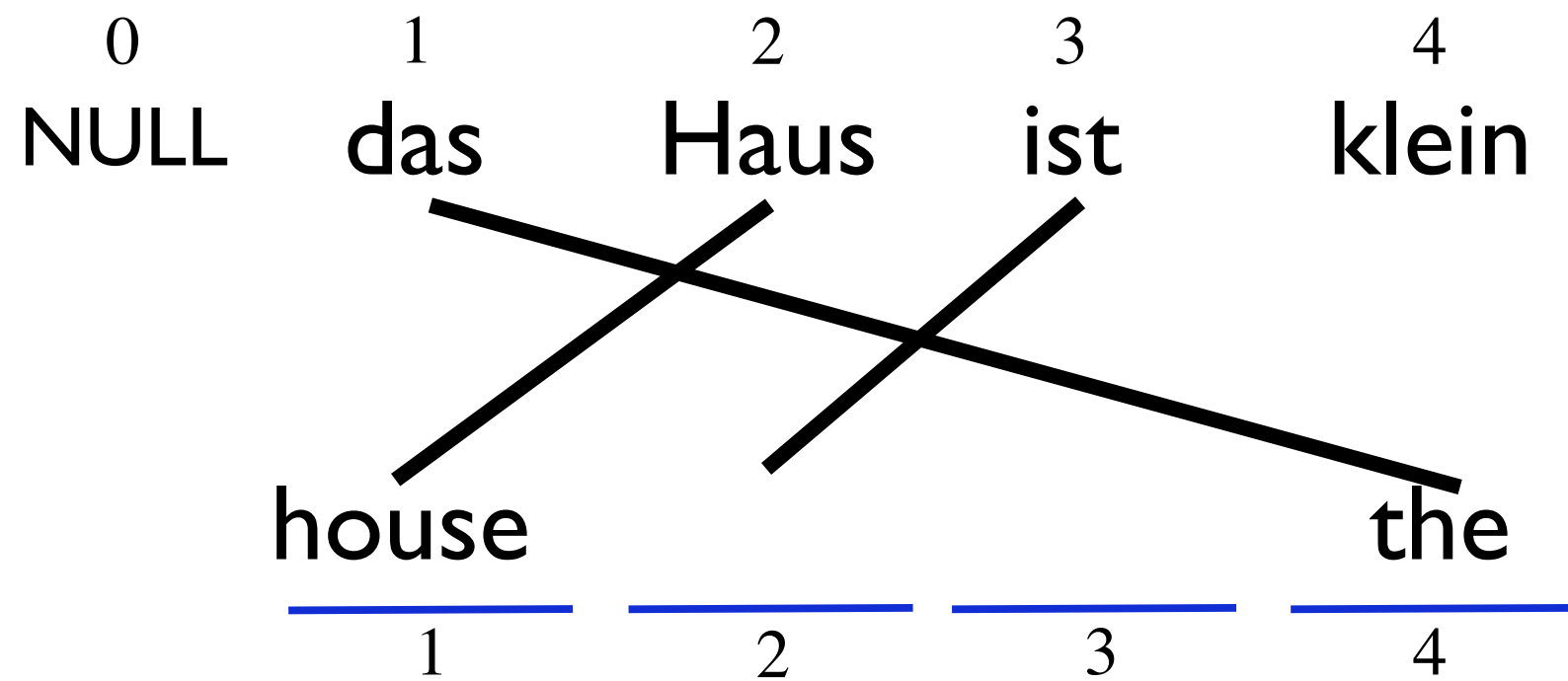
Example



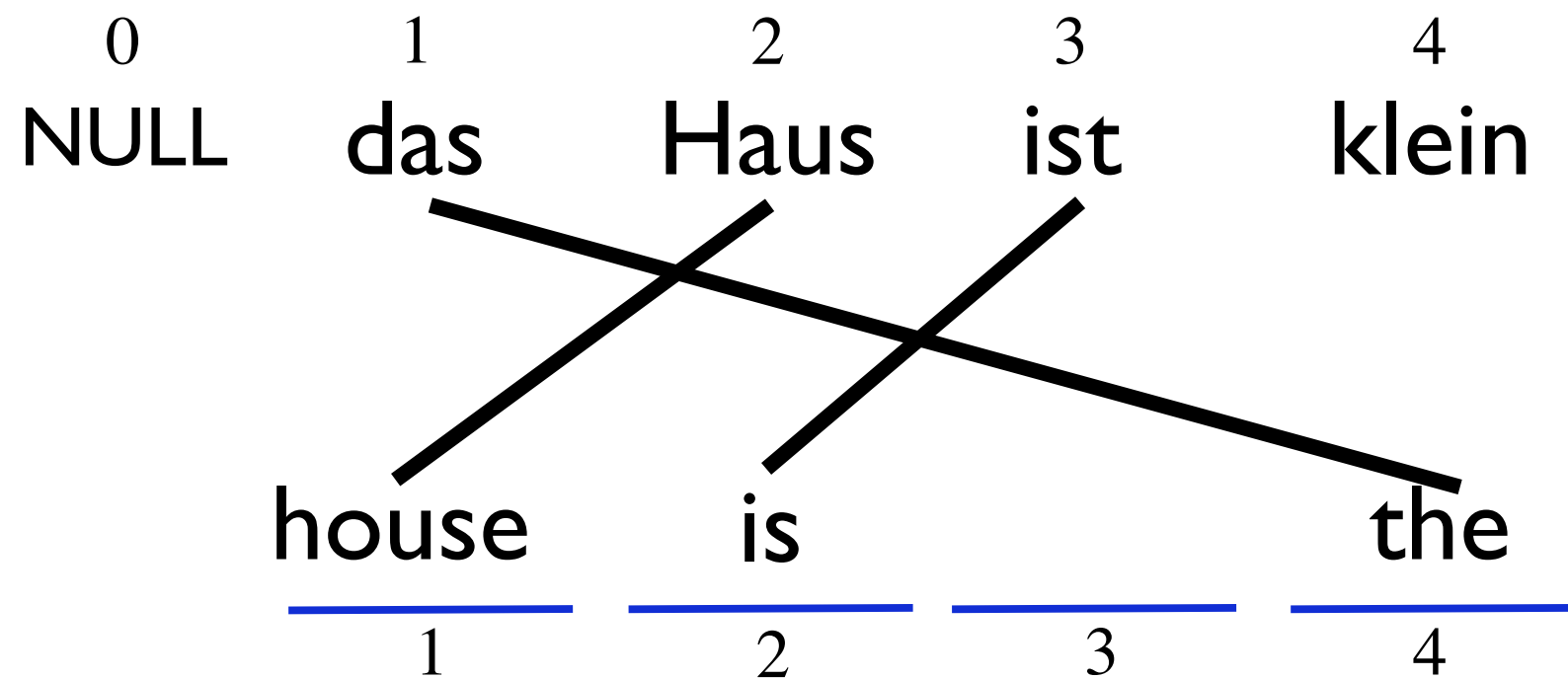
Example



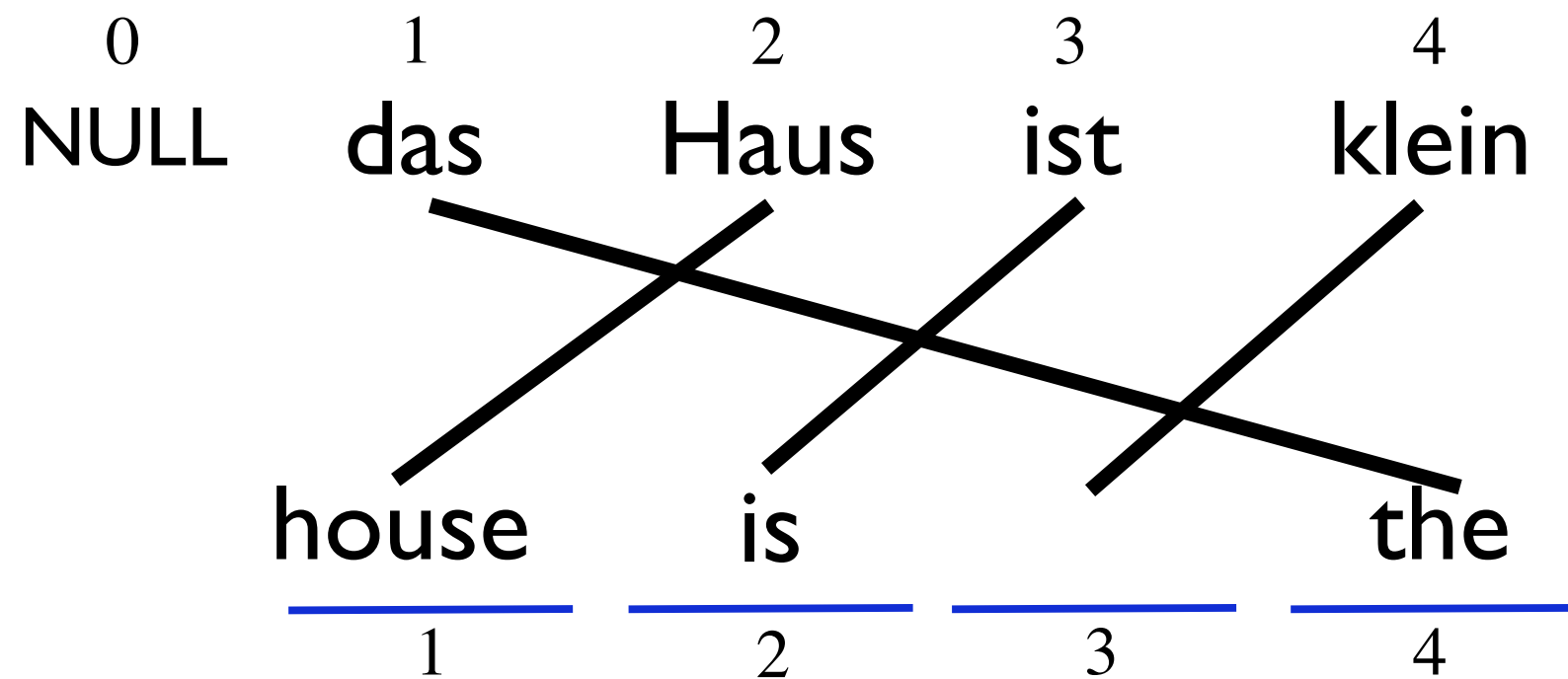
Example



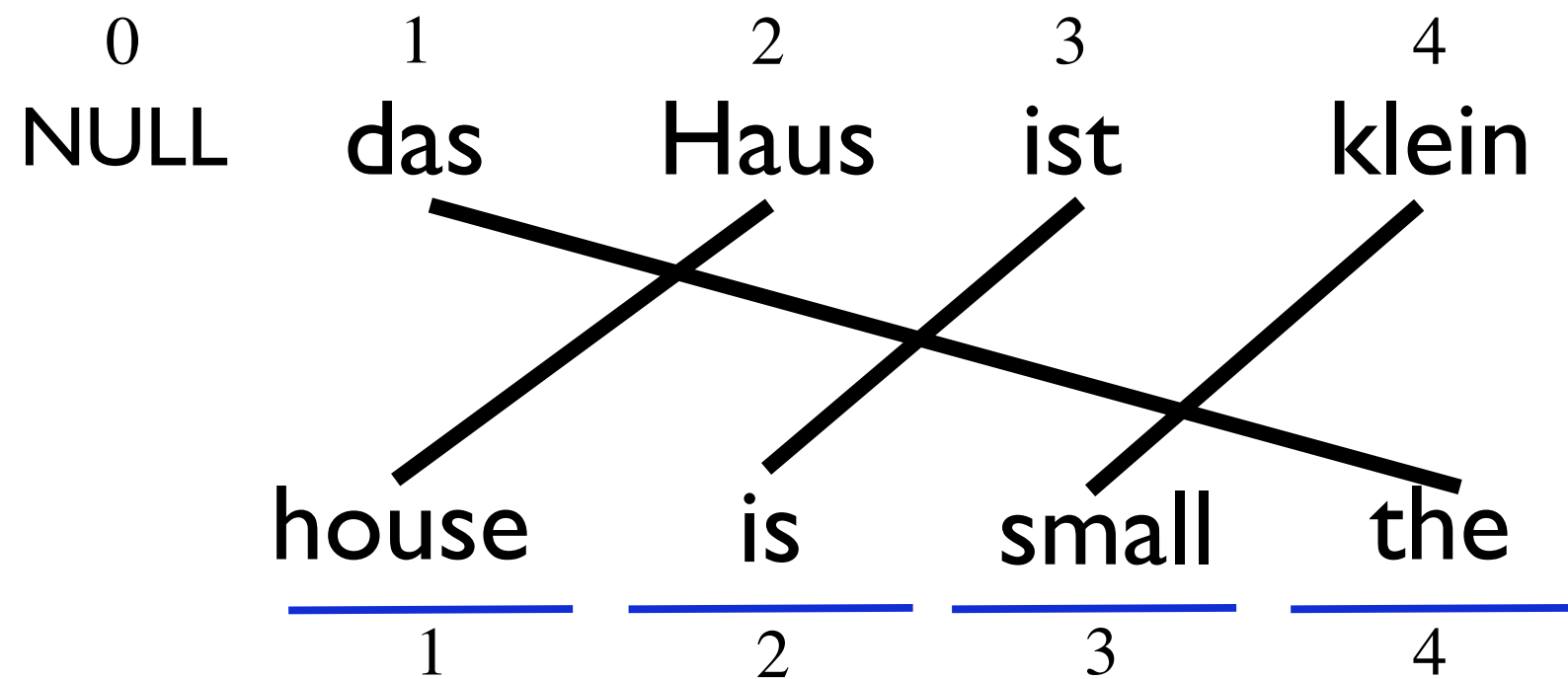
Example



Example



Example



What is this good for?

