

SUPPLEMENTARY MATERIAL

Detecting contact in language trees: a Bayesian phylogenetic model with horizontal transfer

S1 The contacTrees model

The contacTrees model is defined by the following posterior distribution (cf. main article, Equation 5):

$$\begin{aligned} P(T, \mathcal{C}, Z, \beta, \Gamma, \theta_X, \theta_T | X) &\propto \overbrace{P(X | T, \mathcal{C}, Z, \theta_X)}^{\text{likelihood}} \cdot \overbrace{P(Z | \mathcal{C}, \beta)}^{\text{borrowing prior}} \cdot \overbrace{P(\mathcal{C} | T, \Gamma)}^{\text{contact prior}} \\ &\quad \cdot P(T | \theta_T) \cdot P(\beta, \Gamma, \theta_X, \theta_T) \end{aligned} \quad (1)$$

$$(2)$$

In this section, we describe aspects of the likelihood and the contact prior in more detail.

S1.1 Contact Prior

The contact prior is defined by a Poisson process, i.e. a continuous random process with a contact edges appearing over time at a specific rate γ (the contact rate). We propose two models of how this rate can define a prior over the tree: the linear contact prior and the quadratic contact prior. We explain the quadratic contact prior first since it seems mathematically very natural. Later, we will argue why the linear contact prior is more appropriate for most applications.

Quadratic contact prior. Since a contact edge always involves two languages – a donor and a receiver – it seems plausible to assume a specific rate of contact for each pair of languages. Consequently, the number of expected contact edges in a time interval would grow quadratically with the number of lineages at that time. More precisely, we enumerate the intervals between the internal nodes in a tree from the root to the leaves by $i \in \{2, \dots, n\}$, assuming for simplicity that the leaves are all contemporary. Since we start the enumeration at 2, i also represents the number of lineages in each interval. We further denote the length of the i -th interval by τ_i and the contact edges in the i -th interval by $\mathcal{C}_i$. The probability of observing $\mathcal{C}_i$ is then given by:

$$P(\mathcal{C}_i | \gamma, T) = \gamma^{|\mathcal{C}_i|} e^{-\gamma \tau_i i(i-1)} \quad (3)$$

This Poisson process implicitly defines a uniform distribution over the height of the edges and the donor- and receiver-lineages within interval i . Taking the product over all intervals gives the full density of the contact prior:

$$P(\mathcal{C} | \gamma, T) = \prod_{i=2}^n \gamma^{|\mathcal{C}_i|} e^{-\gamma \tau_i i(i-1)} \quad (4)$$

$$= \gamma^{|\mathcal{C}|} \prod_{i=2}^n e^{-\gamma \tau_i i(i-1)} \quad (5)$$

$$= \gamma^{|\mathcal{C}|} e^{-\gamma \sum_{i=2}^n \tau_i i(i-1)} \quad (6)$$

The sum in Equation 6 can be interpreted as the length of all possible pairs of lineages in the tree and we denote it by $L^{(2)} = \sum_{i=2}^n \tau_i i(i-1)$. This allows us to conveniently write the prior as

$$P(\mathcal{C} | \gamma, T) = \gamma^{|\mathcal{C}|} e^{-\gamma L^{(2)}} \quad (7)$$

and the number of contact edges follows the Poisson distribution

$$P(|\mathcal{C}| \mid \gamma, T) = \frac{(\gamma L^{(2)})^{|\mathcal{C}|} e^{-\gamma L^{(2)}}}{|\mathcal{C}|!} \quad (8)$$

The expected value of this Poisson distribution is $\Gamma = \gamma L^{(2)}$. In practice, we recommend to parameterise the prior in terms of the expected number of edges since it avoids a dependency between the number of contact edges and the time scale of the tree:

$$P(|\mathcal{C}| \mid \Gamma) = \frac{\Gamma^{|\mathcal{C}|} e^{-\Gamma}}{|\mathcal{C}|!} \quad (9)$$

Linear contact prior. As the lineages increase with every interval, the quadratic growth prior leads to many very recent contact edges and very few edges close to the root, which may be undesirable. We think the geographical structure of languages (or other types of taxa) leads to the plausible assumption that contact only increases linearly with the number of lineages. Let us assume that two languages need to be near another to be in contact. We suppose further that each language can only be near a certain constant number of other languages. In that case, each language's potential number of contacts is also bounded by that constant. Finally, a constant number of contacts per language implies that the total number of contact edges grows linearly with the number of languages.

The definition of the prior only changes slightly under these different assumptions: instead of $i(i+1)$ possible contacts in the i -th interval, the prior only allows for ai possible contacts for some constant a . The resulting prior is defined as:

$$P(\mathcal{C} \mid \gamma, T) = \prod_{i=2}^n \gamma^{|\mathcal{C}_i|} e^{-\gamma \tau_i a i} \quad (10)$$

$$= \gamma^{|\mathcal{C}|} e^{-\gamma a \sum_{i=2}^n \tau_i i} \quad (11)$$

In this case, the sum in the exponent can be interpreted as the tree length $L = \sum_{i=2}^n \tau_i i$. Again, we can write the expected number of contact edges as a Poisson distribution with mean $\Gamma = \gamma L$:

$$P(|\mathcal{C}| \mid \gamma, T) = \frac{(\gamma L)^{|\mathcal{C}|} e^{-\gamma L}}{|\mathcal{C}|!} = \frac{\Gamma^{|\mathcal{C}|} e^{-\Gamma}}{|\mathcal{C}|!} \quad (12)$$

In the case study, we use the linear contact prior, but both options are available in the `contactTrees` package.

S1.2 Likelihood and word trees

For a `contactTrees` analysis, we have to define the units of borrowing, data in blocks of sites that are assumed to be borrowed together. In the case of lexical phylogenies, we call these blocks words and denote the set of all words by $\mathcal{W}$. For each word $w \in \mathcal{W}$, we denote the corresponding data by X_w . The whole data set is defined by $X = \{X_w \mid w \in \mathcal{W}\}$. The `contactTrees` likelihood can be computed as a product of standard tree likelihoods over all words:

$$P(X \mid T, \mathcal{C}, Z, \theta_X) = \prod_{w \in \mathcal{W}} P(X_w \mid \tau_w, \theta_X), \quad (13)$$

where τ_w is the word tree for word w and θ_X are all parameters of the tree likelihood, including the substitution model parameters and branch rates.

The main article explains that the word trees are derived from the language tree T , the contact edges $\mathcal{C}$ and the borrowing indicators Z . Here we give detailed instruction for how this derivation works. The implementation can be found in <https://github.com/NicoNeureiter/contactTrees/blob/master/src/contacttrees/MarginalTree.java>.