- 1. Evaluate translation quality
- 2. Find the best alignment
 - The IBM models are still used for this, though not as translation models

Finding the Viterbi Alignment

$$\begin{aligned} \mathbf{a}^* &= \arg \max_{\mathbf{a} \in [0,1,\dots,n]^m} p(\mathbf{a} \mid \mathbf{e}, \mathbf{f}) \\ &= \arg \max_{\mathbf{a} \in [0,1,\dots,n]^m} \frac{p(\mathbf{e}, \mathbf{a} \mid \mathbf{f})}{\sum_{\mathbf{a}'} p(\mathbf{e}, \mathbf{a}' \mid \mathbf{f})} \\ &= \arg \max_{\mathbf{a} \in [0,1,\dots,n]^m} p(\mathbf{e}, \mathbf{a} \mid \mathbf{f}) \\ &= p(\mathbf{a} \mid \mathbf{f}) p(\mathbf{e} \mid \mathbf{a}, \mathbf{f}) \end{aligned}$$

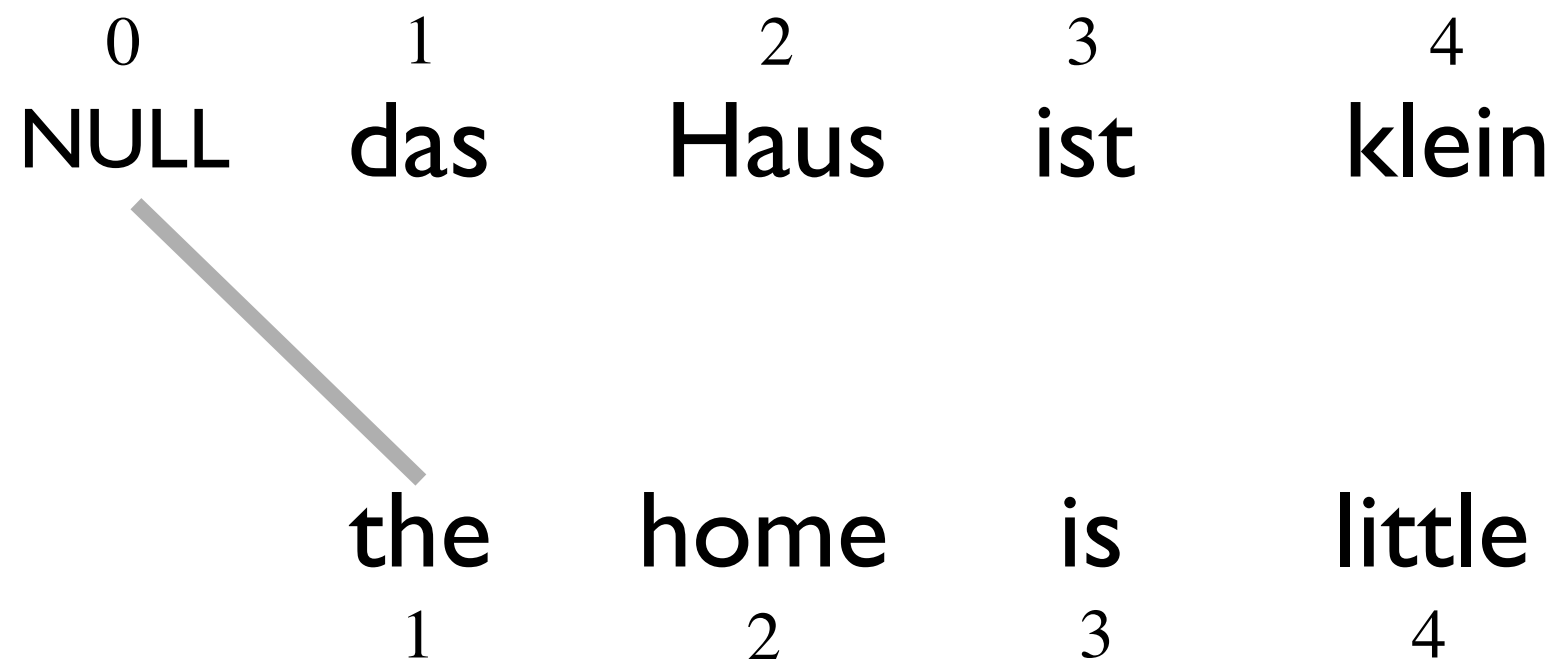
$$\begin{aligned} a_i^* &= \arg \max_{a_i=0}^n \frac{1}{1+n} p(e_i \mid f_{a_i}) \\ &= \arg \max_{a_i=0}^n p(e_i \mid f_{a_i}) \end{aligned}$$

Finding the Viterbi Alignment

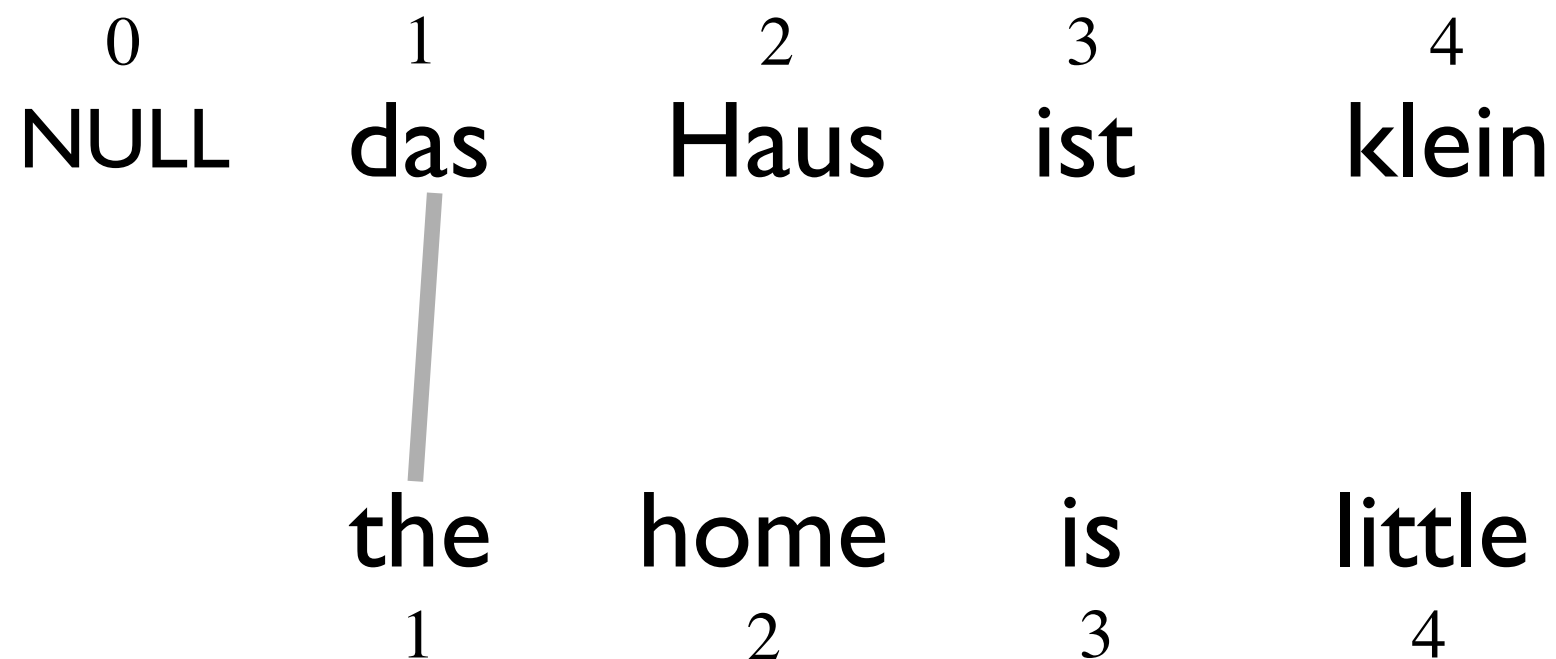
0	1	2	3	4
NULL	das	Haus	ist	klein

the	home	is	little
1	2	3	4

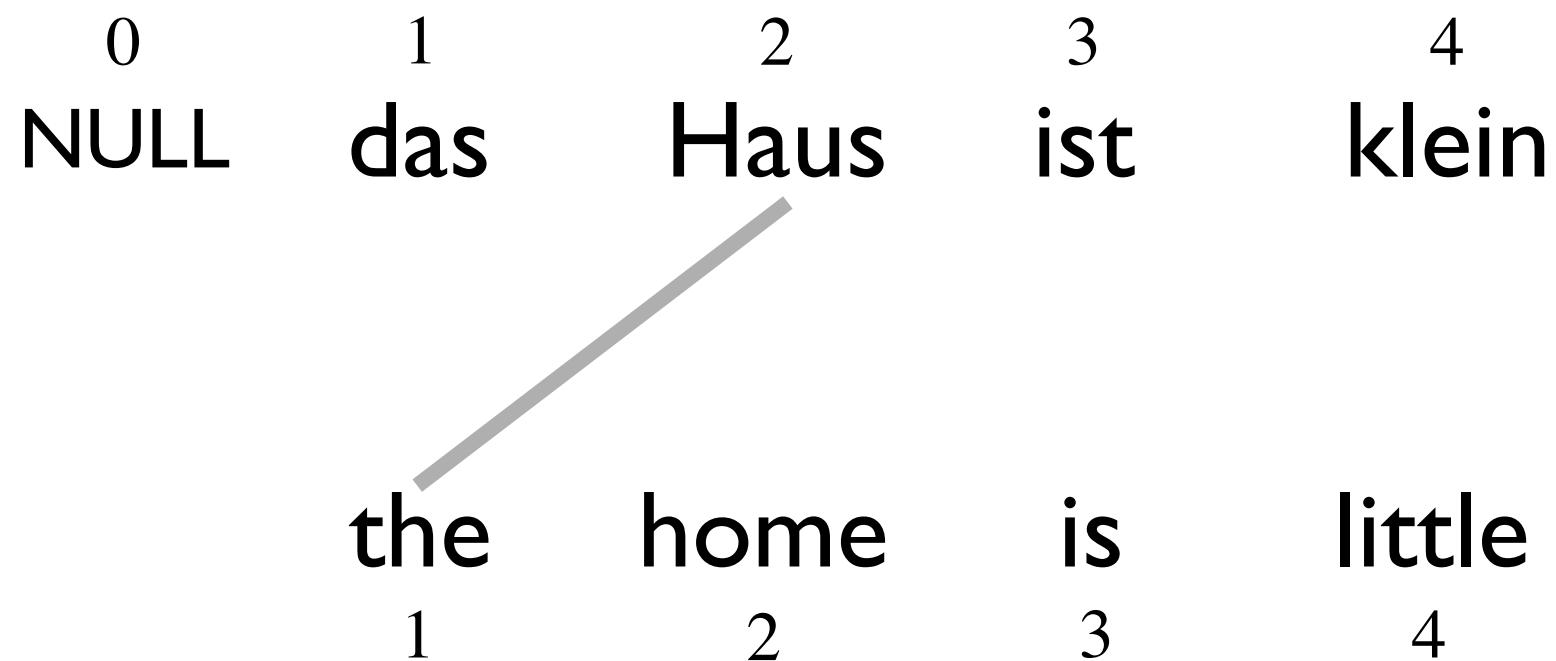
Finding the Viterbi Alignment



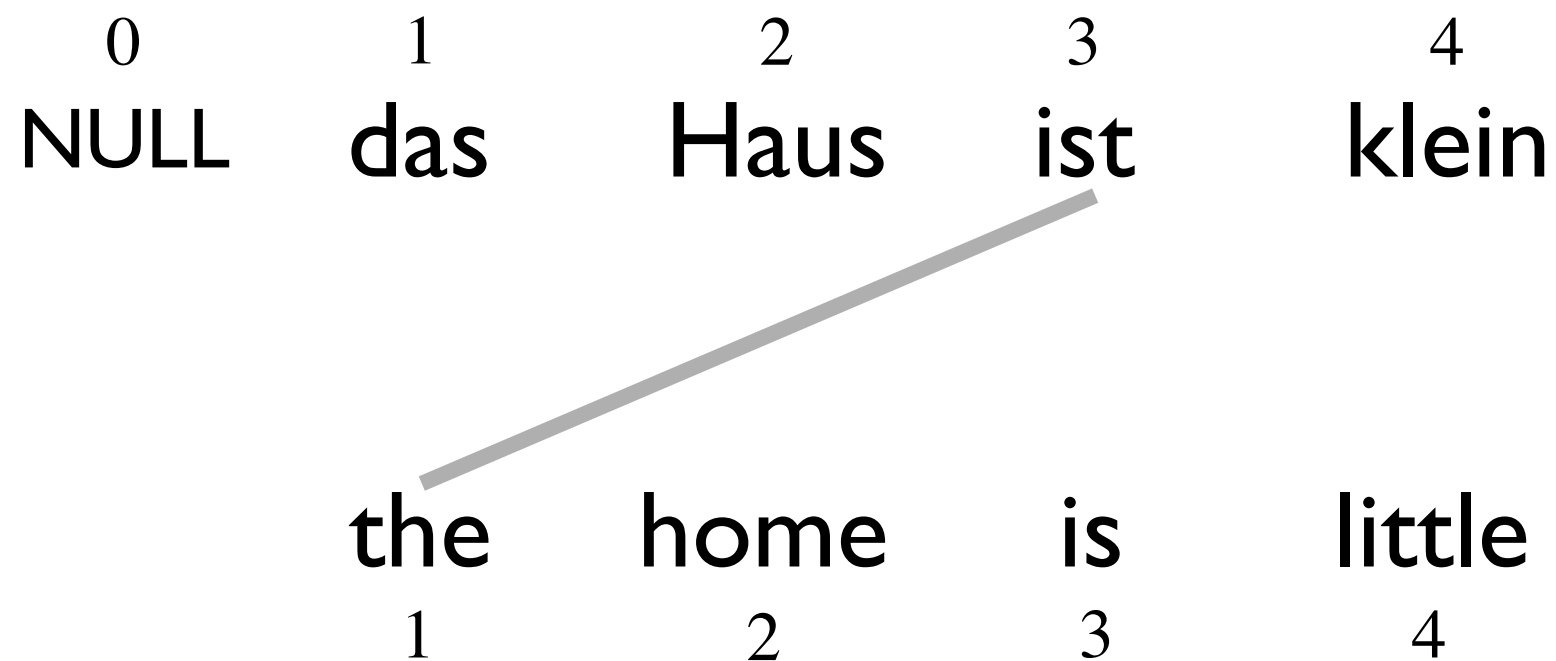
Finding the Viterbi Alignment



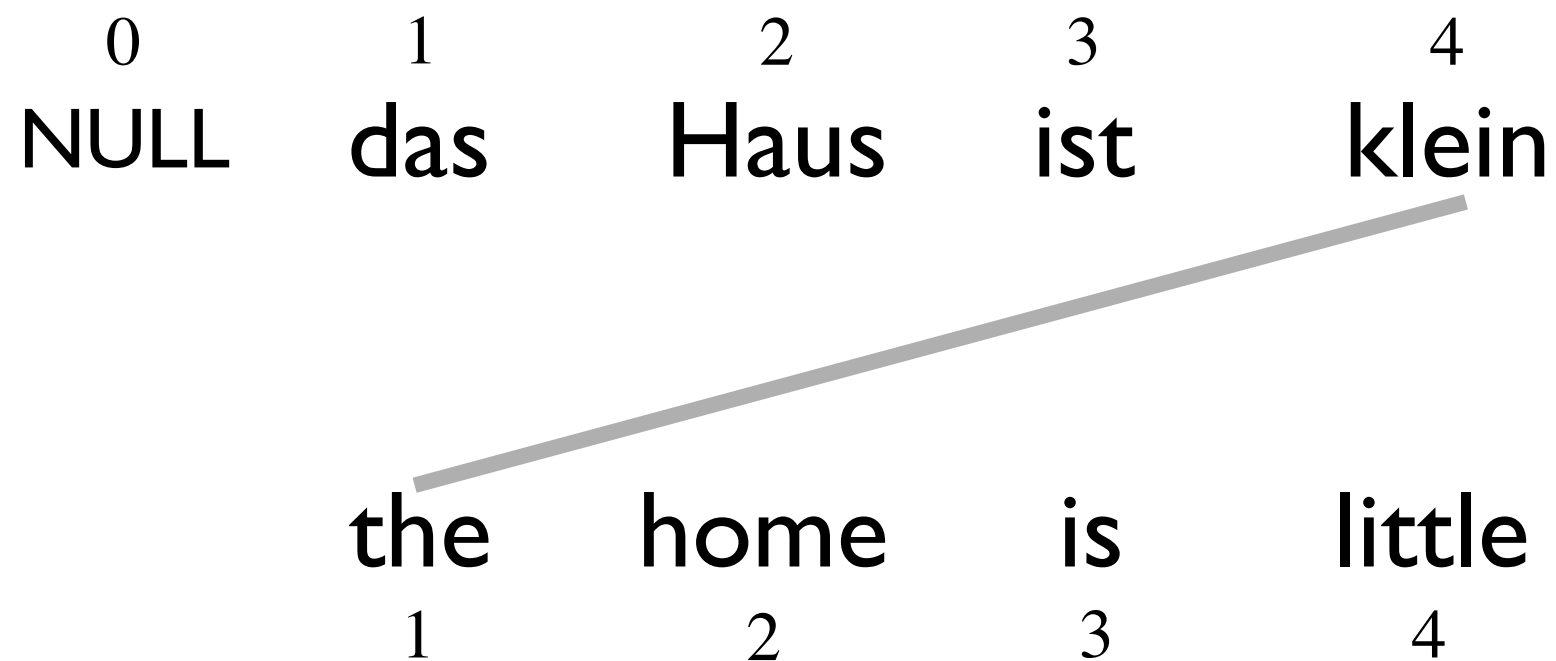
Finding the Viterbi Alignment



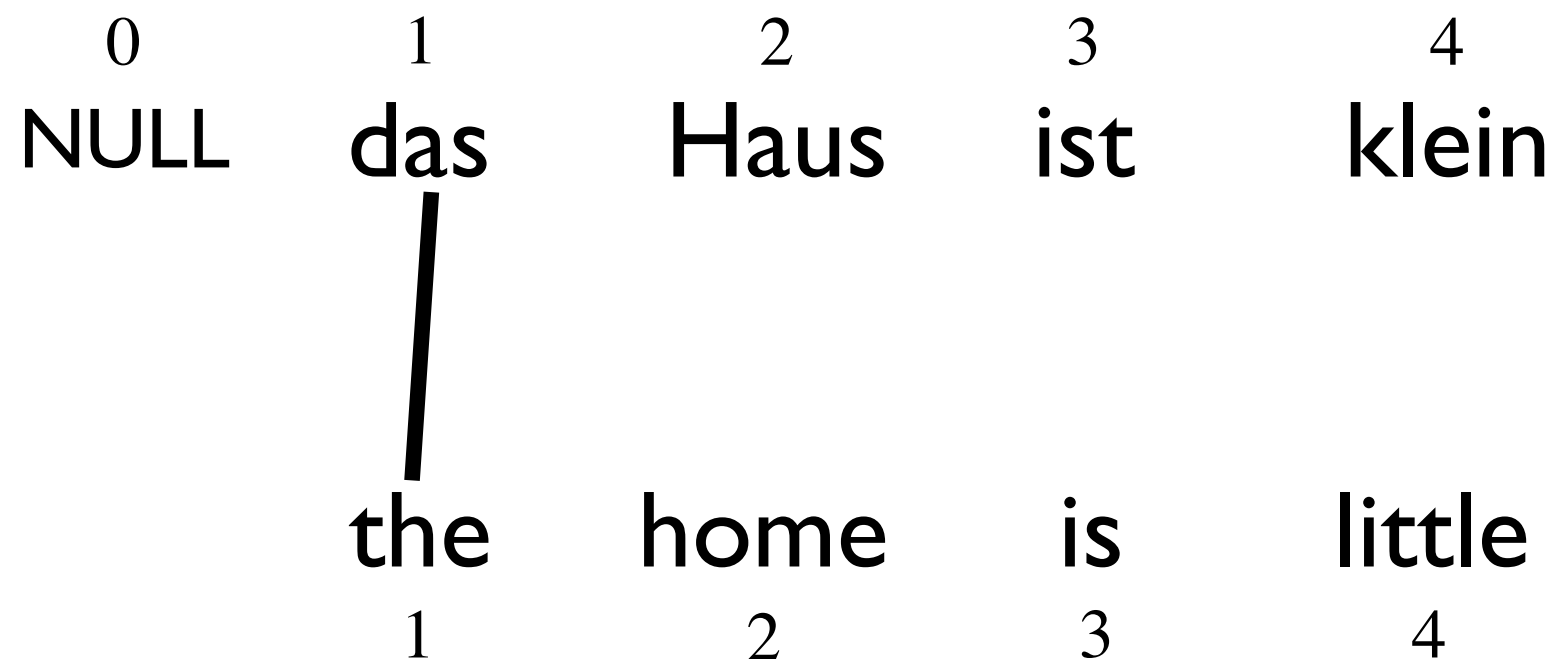
Finding the Viterbi Alignment



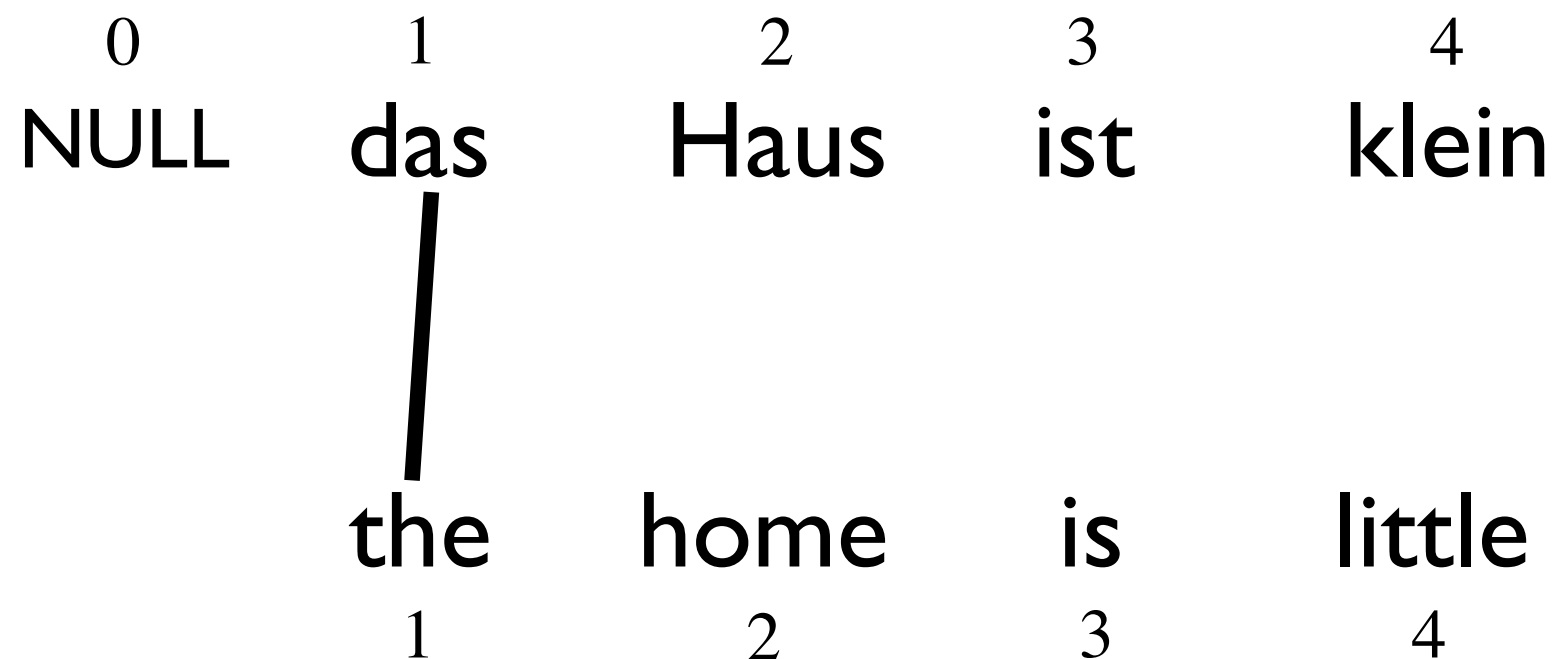
Finding the Viterbi Alignment



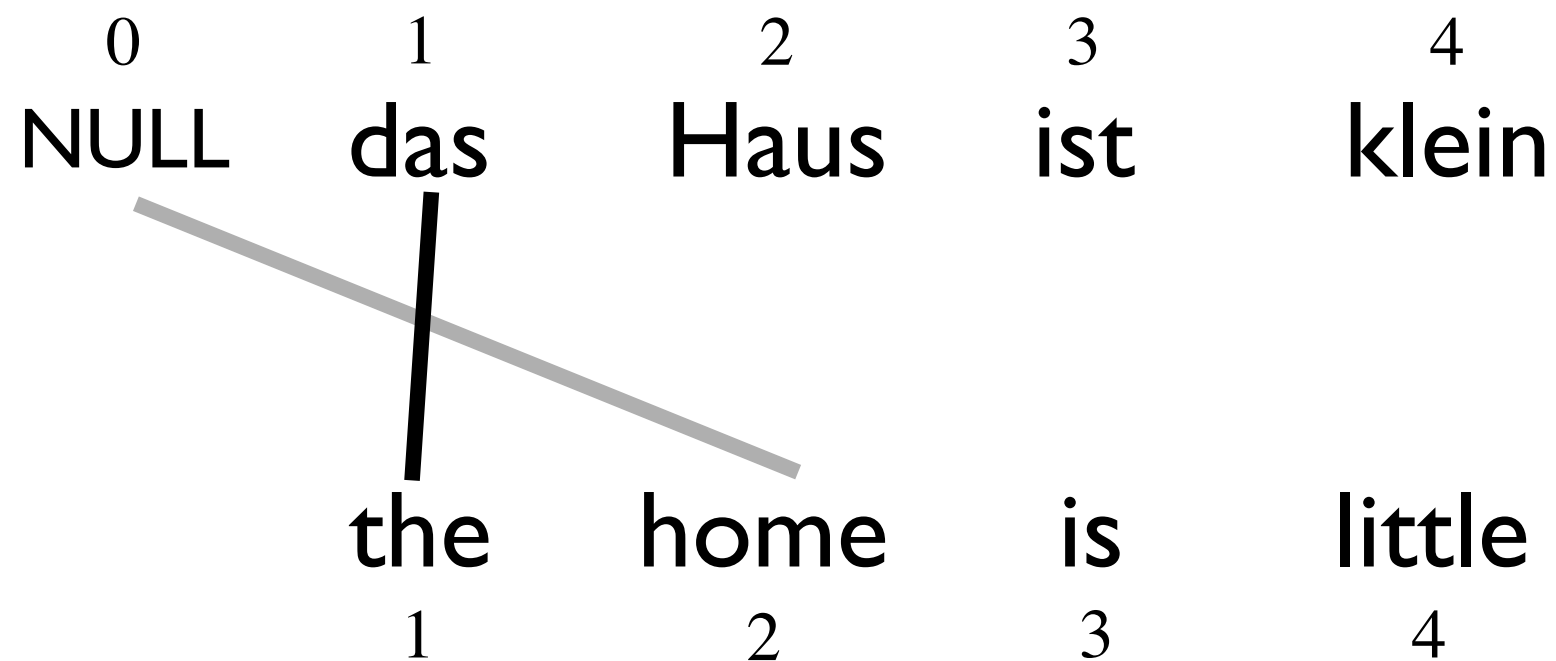
Finding the Viterbi Alignment



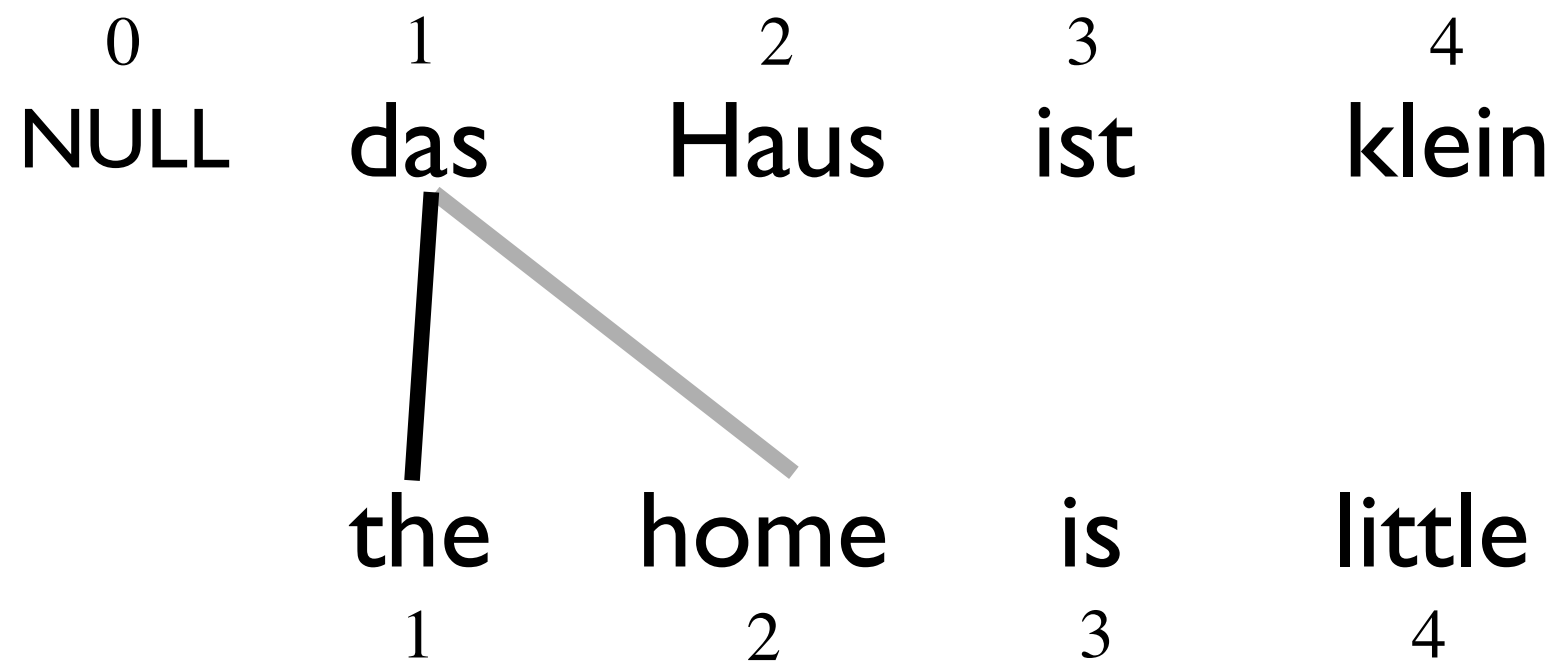
Finding the Viterbi Alignment



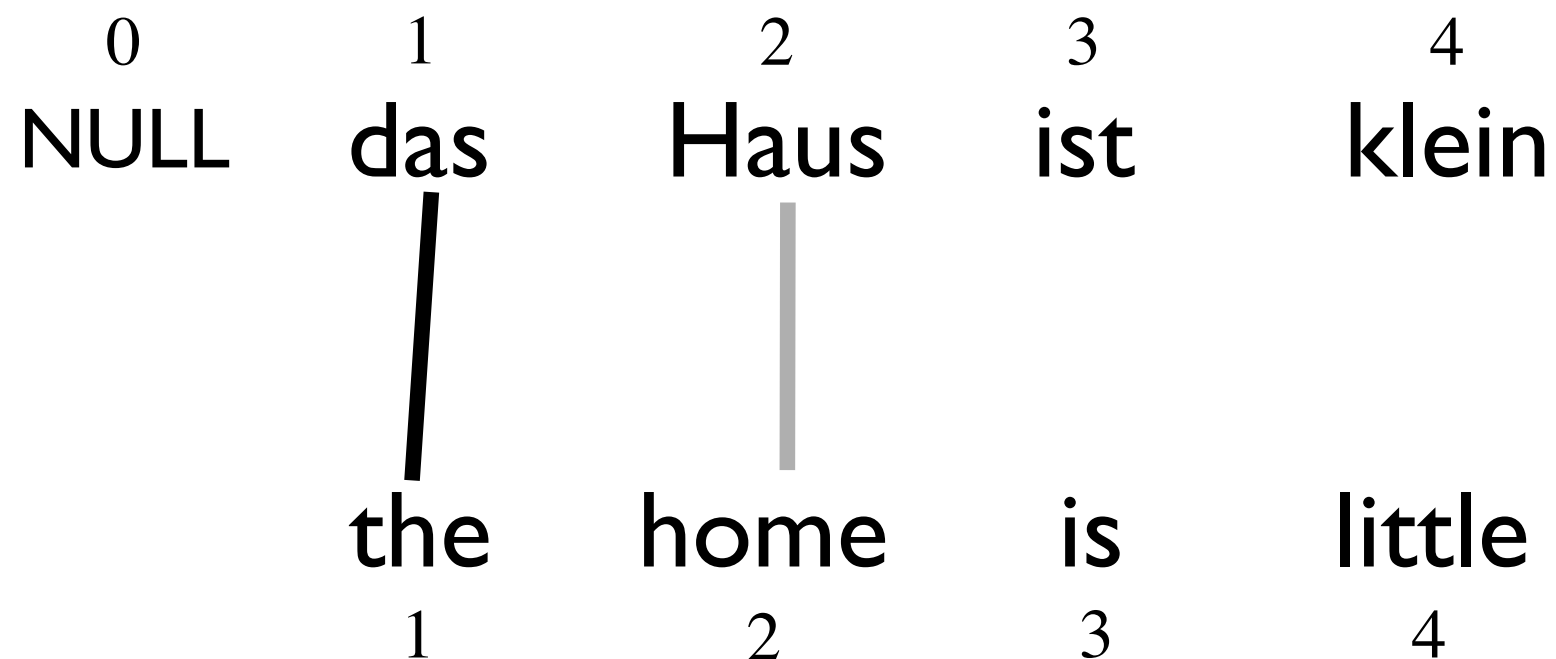
Finding the Viterbi Alignment



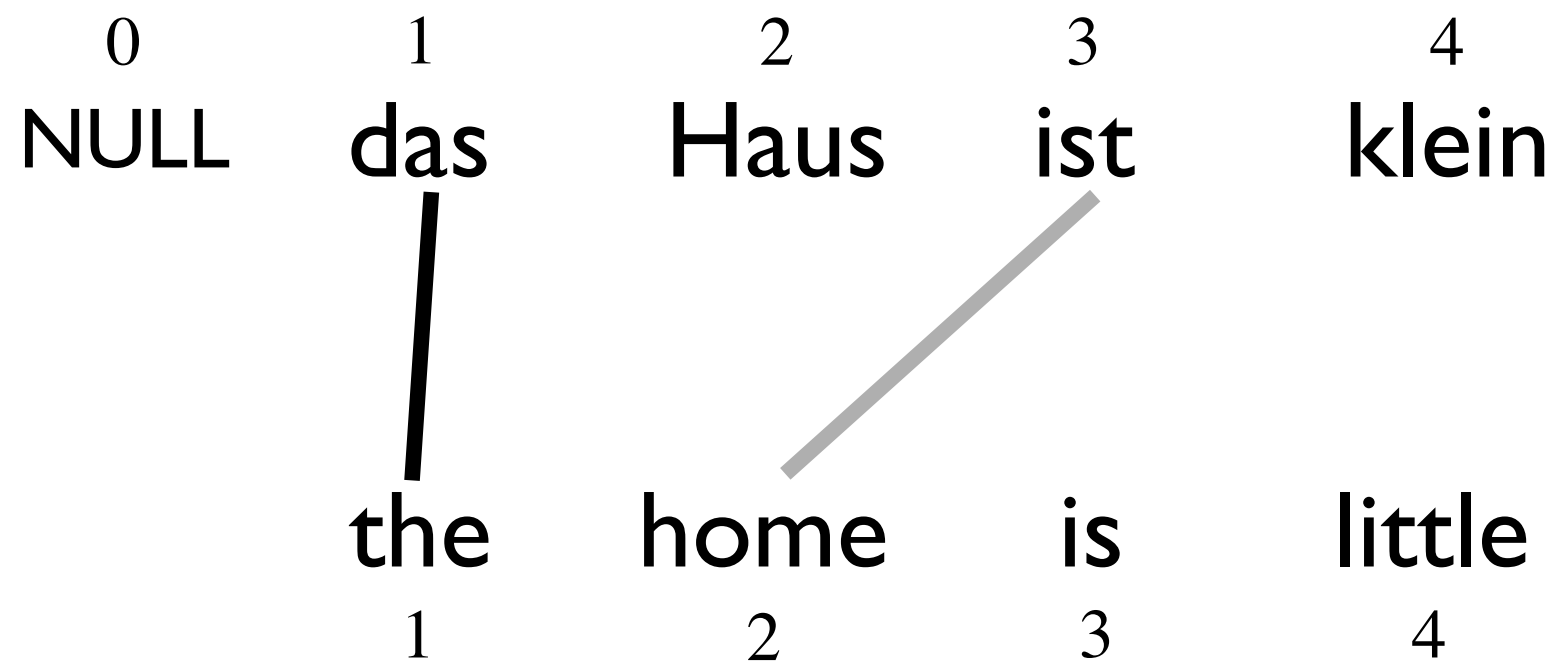
Finding the Viterbi Alignment



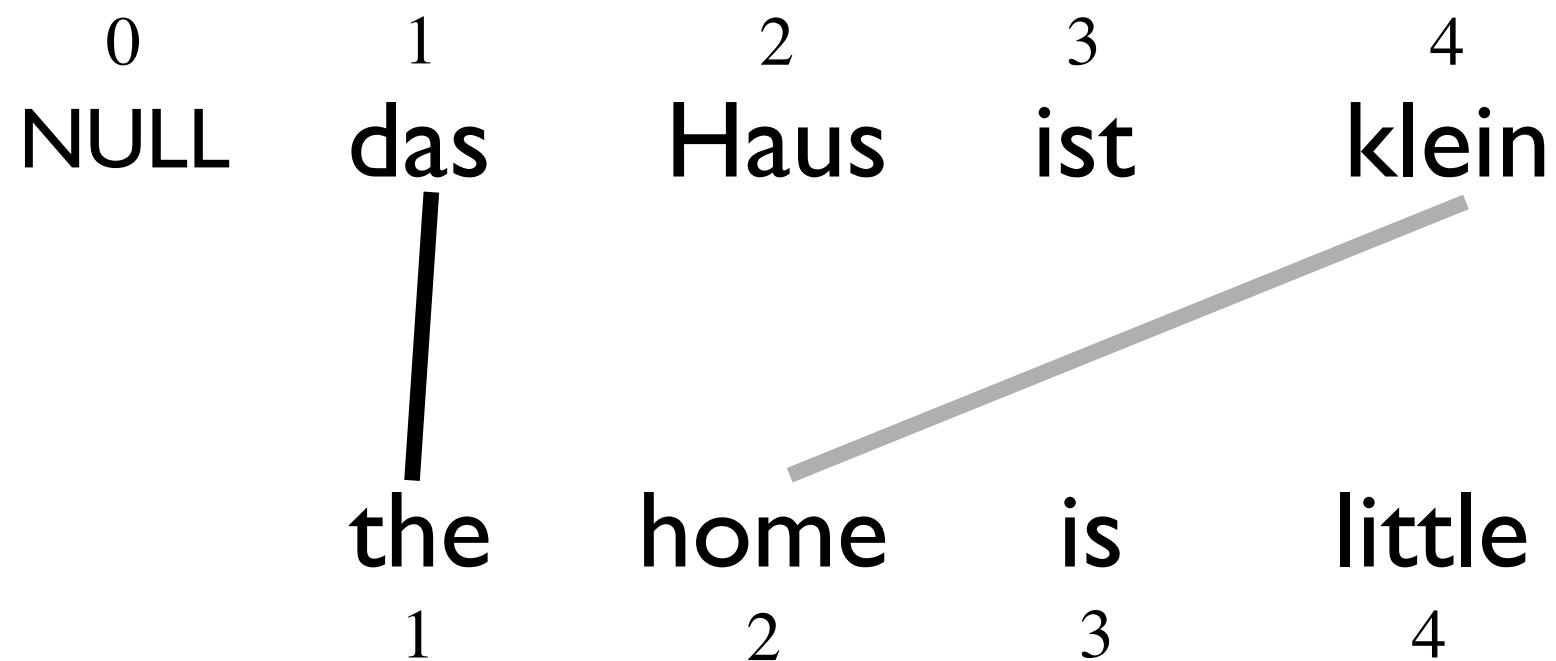
Finding the Viterbi Alignment