The word tree of a word $w \in \mathcal{W}$ is a function of the language tree T , the contact edges $\mathcal{C}$ and the borrowing indicators $Z_{\cdot, w}$. The generative process describing the development of the word w is a simple modification of a standard phylogenetic model. At the root, the word is in a specific unknown state. Over time the state changes in a continuous-time Markov chain, i.e. it has a continuous – potentially varying – chance to switch to a different state. At each node in the tree, a branch splits into two. Both branches continue the Markov process independently. Additionally, w is affected by a set of contact edges $\mathcal{C}_w = \{c \in \mathcal{C} \mid Z_{c, w} = 1\}$. At each contact edge $(l_1, l_2, t) \in \mathcal{C}_w$, the state of w is copied from l_1

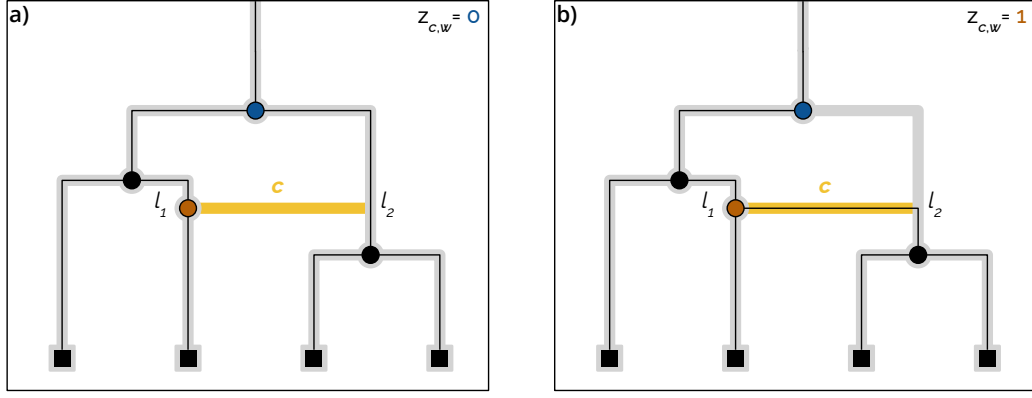


Figure S1: The language tree (bold grey line) and word tree for word w (thin black line) in an example with one contact edge $c = (l_1, l_2, t)$. The edge c defines two possible word trees: a) w is not borrowed at c , and the word tree follows the language tree. b) w is borrowed, and the word tree follows the contact edge c . The changes leading up to the edge on branch l_2 do not affect the data for this word and hence the branch is unobserved for w (grey line with black line missing).

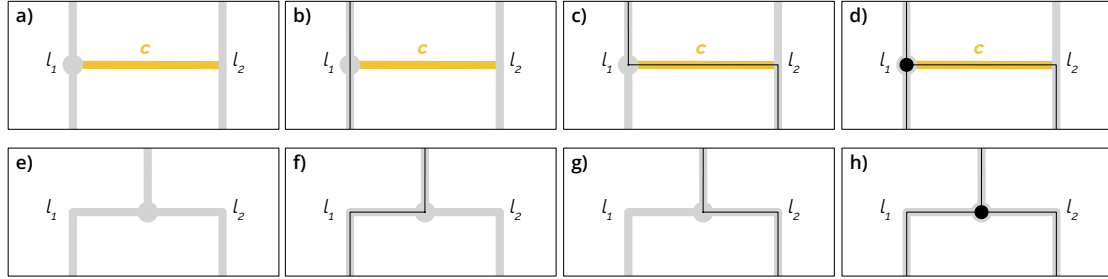


Figure S2: Contact edges and language tree splits both lead to a split in the word tree if the lineages below l_1 and l_2 are observed (d and h). If only the receiver, l_2 , is observed below a contact edge, the contact edge above the donor becomes observed instead (c). Otherwise, the word tree is neither affected by the contact edge nor the language tree split (a,b,e,f,g).

to l_2 at time t . This copied state overshadows everything that happened on the branch l_2 leading up to the contact event. The overshadowed branch above l_2 is replaced by a coalescence with the branch l_1 . Hence, we can still describe the history of a word as a tree, as illustrated in Figure S1.

To construct the word tree τ_w , we need to be aware of possible interactions between contact edges. Vaughan et al. (2017) present an efficient way to do this. The idea is to keep track of all active lineages in the word tree. Starting from the leaves of the tree we go backwards in time through the language tree splits and contact edges sorted by height. Contact edges that do not affect w (i.e. where $Z_{c,w} = 0$) can be entirely ignored for the construction of τ_w . The ones that do affect w (i.e. where $Z_{c,w} = 1$), are treated the same way as a split in the language tree. Both represent a possible split in the word tree. However, this split might be unobserved if at least one of the involved lineages is unobserved, requiring further checks.

Since we order events by height, we always start at the leaves. Each leaf in the language tree is also a leaf in the word tree, which is why they are added to the set of active lineages. At each potential split, i.e. a contact edge or a split in the language tree, we check whether both potential child nodes resulting from the split are part of the active lineages. If both are, the split is added to the word tree. If only one or none of the child nodes is active, there is no split in the word tree. Eventually, the last event that replaces two active lineages by one defines the root of the marginal tree. This procedure is described step-by-step in the pseudo-code in Algorithm 1.

S2 MCMC operators

We use MCMC sampling to draw samples from the posterior distribution. The `contactTrees` model includes additional parameters compared to standard phylogenetic models: The contact edges $\mathcal{C}$ and the