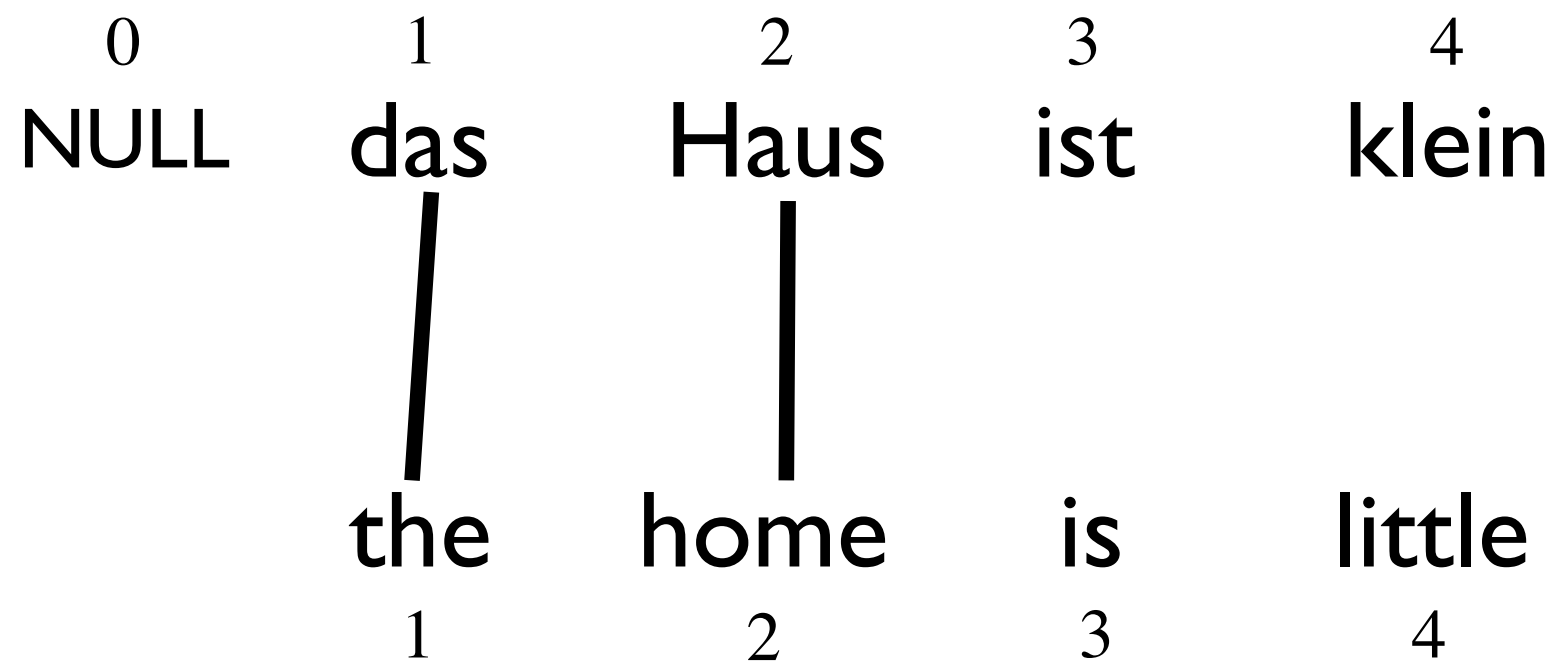
Finding the Viterbi Alignment



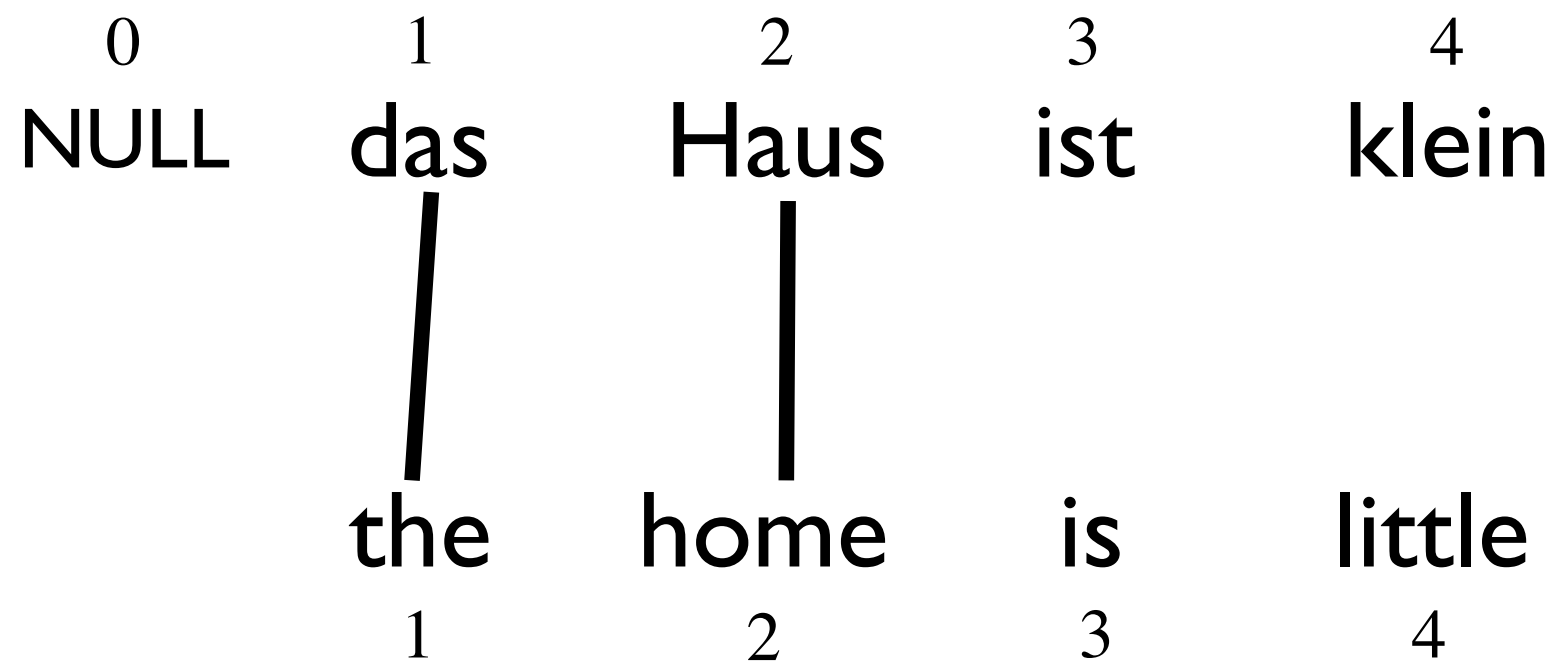
Finding the Viterbi Alignment



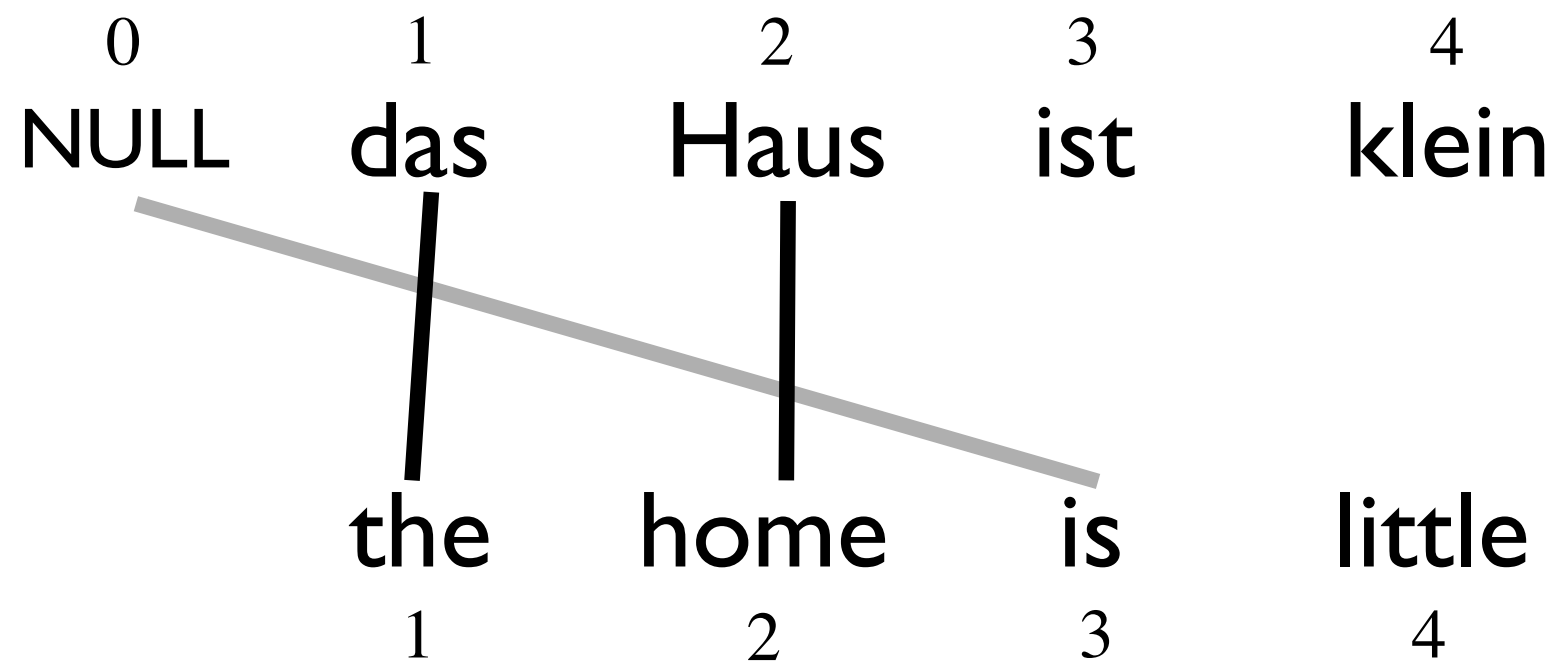
Finding the Viterbi Alignment



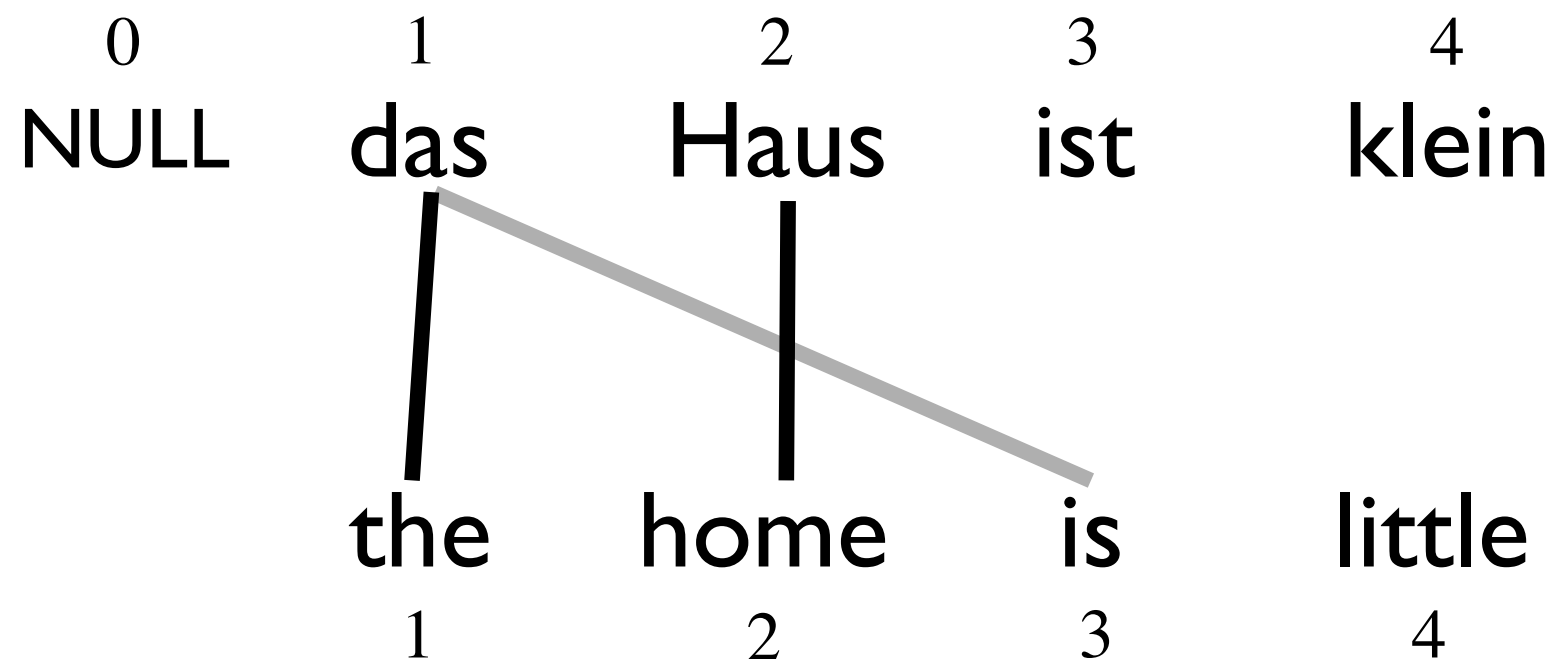
Finding the Viterbi Alignment



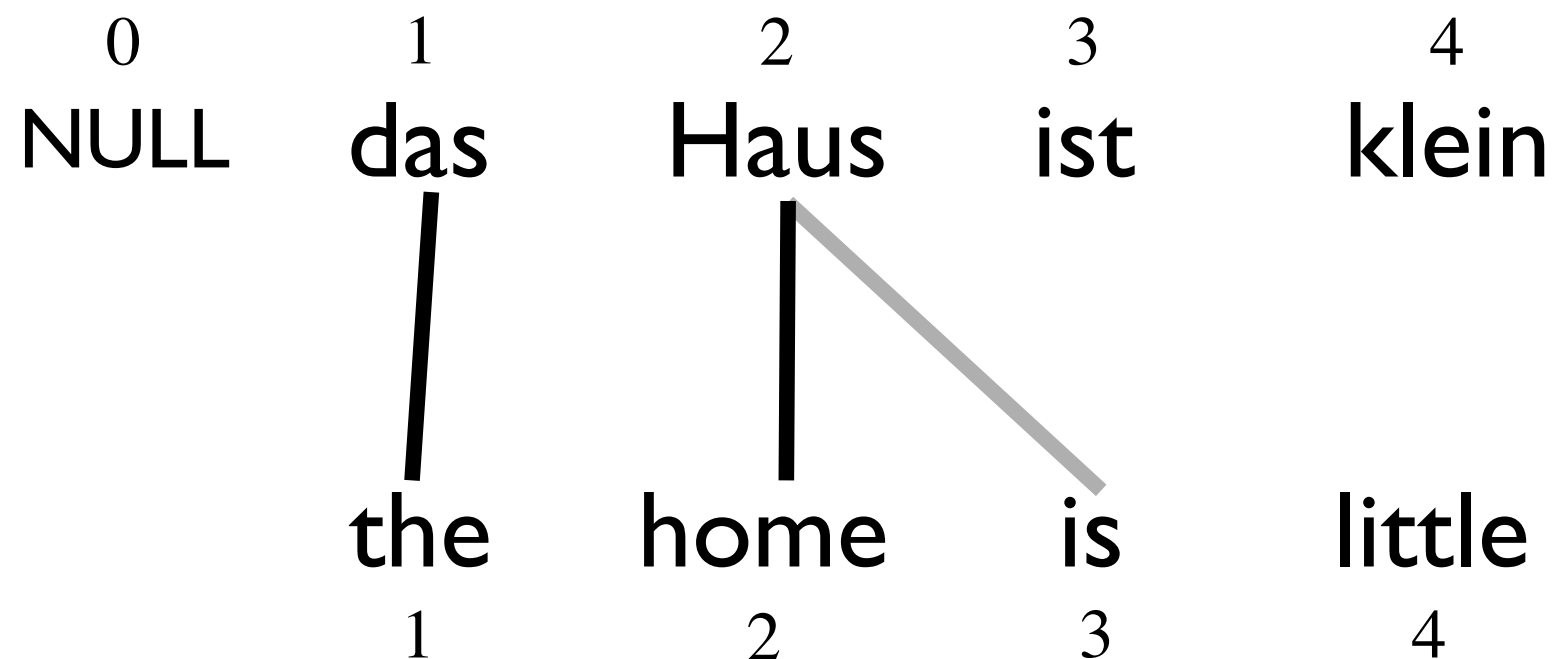
Finding the Viterbi Alignment



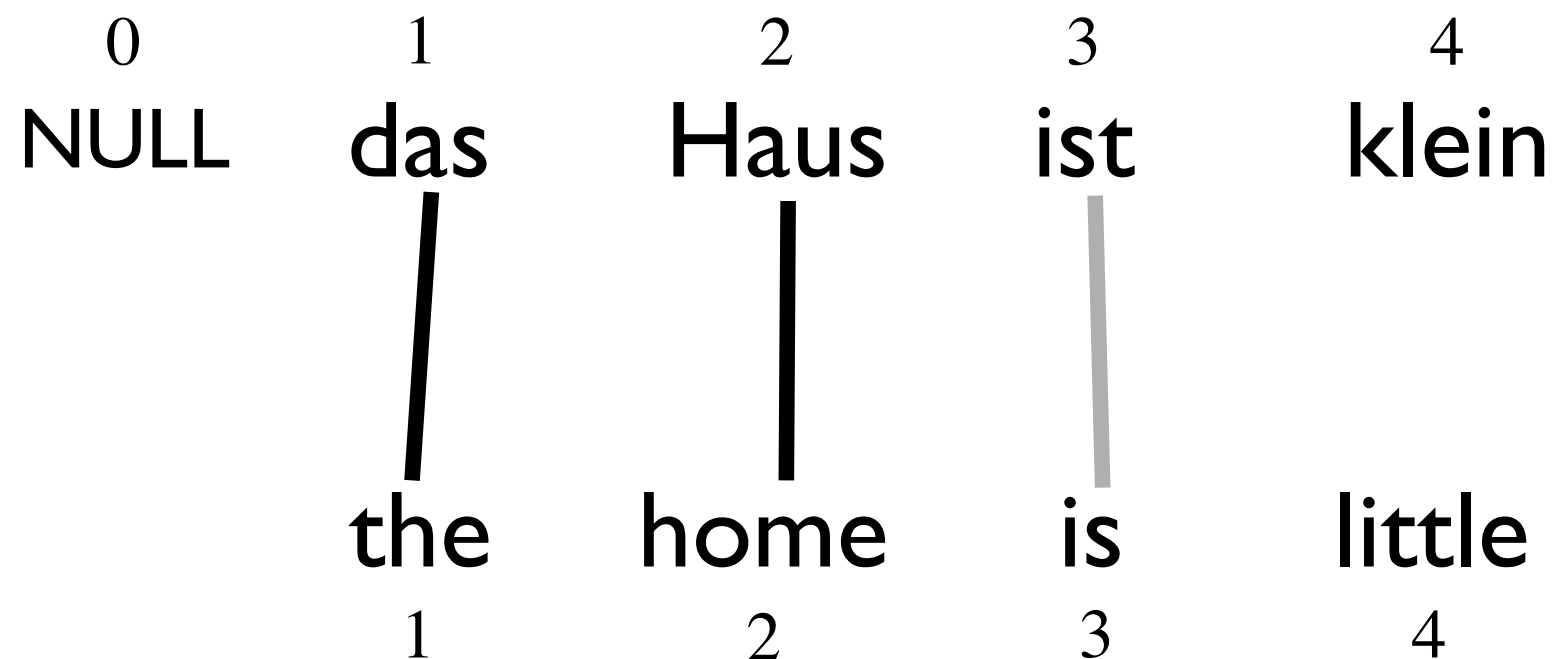
Finding the Viterbi Alignment



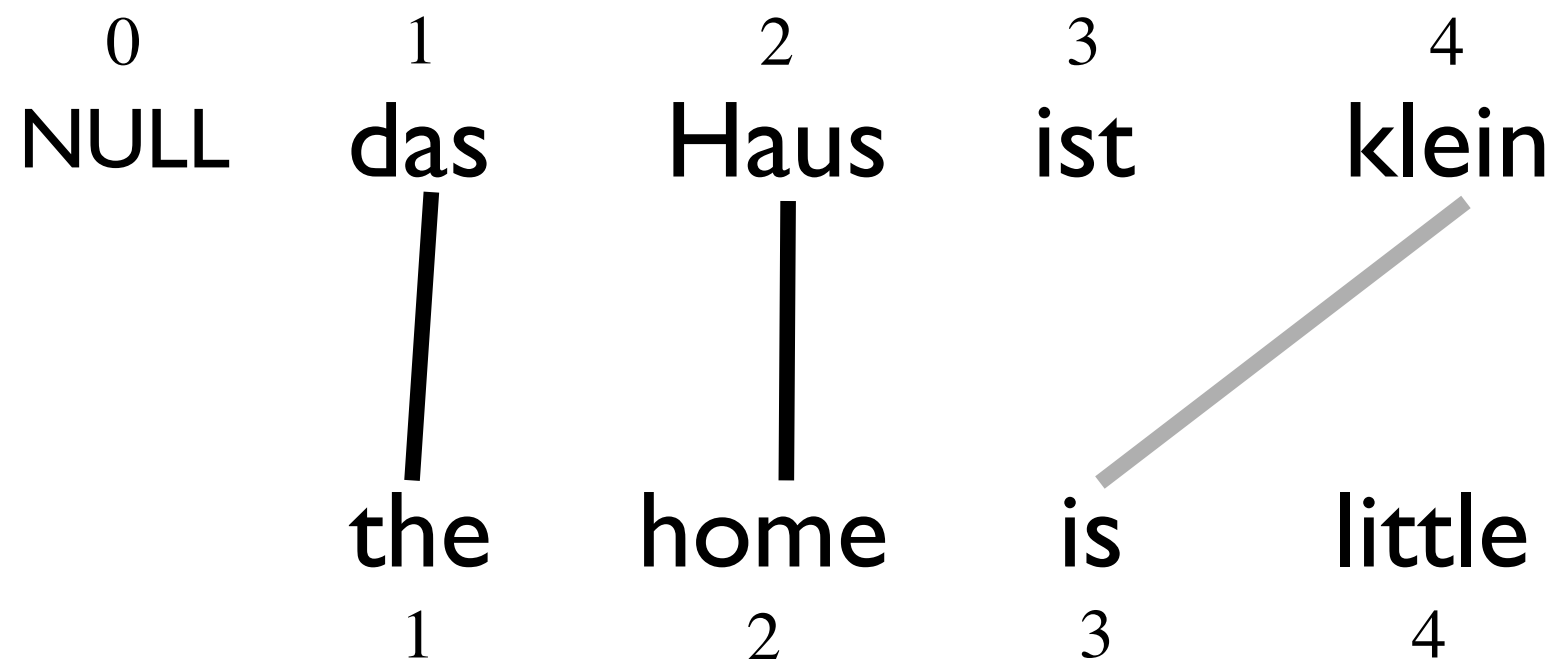
Finding the Viterbi Alignment



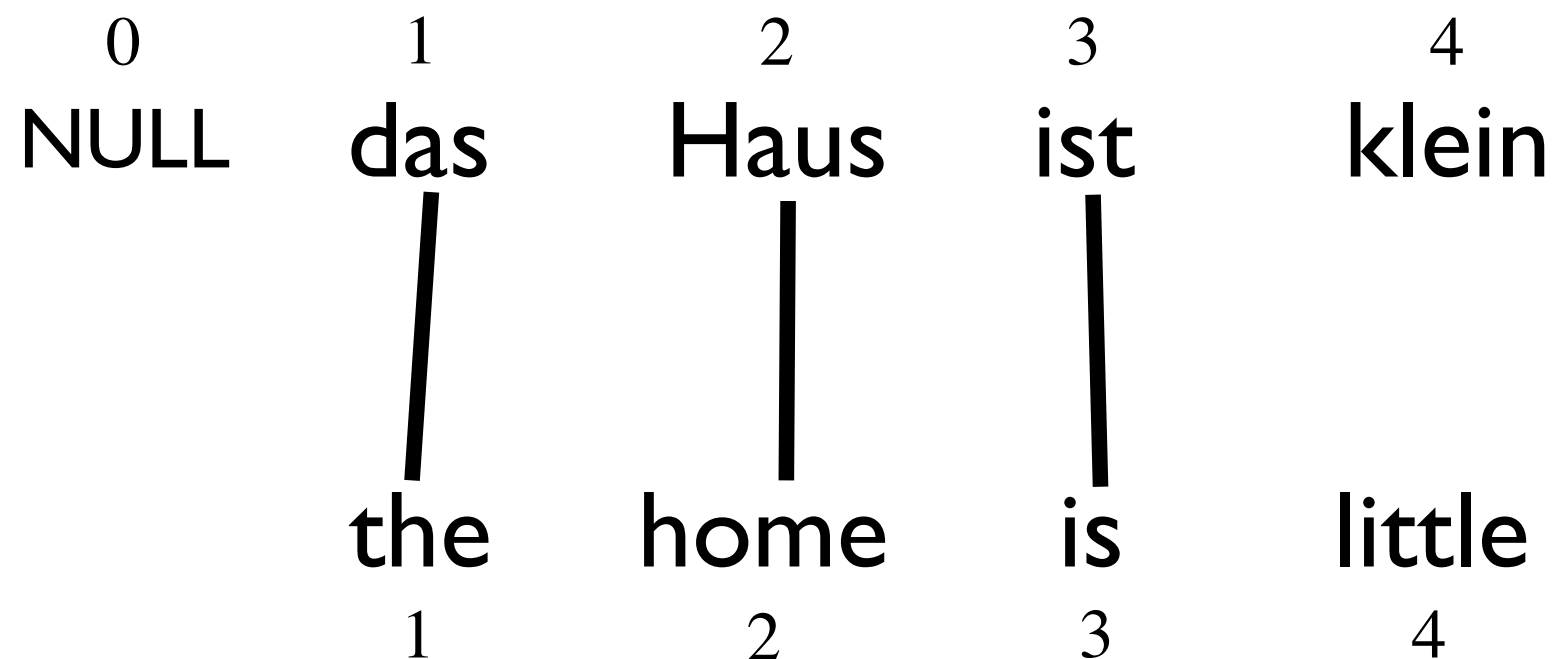
Finding the Viterbi Alignment



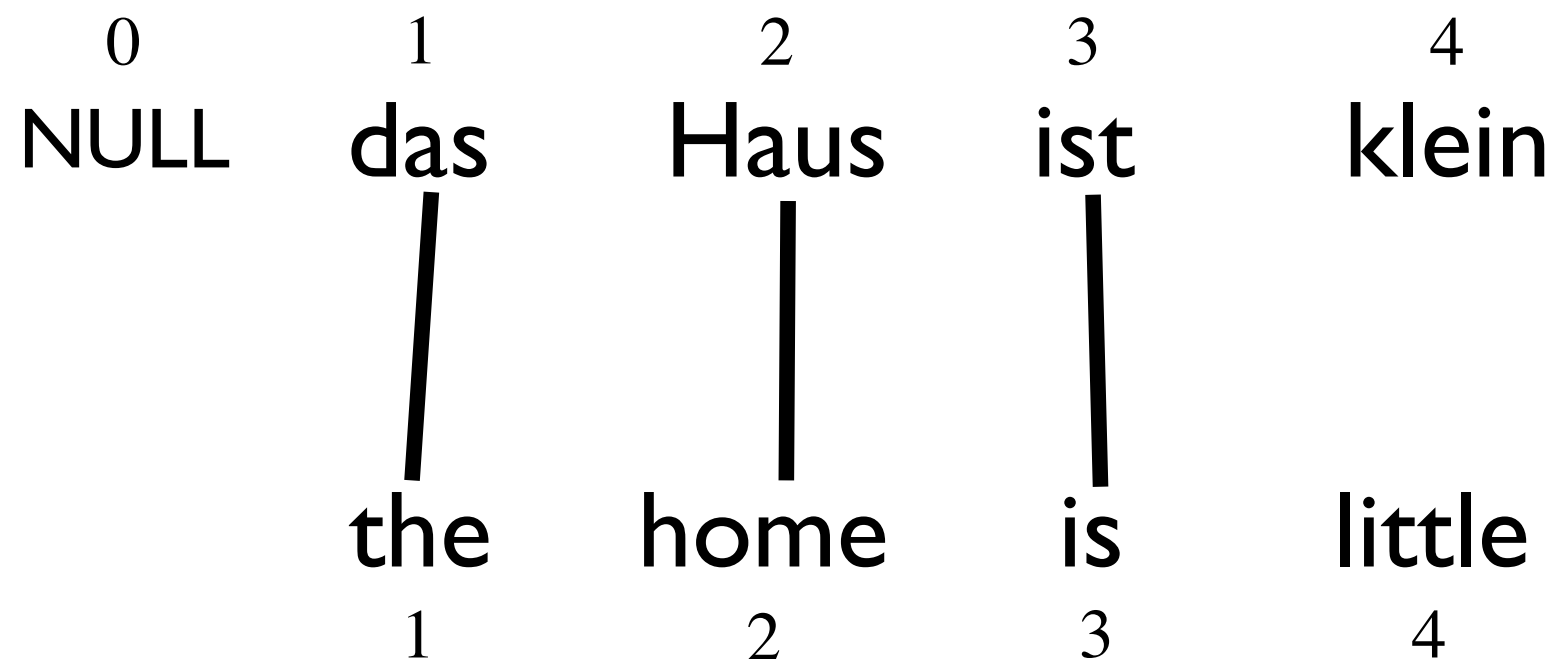
Finding the Viterbi Alignment



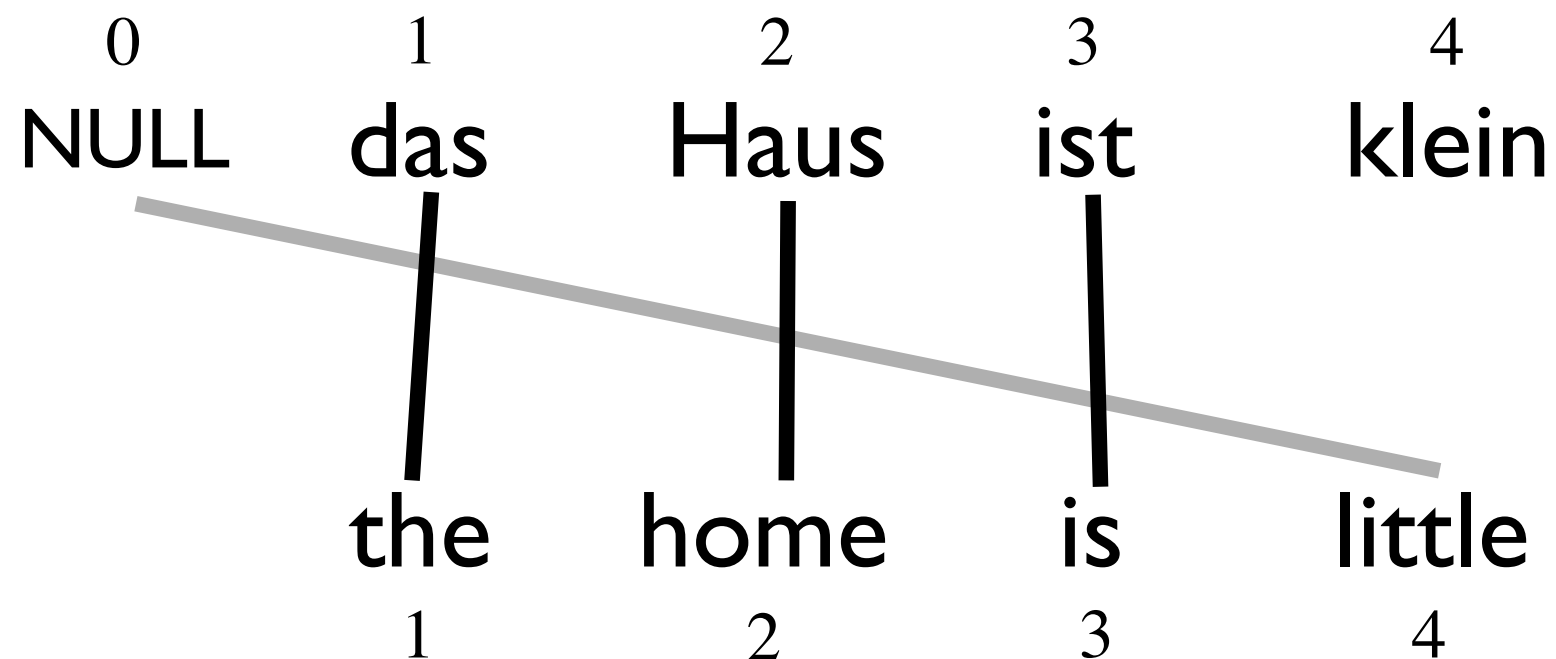
Finding the Viterbi Alignment



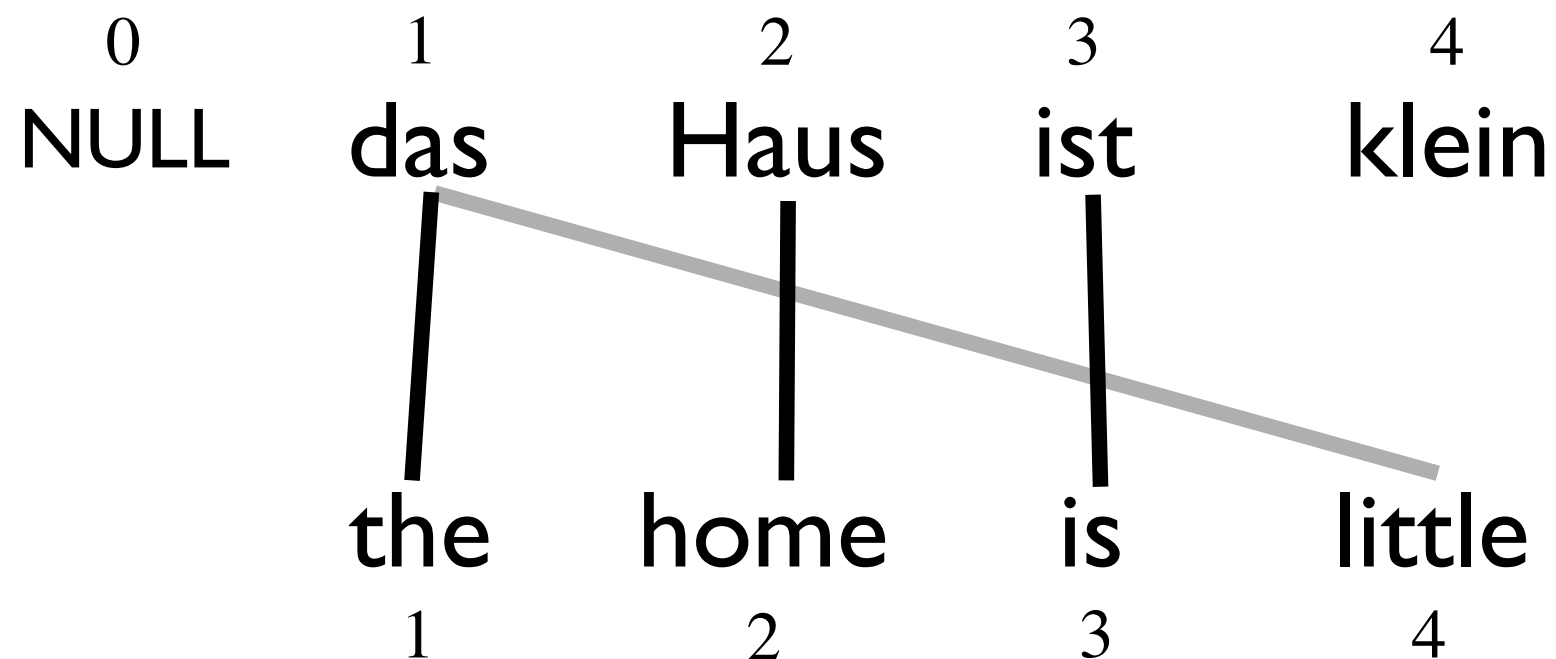
Finding the Viterbi Alignment



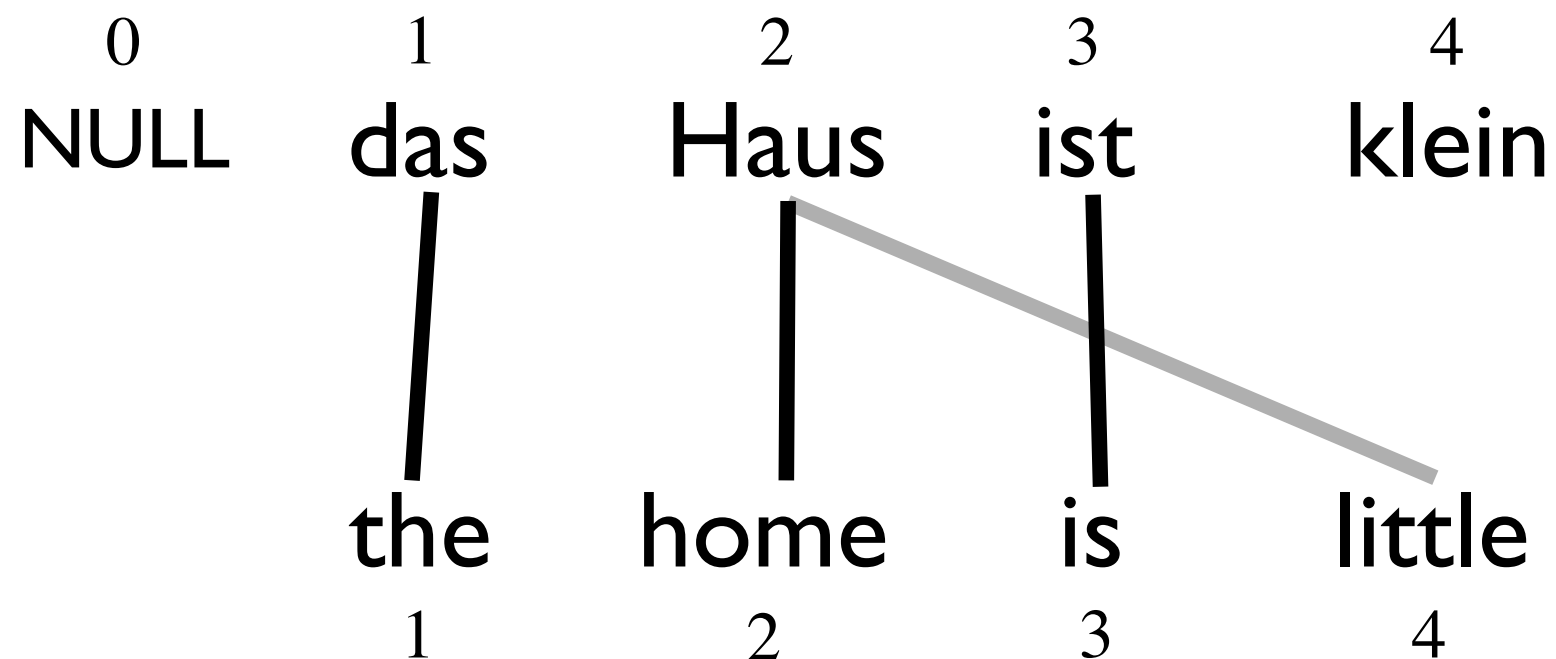
Finding the Viterbi Alignment