Algorithm 1 MarginalTree($T, \mathcal{C}, Z_{\cdot, w}$)

```

1:  $\mathcal{C}_w \leftarrow \{c \in \mathcal{C} \mid Z_{c, w} = 1\}$ 
2:  $\mathcal{E} \leftarrow \mathcal{C}_w \cup \text{nodes}(T)$ 
3:  $\mathcal{A} \leftarrow \{\}$  ▷ Keep track of active lineages in a set of (key, value) pairs.
4: ▷ Keys are nodes in  $T$ ; Values are nodes in  $\tau$ .
5:  $\tau \leftarrow \text{EmptyTree}()$ 
6: for  $e$  in  $\text{sortedByHeight}(\mathcal{E})$  do
7:   if  $e$  is a leaf node then
8:     Add new leaf node  $n$  to  $\tau$ 
9:      $\mathcal{A} \leftarrow \mathcal{A} \cup \{(e, n)\}$ 
10:    continue
11:   else if  $e$  is an internal node then
12:      $(l_1, l_2) \leftarrow \text{children}(e)$ 
13:      $t \leftarrow \text{height}(e)$ 
14:   else if  $e$  is a contact edge then
15:      $(l_1, l_2, t) \leftarrow e$ 
16:   if  $\exists n_1, n_2 \in \text{nodes}(\tau) : (l_1, n_1) \in \mathcal{A} \wedge (l_2, n_2) \in \mathcal{A}$  then
17:     ▷ Both lineages are active (Fig. S2 d and h)
18:     Add new node  $n$  with height  $t$  and children  $n_1, n_2$  to  $\tau$ .
19:      $\mathcal{A} \leftarrow \mathcal{A} \cup \{(e, n)\}$  ▷ Add new active lineage
20:      $\mathcal{A} \leftarrow \mathcal{A} \setminus \{(l_1, n_1), (l_2, n_2)\}$  ▷ Remove old active lineages
21:   else if  $\exists n_1 \in \text{nodes}(\tau) : (l_1, n_1) \in \mathcal{A}$  then
22:     ▷ Only lineage  $l_1$  is active (Fig. S2 b and f)
23:     if  $e$  is a node then
24:        $\mathcal{A} \leftarrow \mathcal{A} \cup \{(e, n_1)\} \setminus \{(l_1, n_1)\}$  ▷ Fig. S2 f
25:   else if  $\exists n_2 \in \text{nodes}(\tau) : (l_2, n_2) \in \mathcal{A}$  then
26:     ▷ Only lineage  $l_2$  is active (Fig. S2 c and g)
27:     if  $e$  is a node then
28:        $\mathcal{A} \leftarrow \mathcal{A} \cup \{(e, n_2)\} \setminus \{(l_2, n_2)\}$  ▷ Fig. S2 g
29:     else if  $e$  is a contact edge then
30:        $\mathcal{A} \leftarrow \mathcal{A} \cup \{(l_1, n_2)\} \setminus \{(l_2, n_2)\}$  ▷ Fig. S2 c
31: return  $\tau$ 

```

borrowing indicators Z . These parameters require custom MCMC operators, which are implemented in the `contactTrees` packages. Further parameters in the priors, like Γ or β , can be estimated too, but they do not require custom operators. Standard MCMC operators for real-valued parameters (like the `ScaleOperator` or the `RealRandomWalkOperator` available in BEAST 2) can be readily used in these cases. The added MCMC operators fall into three categories: 1) tree operators, 2) contact edge operators and 3) borrowing operators. We briefly explain each new operator below.

The MCMC algorithm samples from a target distribution $\pi(\theta)$ by iteratively proposing new states according to proposal distributions $q(\theta'|\theta)$. In our case, θ is the set of all model parameters, and the target distribution is the posterior distribution $P(\theta|X)$ of the `contactTrees` model (main article, Equation 5). The proposal distribution is chosen randomly from a set of transition operators in BEAST 2. Each proposal is accepted with probability $\alpha(\theta'|\theta)$ and rejected otherwise, ensuring that this iterative process converges to the target distribution. The acceptance probability is defined by:

$$\alpha(\theta'|\theta) = \min \left(1, \frac{\pi(\theta')}{\pi(\theta)} h(\theta'|\theta) \right) \quad (14)$$

Here, $h(\theta'|\theta)$ is the Hastings-Green factor (HGF) for the proposal distribution. The HGF compensates for potentially asymmetric proposal probabilities.

As a basis for this section, we introduce some additional formal notation. As defined in the main article, `contactTrees` seeks to find a language tree T and contact edges $\mathcal{C}$. The language tree $T = (V, E, \mathbf{t})$ is defined by nodes V , direct edges between the nodes E – where $\langle l_1, l_2 \rangle \in E$ implies that l_1 is the parent of l_2 in the tree – and node heights $\mathbf{t} = \{t_l | l \in V\}$.

S2.1 Borrowing operators

Borrowing operators pick a random contact edge $c \in \mathcal{C}$ and randomly change the corresponding borrowing events $Z_{c,w}$ for some words $w \in \mathcal{W}$. If $\mathcal{C}$ is empty, the operator is immediately rejected. Otherwise, there are two variants for re-sampling $Z_{c,w}$: 1) a naive random walk operator which samples each $Z_{c,w}$ from the prior or 2) a Gibbs operator which samples $Z_{c,w}$ from the conditional posterior distribution.

Naive borrowing operator. The simplest way to propose a new $Z_{c,w}$ is to sample it from the prior:

$$q(Z'_{c,w}|\theta) = P(Z'_{c,w}) = \beta^{Z'_{c,w}} (1 - \beta)^{1-Z'_{c,w}} \quad (15)$$

Since the operator is symmetric this way, the HGF is $h(Z'_{c,w}|Z_{c,w}) = 1$.

Gibbs borrowing operator. Due to the high dimensionality of Z , the naive borrowing operator is highly inefficient and will need many attempts to propose suitable borrowings. We introduce a Gibbs operator for this purpose, which samples $Z_{c,w}$ from the conditional posterior distribution:

$$q(Z'_{c,w}|\theta) = P(Z'_{c,w} | T, \mathcal{C}, X, Z_{\bar{c},w}, \theta_X) = \frac{P(X | T, \mathcal{C}, Z'_{c,w}, \theta_X)}{P(X | T, \mathcal{C}, Z'_{c,w}, \theta_X) + P(X | T, \mathcal{C}, 1 - Z'_{c,w}, \theta_X)} \quad (16)$$

Here $Z'_{c,w}$ is the new proposal for the borrowing indicator of word w at edge c , $Z_{\bar{c},w}$ is the old state of the borrowing indicators for all other edges $\bar{c} = \mathcal{C} \setminus \{c\}$, and $Z'_{\bar{c},w}$ is the combination of the old state for $\bar{c}$ and the proposed new state for c , i.e. $Z'_{c,w}$ and $Z_{\bar{c},w}$ together. The probabilities in the right-hand side of Equation 16 can be computed as a tree likelihood for the two different word trees resulting from $Z'_{c,w} = 0$ and $Z'_{c,w} = 1$ (cf. Equation 1). Since the proposal directly samples from the conditional posterior distribution, this operator will always be accepted. If we want to express this in an HGF, it would have the following form (and would cancel out with the posterior ratio in Equation 14):

$$h(Z'_{c,w}|Z_{c,w}) = \frac{P(X | T, \mathcal{C}, Z'_{c,w}, \theta_X)}{P(X | T, \mathcal{C}, Z_{c,w}, \theta_X)} \quad (17)$$

S2.2 Contact edge operators

Sampling the contact edges $\mathcal{C}$ requires custom operators to add and remove contact edges and, for sampling efficiency, to move contact edges locally. Many of these operators are inspired by Vaughan et al. (2017), but require adaptations due to differences between the Bacter and the `contactTrees` model.

Naive add/remove contact edge operator. This operator randomly adds or removes a contact edge with equal probability. To remove an edge, the operator simply chooses one $c \in \mathcal{C}$ uniformly at random and deletes it from $\mathcal{C}$. To add an edge, the operator draws a new contact edge from the contact prior. For a linear contact prior the starting point of the edge (the donor) is chosen uniformly at random along the language tree. Subsequently, the end point of the edge (the receiver) is chosen uniformly at random from all other lineages at the same height. For the quadratic contact prior, the starting and end point are sampled together from the volume of all possible synchronous starting and end points on the language tree. Finally, adding a new contact edge also adds a dimension to the borrowing indicators, which also needs to be sampled. For a new contact edge c^* , the naive add/remove contact edge operator samples the new indicators $Z_{c^*,w}$ from the prior as in the naive borrowing operator. Let the new set of contact edges be $\mathcal{C}^+ = \mathcal{C} \cup \{c^*\}$. Together this yields the following HGF:

$$h_{\text{add}}(\mathcal{C}^+, Z_{c^*,.} \mid \mathcal{C}) = \underbrace{\frac{1}{|\mathcal{C}^+|}}_{\text{removal probability}} \cdot \underbrace{L \cdot (i-1)}_{\text{inverse of edge attachment probability}} \cdot \underbrace{\frac{1}{\prod_{w \in \mathcal{W}} \beta^{Z_{c^*,w}} (1-\beta)^{1-Z_{c^*,w}}}}_{\text{inverse of borrowing probability}} \quad (18)$$

The first term is the removal probability, i.e. the probability of choosing the edge c^* in a removal step. The second term is the inverse of the probability density of placing the edge $c^* = (l_1, l_2, t)$ at its chosen location. L denotes the tree length (i.e. the possible choices for the donor l_1 and height t) and i is the number of lineages at height t (i.e. $(i-1)$ is the number of possible choices for l_2). The third term is the inverse of the borrowing probability, defined by the borrowing prior. Note that the edge attachment probability assumes a linear contact prior. For a quadratic contact prior, the second term would be $L^{(2)}$ (the length of all possible pairs of lineages in the tree, as defined in Section S1.1). The HGF of a removal step is the inverse of h_{add} :

$$h_{\text{remove}}(\mathcal{C} \mid \mathcal{C}^+, Z_{c^*,.}) = \frac{1}{h_{\text{add}}(\mathcal{C}^+, Z_{c^*,.} \mid \mathcal{C})} = \frac{|\mathcal{C}^+| \prod_{w \in \mathcal{W}} \beta^{Z_{c^*,w}} (1-\beta)^{1-Z_{c^*,w}}}{L \cdot (i-1)} \quad (19)$$

Semi-Gibbs add/remove contact edge operator. The naive add/remove contact edge operator is very often rejected due to the inefficient sampling of $Z_{c^*,.}$. In most cases, the proposed borrowings will not fit the contact edge. Hence, we also apply the idea of the Gibbs borrowing operator when adding/removing edges: we sample borrowings $Z_{c,.}$ of the new edge c from the conditional posterior distribution. The resulting Hastings-Green factor is:

$$h_{\text{add}}(\mathcal{C}^+, Z_{c^*,.} \mid \mathcal{C}) = \frac{L \cdot (i-1)}{|\mathcal{C}^+|} \prod_{w \in \mathcal{W}} \left(\frac{\text{P}(X_w \mid T, \mathcal{C}^+, Z_{.,w}, Z_{c^*,w}, \theta_X)}{\text{P}(X_w \mid T, \mathcal{C}, Z_{.,w}, \theta_X)} \frac{1}{\beta^{Z_{c^*,w}} (1-\beta)^{1-Z_{c^*,w}}} \right) \quad (20)$$

In contrast to the Gibbs borrowing operator, we can not directly accept this operator. Since the first step of adding an edge does not follow the conditional posterior distribution, the HGF and the posterior ratio do not cancel out in the acceptance probability. However, sampling the borrowing indicators from the posterior generates a sensible set of borrowings for the new edge and vastly increases the acceptance probability.

Contact edge slide. The *contact edge slide* operator locally explores the space of possible contact edges by moving an existing edge up or down in the language tree, possibly beyond internal nodes. The operator starts by picking a random contact edge $c = (l_1, l_2, t) \in \mathcal{C}$. If $\mathcal{C}$ is empty, the operator is rejected immediately. Then it draws an offset $\delta_t \in [-s \cdot h, s \cdot h]$ uniformly at random, where h is the tree height and s is a tuning parameter of the operator. This offset is added to the node height: $t' = t + \delta_t$. The operator is rejected if $t' \notin [0, h]$. Changing the height of a contact edge might also require a change in the lineages. If the new height is above the parent of l_1 , the new donor will be the ancestor of l_1 at height t' . If the new height is below l_1 , the new donor is chosen uniformly at random from all descendants of l_1 at height t' . The same holds for the receiver l_2 . Let $\kappa(l, t)$ be the number of lineages at height t that are direct ancestors or descendants of node l . The HGF of the contact edge slide is:

$$h((l'_1, l'_2, t') \mid (l_1, l_2, t)) = \frac{\kappa(l_1, t') \kappa(l_2, t')}{\kappa(l'_1, t) \kappa(l'_2, t)} \quad (21)$$

Contact edge flip. This is another operator for local exploration of the search space that simply proposes to flip the direction of a contact edge. That is, it picks a random edge $(l_1, l_2, t) \in \mathcal{C}$ and proposes to replace it by the edge (l_2, l_1, t) . Since this operator is symmetric, the HGF is always 1.

Contact edge hop. This operator chooses a random edge $(l_1, l_2, t) \in \mathcal{C}$ and proposes to change one of its attachment points l_1 or l_2 by a different lineage. The proposed new edge will either have a new donor, i.e. (l'_1, l_2, t) , or new receiver, i.e. (l_1, l'_2, t) . In the simplest form, this operator can choose any lineage at height t as the new attachment point for l_1 or l_2 . The resulting proposal is symmetric and has an HGF of 1.

However, randomly proposing a new lineage can drastically change the role of a contact edge, and we empirically found this operator to give very poor performance, i.e. it was mostly rejected. Hence, we introduce different variants of the contact edge hop:

- **Local contact edge hop:** To restrict the operator to sample more locally in the parameter space we limit the distance between the original attachment point l (either l_1 or l_2) and the proposed attachment point l' . Let $d(u, v) : V \times V \rightarrow \mathbb{N}$ be a function that counts the nodes on the path from u to v in the language tree. Further, let $d_{\max}$ be a tuning parameter that defines the maximum distance between the old and the new attachment point. The local contact edge hop operator chooses the new attachment points from the candidate set $c(l, d_{\max}) = \{l' | d(l, l') \leq d_{\max}\}$. Since the number of candidate points is not always symmetric – $|c(l, d_{\max})|$ is not generally the same as $|c(l', d_{\max})|$ – this changes the HGF to:

$$h((l'_1, l_2, t) | (l_1, l_2, t)) = \frac{c(l_1, d_{\max})}{c(l'_1, d_{\max})} \quad (22)$$