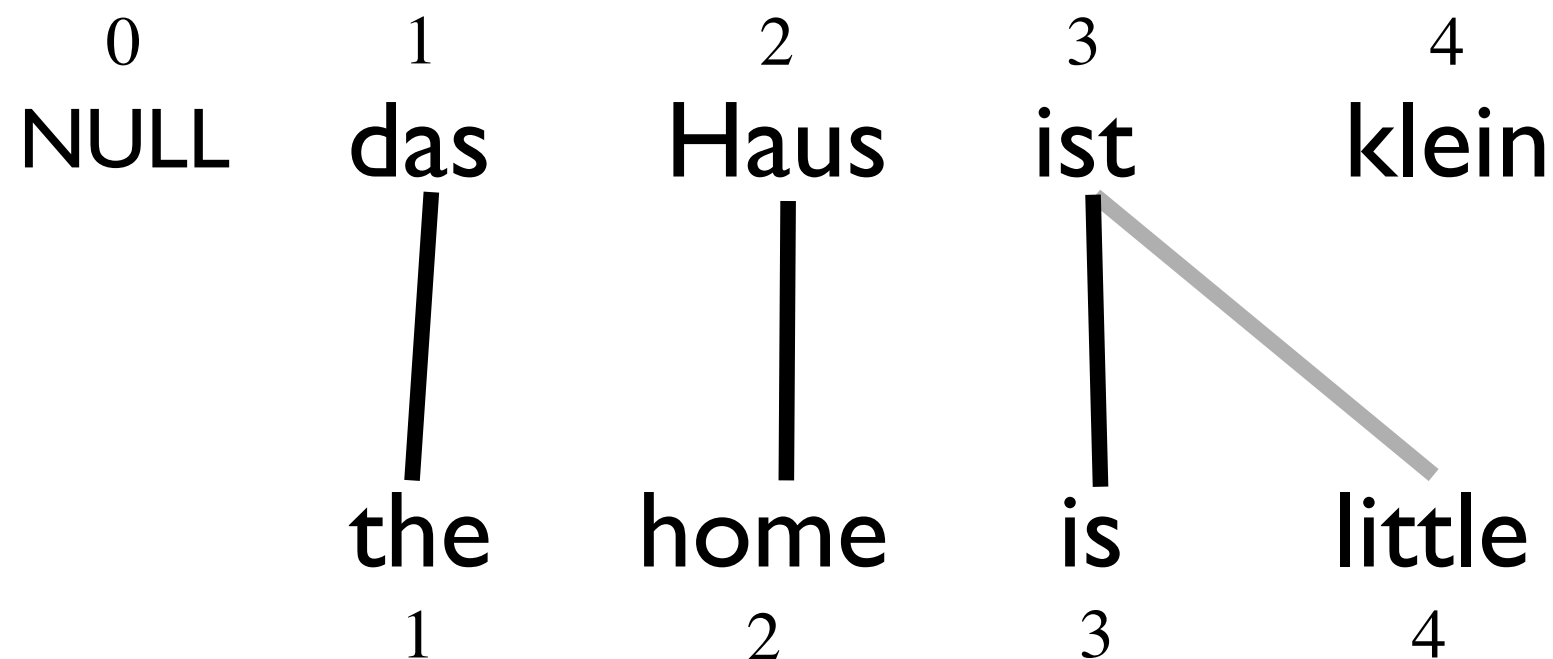
Finding the Viterbi Alignment



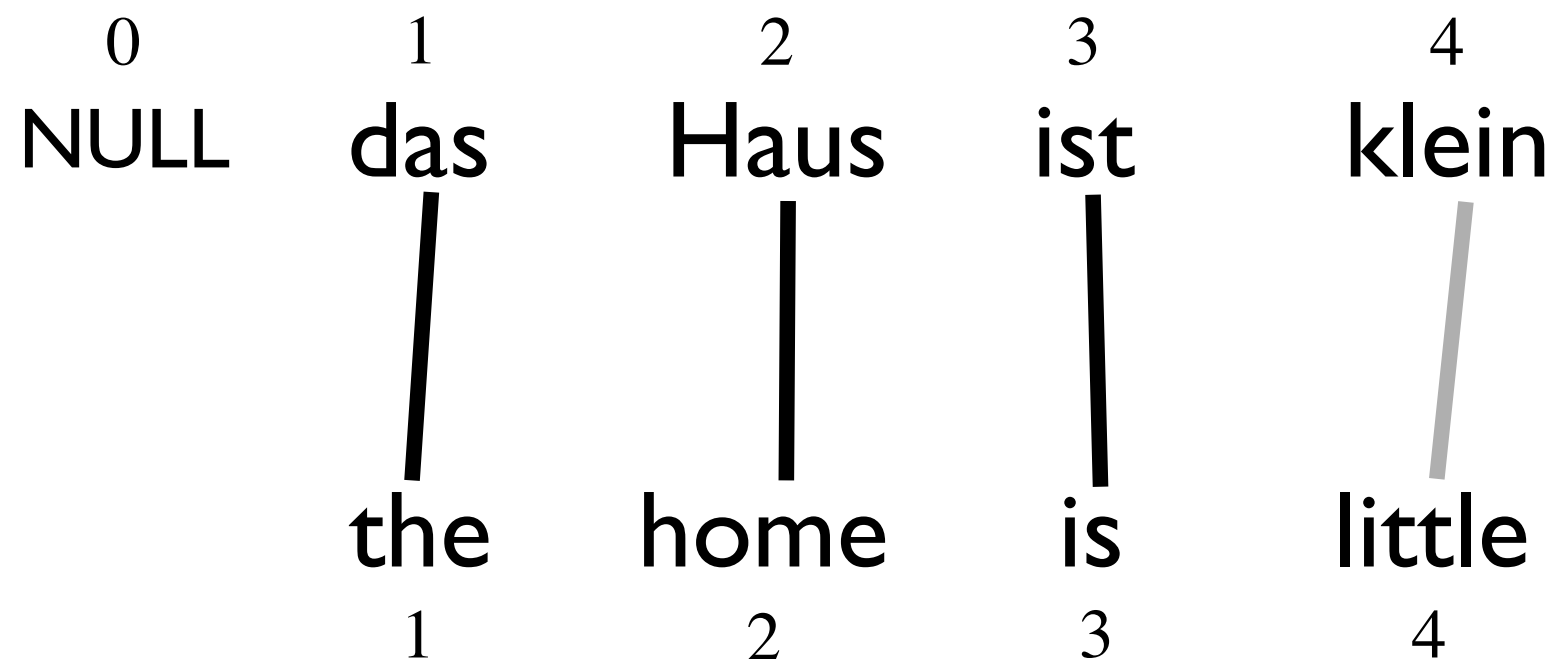
Finding the Viterbi Alignment



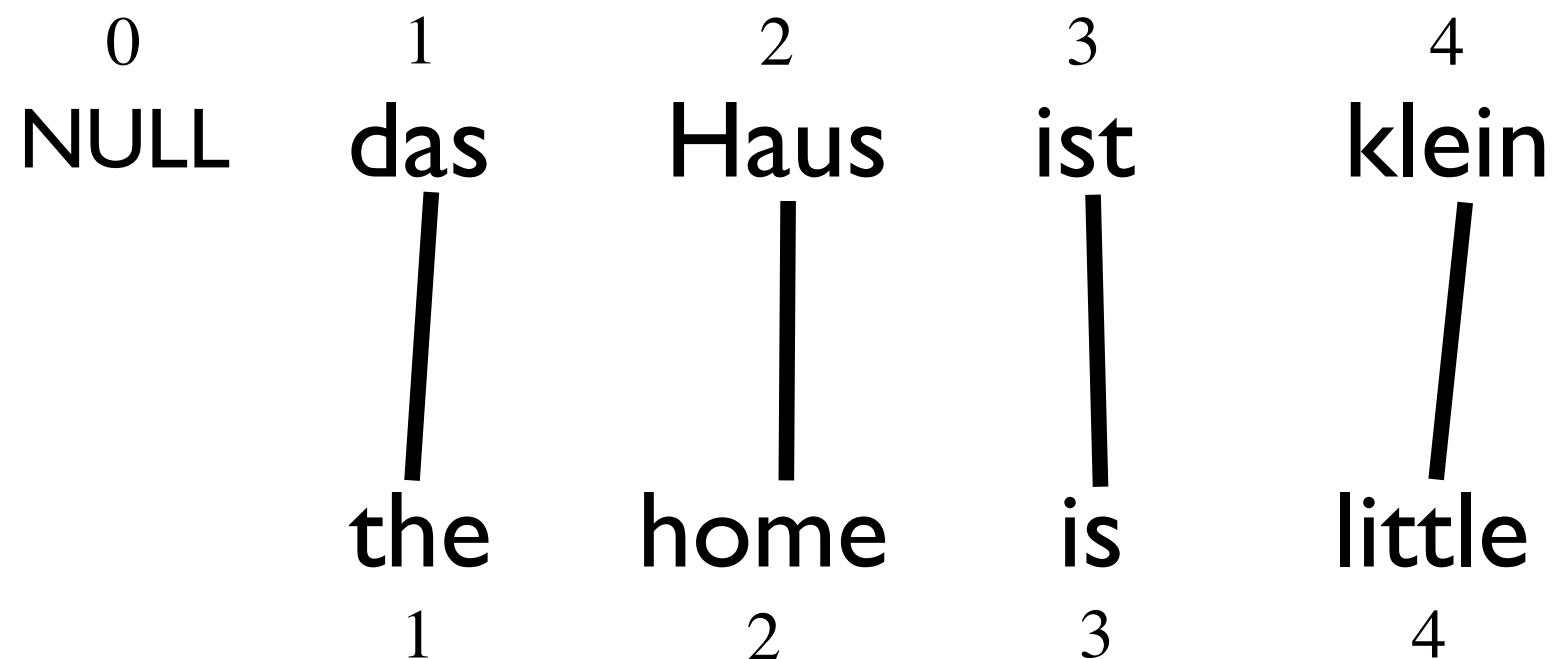
Finding the Viterbi Alignment



Finding the Viterbi Alignment



Finding the Viterbi Alignment



Learning Lexical Translation Models

- How do we learn the parameters $p(e | f)$
- “Chicken and egg” problem
 - If we had the alignments, we could estimate the parameters (MLE)
 - If we had parameters, we could find the most likely alignments



How to learn?

- Training data: tons of $(\mathbf{e}, \mathbf{f})$ pairs. Want to learn *theta*: lexical translation probabilities.

- If knew the alignments $\mathbf{a}$, MLE would be easy!

$$\max_{\theta} \sum_{(\mathbf{e}, \mathbf{f})} \log p(\mathbf{e} \mid \mathbf{a}, \mathbf{f}; \theta)$$