- **Contact edge donor hop:** Empirically, there seems to be more uncertainty about the donor of a contact edge than the receiver, at least in our case study. This also matches our intuition since the data can clearly show which languages obtained a set of loanwords, but the source of the loanword will always be more uncertain. For example, the contact edge from Norman French to English very clearly has English as a receiver since other West-Germanic languages do not share the same loanwords. However, other Gallo-Romance languages could plausibly serve as a donor since they share many of the cognates borrowed into English. As a result, most contact edge hops that attempt to change the receiver are rejected, while changes to the donor are more likely accepted. For this reason, we introduce a version of the contact edge hop that only changes the donor language and achieves a higher acceptance rate. The HGF remains unchanged.
- **Contact edge hop + Gibbs sample borrowings:** Since changing the attachment point can significantly impact the meaning of a contact edge, the borrowings' context for this edge changes too. Hence, it makes sense to propose new borrowings $Z_{c..}$. We introduce a modified contact edge hop, similar to the Semi-Gibbs add/remove contact edge operator. In this version, $Z_{c..}$ is drawn from the conditional posterior distribution after c is changed according in the *hop* step. Similar to Equation 20, the HGF changes by a factor of $\frac{P(X | T, C', Z, \theta_X) P(Z | C', \beta)}{P(X | T, C', Z', \theta_X) P(Z' | C', \beta)}$.

These modifications can be combined, and different versions can be included in the same MCMC run.

S2.3 Tree Operators

BEAST 2 offers many operators to explore the space of phylogenetic trees. The ideas of these operators are also applicable to the language tree in a *contactTrees* analysis. However, whenever a tree operator moves a branch in the language tree, this might change the attachment points of a contact edge in different ways. Any tree operator used for a *contactTrees* analysis needs to take care of these interactions and adapt the HGF accordingly. The *contactTrees* packages include a range of these adapted tree operators, which we will describe in this section. These operators are very similar to those developed in the *Bacter* package (Vaughan et al., 2017), but the HGFs had to be adapted to the different priors and the strictly horizontal contact edges in the *contactTrees* model. In the following paragraphs, we briefly recap these operators and state the adapted HGFs.

Expand/collapse. We construct other tree operators based on two basic operations: *expand* and *collapse*. In both operations, a subtree with root $u \in V$ is moved from one branch to another. Let $u_p \in V$ be the parent of u , u_s its sibling and u_g its grandparent (Figure S3). The expand/collapse move chooses a target branch. The node $v \in V$ and its parent, $v_p \in V$, define the upper and lower end of the target branch. After reattaching the subtree at the target branch below v , the edges E change in the following way:

$$E' = E \setminus \{\langle u_g, u_p \rangle, \langle u_p, u_s \rangle, \langle v_p, v \rangle\} \cup \{\langle u_g, u_s \rangle, \langle v_p, u_p \rangle, \langle u_p, v \rangle\} \quad (23)$$

The new attachment point is generally not at the same height as the old one. If it is higher, i.e. $t'_{u_p} > t_{u_p}$, the operation is called *expand*, if it is lower, i.e. $t'_{u_p} < t_{u_p}$, it is called *collapse*.

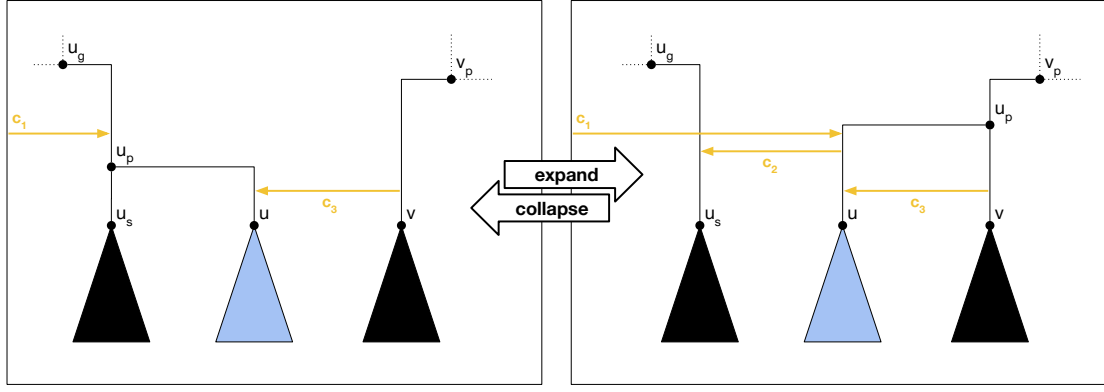


Figure S3: The expand and collapse operators move the attachment of a subtree u from a branch above u_s to a branch above v . In an expand move, the new attachment point is above the old one ($t'_{u_p} > t_{u_p}$), while in a collapse move, it is below the old one ($t'_{u_p} < t_{u_p}$).

The expand and collapse operators necessarily affect some contact edges. In a collapse move, any contact edge $(l_1, l_2, t) \in \mathcal{C}$ attached to u (i.e. $l_1 = u$ or $l_2 = u$) with height $t > t'_{u_p}$ has to be reattached since the lineage above u ends at t'_{u_p} . Instead, the attachment point is simply moved to the ancestor of u at height t (the edge c_1 in Figure S3). In case this move leads to the donor and receiver being equivalent, c is removed (the edge c_2 in Figure S3). For each collapse move, there must be a converse expand move to revert the changes of the collapse operator. This reversibility is necessary for an MCMC operator to be accepted. For this reason, the reattachment of the contact edges needs to be reversible too. Hence, any contact edge attached to an ancestor of u at height $t \in [t_{u_p}, t'_{u_p}]$ is reattached to u itself (the edge c_1 in Figure S3). We denote the set of edges that could be moved in this way by:

$$\mathcal{C}_{\text{movable}} = \left\{ (l_1, l_2, t) \in \mathcal{C} \mid t \in [t_{u_p}, t'_{u_p}] \wedge (l_1 \in \text{ancestors}(u) \vee l_2 \in \text{ancestors}(u)) \right\}$$

To ensure reversibility of removing degenerate edges in the collapse move, we randomly sample new contact edges $\mathcal{C}_{\text{added}}$ between u and v at a random height $t \in [t_{u_p}, t'_{u_p}]$ from the contact prior in every expand move. The resulting contribution to the HGF of an expand move is:

$$h(T', \mathcal{C}' | T, \mathcal{C}) = \left(2^{-|\mathcal{C}_{\text{movable}}|} \gamma^{|\mathcal{C}_{\text{added}}|} \exp(-2 \Delta \gamma) \prod_{c \in \mathcal{C}_{\text{added}}} P(Z_c | \beta) \right)^{-1} \quad (24)$$

Where $\Delta = t'_{u_p} - t_{u_p}$ is the time difference between the old and the new attachment point. For a collapse move, $\mathcal{C}_{\text{movable}}$ are all edges that are attached to an ancestor u at height $t \in [t'_{u_p}, t_{u_p}]$ after the collapse move, and $\mathcal{C}_{\text{removed}}$ are all edges that are removed during the collapse move. The corresponding HGF factor contribution for a collapse move is:

$$h(T', \mathcal{C}' | T, \mathcal{C}) = 2^{-|\mathcal{C}_{\text{movable}}|} \gamma^{|\mathcal{C}_{\text{removed}}|} \exp(-2 \Delta \gamma) \prod_{c \in \mathcal{C}_{\text{removed}}} P(Z_c | \beta) \quad (25)$$

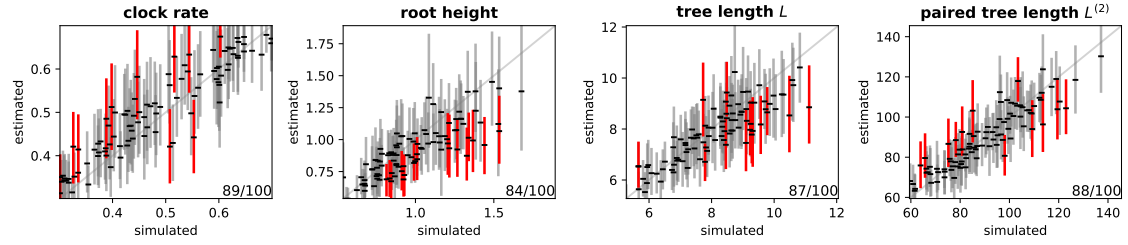
The expand and collapse moves can be combined to form different tree operators. With the adapted expand and collapse move, these tree operators can be constructed in the same way as described in (Vaughan et al., 2017). We implemented this for the *uniform operator*, the *subtree exchange operator* and the *Wilson-Balding operator*.

S3 Simulation study

S3.1 Coverage of other model parameters

Section 3 of the main article shows the coverage of the 95% credible intervals for the tree height and the clock rate in the simulation study. Of course, the tested models also have other parameters that can be compared. Figure S4 shows all sampled parameters for the noCT and CT models. The parameters for the CT model all show a coverage between 91 and 99, indicating that the inferred posterior distributions are well-calibrated.

a)



b)

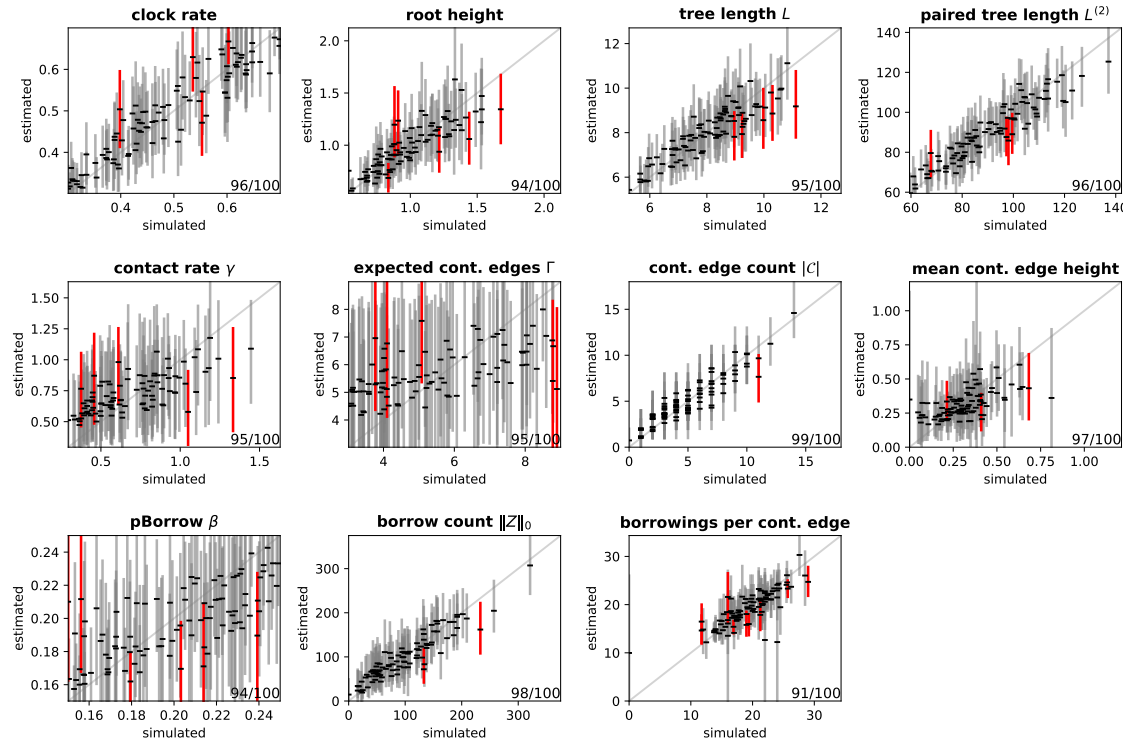


Figure S4: All logged parameters for a) the noCT and b) the CT model. Each plot shows simulated (x-axis) and reconstructed (y-axis) values for one parameter. The bars show the 95% credible interval of the reconstructed values, with grey indicating that the simulated value is within the interval and red that it is outside.

S3.2 Tree height bias and calibration points

In a phylogenetic study with real-world data, the clock rate of a phylogenetic model is usually calibrated through ancient languages or priors on the height of internal nodes. Section S4.3 explains how we do this in the case study. For the simulation study, we wanted to be close to this setting and, thus, calibrated the clock through strong priors on the height of three internal nodes. Since internal nodes are otherwise not well defined a priori, we put monophyly constraints on three groups of taxa: {1, 2, 3}, {4, 5, 6} and {7, 8}. The taxa were simply named by numbers 1 through 20. For each of these monophyletic clades, we define a normal distribution with $\sigma = 0.02$ as a prior on the height of the most recent common ancestor. The prior mean for the three clades is 0.35, 0.4 and 0.2, respectively.

The clock rate follows a uniform prior in the interval [0.3, 0.7]. Within this interval, the clock rate is estimated based on the data and the internal node priors. This setting is close to an actual phylogenetic case study, but it makes it more difficult to predict the effect of undetected contact on the reconstruction.