- But $\mathbf{a}$ is a latent variable. The marginal log-likelihood needs to sum-out the alignments. No closed form.

$$\max_{\theta} \sum_{(\mathbf{e}, \mathbf{f})} \log p(\mathbf{e} \mid \mathbf{f}; \theta)$$

$$= \sum_{(\mathbf{e}, \mathbf{f})} \log \left(\sum_{\mathbf{a}} p(\mathbf{a} \mid \mathbf{f}; \theta) \times p(\mathbf{e} \mid \mathbf{a}, \mathbf{f}; \theta) \right)$$

EM Algorithm

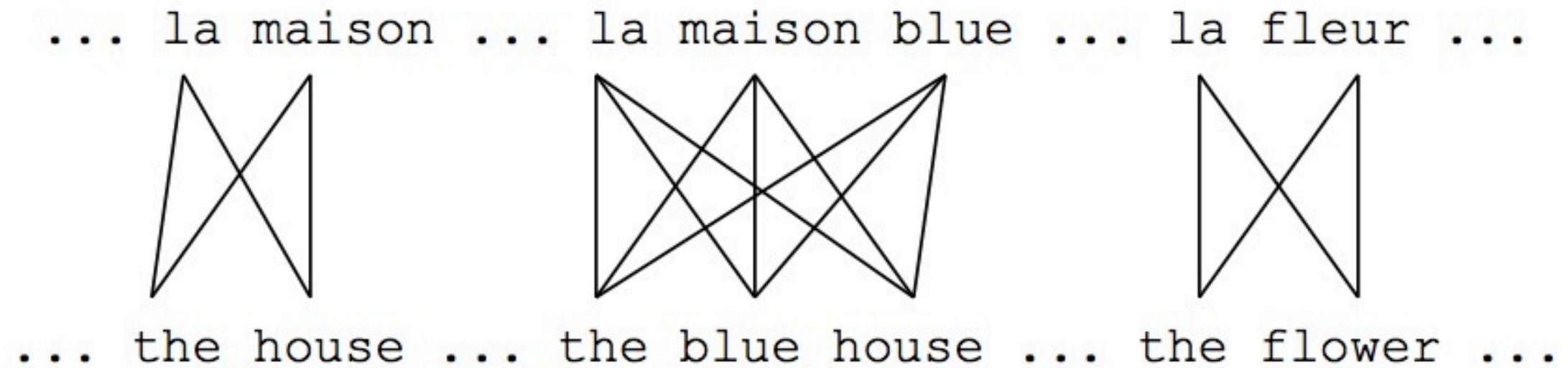
- pick some random (or uniform) parameters
- Repeat until you get bored (~ 5 iterations for lexical translation models)

- using your current parameters, compute “expected” alignments for every target word token in the training data

$$p(a_i \mid \mathbf{e}, \mathbf{f}) \quad (\text{on board})$$

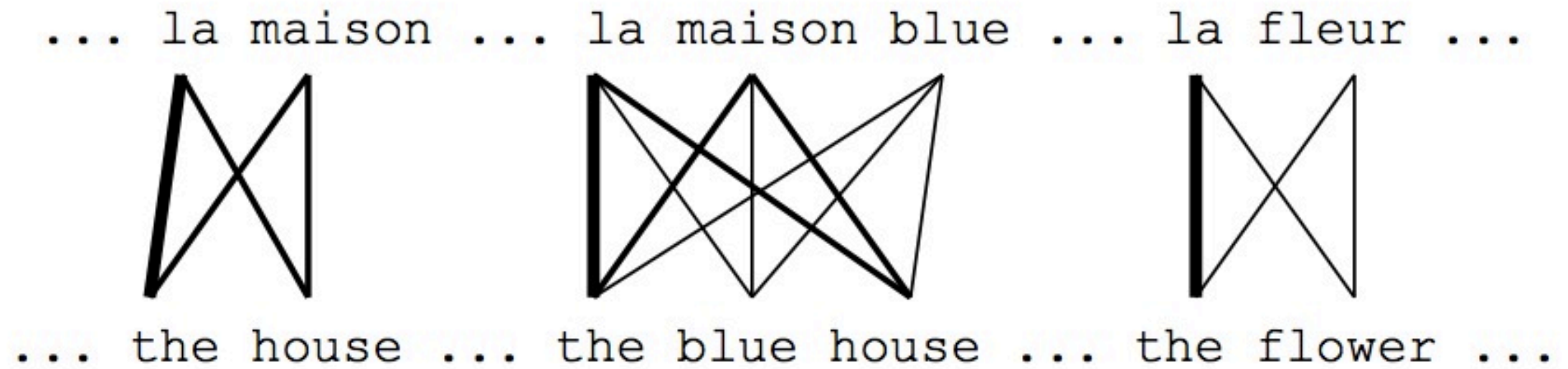
- keep track of the expected number of times f translates into e throughout the whole corpus