Let us consider the effect of contact on a phylogenetic analysis with a strict tree model. Contact can affect the number of changes along the tree in different ways. Borrowings between closely related languages, particularly those in neighbouring branches, makes the involved languages more similar and the phylogenetic pattern more stable. Borrowings from distantly related languages introduce traits that are untypical for the receiver clade. If the clock rate was fixed, the former would increase the inferred tree height, while the latter would decrease it. However, in actual studies, the clock rate is not fixed but is inferred from extinct languages and internal node calibrations (see Section S4.3). Consequently, the inferred clock rate depends on the number of changes between sampled languages at different times and calibration nodes. If these changes are increased/reduced by contact, the clock rate will be increased/reduced too. Finally, the tree height depends on the number of changes from the root to the leaves and calibration points. Hence, we do not have a clear theoretical expectation regarding the effect of undetected contact on the inferred root height. In the simulation study, we find a bias towards lower root heights, but this depends on the distribution of contact edges – whether there is more contact between closely or distantly related languages – and the distribution of calibration nodes relative to the contact edges – whether branches below or above calibration points are more affected by contact.

S4 Case study: Celtic, Germanic, Romance

S4.1 Data

The case study is based on the IELex data set, which contains form-meaning traits (introduced in the main article, Section 2.4.1) for 1419 forms for 206 meaning classes in 39 Celtic, Germanic and Romance languages. We initially took the data set from the supplementary information of Chang et al. (2015) and adapted the data in three ways:

- We corrected clear coding errors.
- We added missing loanword annotations.
- We collected and added data for Medieval Latin.

The adapted data set in CSV format and scripts to parse the data and generate BEAST 2 XML files are available on <https://github.com/NicoNeureiter/contactTrees-IndoEuropean>.

Latin is a special case among Indo-European languages. It was widely spoken in ancient Rome and evolved into today's Romance languages. Despite developing into many descendant languages, Latin itself was still used in science, medicine and the Catholic Church. Over time, this preserved Latin still changed, and these changes are documented in written sources. Even though this later form of Latin would not be considered a living language anymore – due to the lack of native speakers – it still had a significant impact on other languages. Many documented loanwords in Indo-European languages were borrowed from Latin long after antiquity, in particular after our sampling date of 2050 YBP. We also included later stages of Latin in our data set, to give the *contactTrees* model a chance to reconstruct such borrowings. We collected additional form-meaning traits for medieval Latin, which is notably different from ancient Latin. We estimate that the collected traits reflect the state of Latin about 1000 YBP. Even after that, Latin continued to be used, but to our knowledge, without significant changes in the core vocabulary. Hence, we include a duplicate of medieval Latin as a contemporary leaf node to allow later borrowings from Latin into other languages. We include *Latin_M* (1000 YBP) and *Latin_preserved* (0 YBP) for the medieval and contemporary taxa, respectively. To ensure that the three Latin nodes form a chain, we introduce two monophyletic constraints for {*Latin_M*, *Latin_preserved*} and {*Latin*, *Latin_M*, *Latin_preserved*} with a strong prior on the internal node height (Table S1).

Including two taxa with the same data 1000 years apart would drastically affect the estimates of the branch rate model. Since there are no changes between `Latin_M` and `Latin_preserved`, the clock rate would be pulled down, and the standard deviation of the branch rate distribution would be pulled up. We modify the relaxed clock specifically for this purpose: the `FreezableClock` makes it possible to specify *frozen taxa*. On the branch above a frozen taxon, the branch rate is fixed to 0. By specifying `Latin_preserved` as a frozen taxon, we ensure that it does not affect the parameters of the branch rate model.

S4.2 Model configuration

This section describes the model set-up for the case study in more detail, including the choice of priors, the substitution model, the branch rate model, the model for site rate variation and the time calibrations.

Tree prior. As a tree prior, we use a birth-death skyline model (Stadler et al., 2013) with the following configuration:

The net diversification rate is estimated with a log-normal prior distribution. The mean of the log-normal distribution is set to $\frac{\log(100)}{10000} = 0.00046$, based on the idea that over 10000 years a language family could diversify roughly into 100 surviving languages. The standard deviation is set to $\log(10) = 2.3026$, i.e. covering one order of magnitude up or down.

The relative death rate (turnover) is estimated with a uniform prior between 0 and 1. We know that some languages go extinct ($\text{turnover} > 0$), and we assume that languages historically diversified quicker than they went extinct ($\text{turnover} < 1$). Apart from that, we do not have any strong prior information on the turnover rate, hence the uniform prior distribution.

The sampling proportion is estimated in two intervals: Before and after 2500 YBP, based on the idea that historical records are more nearly complete for more recent languages, i.e. we expect a higher sampling rate after 2500 YBP.

The sampling probability of extant languages (ρ) is set to a constant 0.5.

Contact prior. We use a *linear growth* contact prior (see Section S1.1) with the expected number of contact edges set to $\Gamma = 0.25$. This regularising prior does not reflect our expectation regarding the number of contact edges but is intentionally lower to reduce the number of edges in the posterior results. While a strict tree model is naturally constrained to match tree-like patterns in the data, adding edges between branches allows the model to fit very general distributions. The regularising prior prevents overfitting.

Borrowing prior. The borrowing prior puts a constant probability β on each word to be borrowed at each contact edge. In theory, β could be estimated, which we demonstrated in the simulation study. Still, to keep the model simple and reduce the mixing time, we kept β fixed at 0.1 in this study.

Substitution model. The substitution model describes the continuous change of each trait over time. Different models have been proposed for lexical data. For etymon traits, a Dollo model seems a sensible choice. Since a common etymological root defines cognates, they are per definition almost never introduced twice in history. In a strict tree model, borrowing poses an exception to this since the word is introduced in the receiver language when it was already present in the donor language, violating the Dollo assumption. The `contactTrees` model builds exactly on this property to find deviations from the tree-like history of a word and would, most likely, make a Dollo model a good fit for etymon traits.

For root meaning traits, there are other causes for multiple innovations of the same word. Since traits are defined in meaning classes, an innovation is not necessarily a new word form but can also be a semantic change. These changes of meaning – or of the most common word used for a meaning – are much more volatile than etymological change. We also discuss these parallel innovations in the main article in Sections 4.2 and 4.3. Hence, a Dollo model seems to pose excessively strict assumptions for root meaning traits. Instead, we use a binary covarion model (Tuffley and Steel, 1998), an extension of a binary CTMC model with a continuous rate for each trait to change from 0 to 1 and 1 to 0. In a CTMC model, the rates are defined by the *clock rate* and the *state frequencies*, determining what fraction of traits would be present or absent after an infinite time of change. The covarion model additionally introduces a hot and a cold state. In the cold state, the rate of change is reduced by a factor of $\alpha \in [0, 1]$,

representing a time where a trait is particularly stable. Each trait changes between the hot and the cold state at a rate of s . The covarion model gives additional flexibility to the substitution process and has been found to fit linguistic data in previous phylogenetic studies (Bouckaert et al., 2012, 2018). The covarion parameters – α and s – are unknown and need to be sampled in the analysis. We use a uniform prior between 0 and 1 for α and an exponential prior with mean 0.2 for s .