- keep track of the expected number of times that f is used as the source of any translation
- use these expected counts as if they were “real” counts in the standard MLE equation

EM for Model I



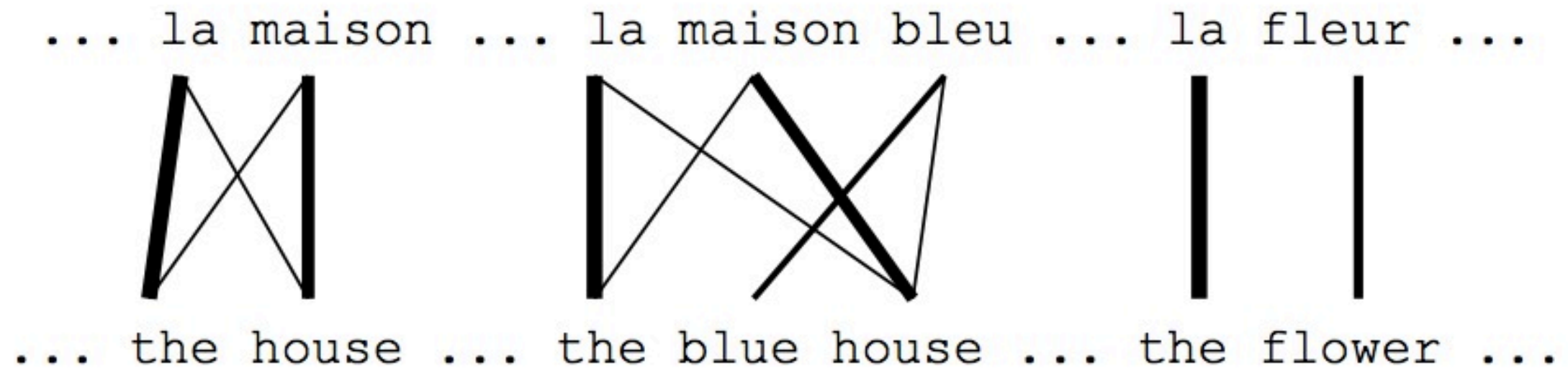
- Initial step: all alignments equally likely
- Model learns that, e.g., **la** is often aligned with **the**

EM for Model I



- After one iteration
- Alignments, e.g., between **la** and **the** are more likely

EM for Model 1



- After another iteration
- It becomes apparent that alignments, e.g., between **fleur** and **flower** are more likely (pigeon hole principle)

EM for Model 1

... la maison ... la maison bleu ... la fleur ...
/ | | X | |
... the house ... the blue house ... the flower ...

- Convergence
- Inherent hidden structure revealed by EM

EM for Model 1

... la maison ... la maison bleu ... la fleur ...
/ | | X | |
... the house ... the blue house ... the flower ...



$p(\text{la}|\text{the}) = 0.453$
 $p(\text{le}|\text{the}) = 0.334$
 $p(\text{maison}|\text{house}) = 0.876$
 $p(\text{bleu}|\text{blue}) = 0.563$
...

- Parameter estimation from the aligned corpus

Convergence

das Haus
the house

das Buch
the book

ein Buch
a book

<i>e</i>	<i>f</i>	initial	1st it.	2nd it.	3rd it.	...	final
the	das	0.25	0.5	0.6364	0.7479	...	1
book	das	0.25	0.25	0.1818	0.1208	...	0
house	das	0.25	0.25	0.1818	0.1313	...	0
the	buch	0.25	0.25	0.1818	0.1208	...	0
book	buch	0.25	0.5	0.6364	0.7479	...	1
a	buch	0.25	0.25	0.1818	0.1313	...	0
book	ein	0.25	0.5	0.4286	0.3466	...	0
a	ein	0.25	0.5	0.5714	0.6534	...	1
the	haus	0.25	0.5	0.4286	0.3466	...	0
house	haus	0.25	0.5	0.5714	0.6534	...	1

Now what?

- Decoding: find the best translation for a sentence
- Alignments: find the best alignment
- Evaluation?
- Other approaches?