Branch rate model. We use a log-normal relaxed clock model for branch rate variation. Each branch has a specific branch rate modelled as independent samples from a log-normal distribution. The log-normal is defined by a mean, the *clock rate*, and standard-deviation in log-space, σ . The clock rate is estimated with a uniform distribution between 0 and 1 as a prior. One can also estimate the standard deviation, but in the main analysis we fixed it at $\sigma = 0.5$, which ensures the comparability of the tree height in the different models. When σ is not fixed, the CT model estimates a higher standard deviation than the noCT model, which allows it to adapt the length of the first branches from the root to the proto-languages of Celtic, Germanic and Romance to fit the tree prior. In the case of Celtic-Germanic-Romance, these first branches seem to be much longer than the later ones, with only three splits before 2500 YBP, in contrast to 35 splits after that, in a time span of similar length. The estimated birth rate adapts to the recent, more rapid diversification, pulling down the length of the first branches. In a strict clock model, the differences between the different branches overshadow the prior. In a relaxed clock model, the tree can adapt to the tree prior by estimating higher rates and shorter lengths on the first branches. The variance of the branch rate distribution is an important determinant of this effect: larger values of σ allow for larger rate deviations in the first branches. Hence, having the same value for σ is crucial to ensure that the estimated root height is comparable across different settings in the CT and noCT model.

Site rate variation. It is common practice in phylogenetics to allow for site rate variation, i.e. different parts of the data are modelled to have different evolutionary rates. Site rate variation has an empirical foundation for lexical data where observations show that different words have evolved at different rates (Pagel et al., 2007). A standard model for site rate variation is *gamma rate heterogeneity*, where the rates of each site are modelled as independent samples from a gamma distribution (Yang, 1994). Using a discrete approximation, one can integrate out the site specific rates, avoiding too many additional parameters in the model. However, the approximation adds additional cost to the evaluation of the likelihood. If four rates approximate the gamma distribution, we need to evaluate the tree likelihood four times. Furthermore, the independence assumption might not be appropriate for root-meaning traits since different forms for a meaning likely evolve at similar rates. We introduce three rate categories – *slow*, *medium* and *fast* – and assign them to different meaning classes. The rate of a meaning class depends on the number of forms for that meaning in our data set: meanings with 1 to 5 forms are in the *slow* category, meanings with 6 to 9 forms are in the *medium* category and meanings with more than 9 forms are in the *fast* category. A meaning that is more susceptible to change is expected to develop more forms. Hence, the number of forms seems like a reasonable proxy for the assignment. We estimate a rate parameter for each category independently with the same prior distribution. The posterior distributions of the three rates confirms our initial expectation that the rate increases with the number of forms per meaning (Figure S5).

S4.3 Calibrations

To allow chronological interpretations, we want to calibrate the phylogenetic clock model. We use ancient languages and priors on the height of internal nodes to estimate the clock rate in changes per year and the tree height in years. Calibration deserves some discussion, as it was done in varying ways in previous phylogenetic studies. In most studies so far (Bouckaert et al., 2012; Gray and Atkinson, 2003), ancient languages were included as leaf nodes with a non-zero height, given by their present age. Modelling ancient languages as leaf nodes implies that they are related to the current language but are not their direct ancestor. Chang et al. (2015) argue that some of the sampled ancient languages in their data set – e.g. Old Irish, Old High German and Latin – are sampled ancestors of the current languages rather than separate leaf nodes. They enforce this by putting a very narrow prior on the length of the branches leading to these languages, which pulls the length towards zero. A branch of length zero would yield the same likelihood as a sampled ancestor. However, this hard constraint of forcing ancient languages to be sampled ancestors is debatable, as it is difficult to ensure that the sampled ancient language reflects the true ancestor of the current languages. Instead, the sample could reflect a dialect or sociolect which differs from the form that became the basis of the current languages. At the same time, these languages are not random extinct non-ancestral languages, either. Old Irish, Old English and Latin are

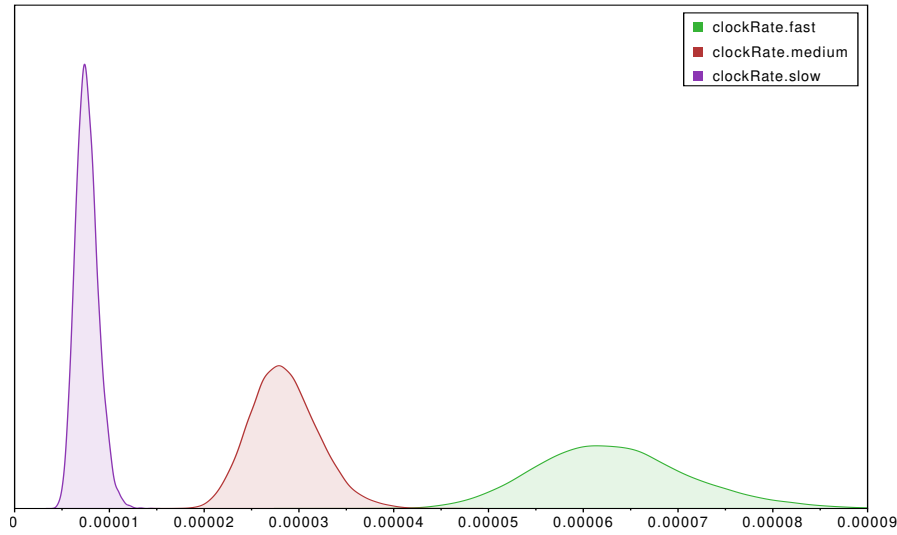


Figure S5: Posterior distribution of the three rate categories *slow*, *medium* and *fast*.

Internal node calibrations

Brythonic	Normal(mean=1550, sigma=100)
French Iberian	Normal(mean=1400, sigma=100)
West Germanic	Normal(mean=1550, sigma=100)
North West Germanic	Normal(mean=1875, sigma=150)
Celtic	Uniform(lower=1700, upper=10000)

Sampled almost-ancestors

Old Irish	Normal(mean=1250, sigma=100)
Old Norse	Normal(mean=775, sigma=100)
Old High German	Normal(mean=1050, sigma=100)
English	Normal(mean=1000, sigma=100)
Latin	Normal(mean=2050, sigma=150)

Chain of Latin variants

Latin + Latin_M + Latin_preserved	Normal(mean=2050, sigma=10)
Latin_M + Latin_preserved	Normal(mean=1000, sigma=10)

Table S1: Priors on tree nodes for: internal node calibrations, sampled almost-ancestors and custom priors to enforce a chain of different stages of Latin (ancient, medieval and “preserved” or at present).

coded to with the intention to include sampled ancestors of Goidelic, English and Romance, respectively. Even if the data reflects variants of these ancestors, we would expect them to be close relatives of the actual ancestor languages. We call such languages *sampled almost-ancestors* and define a prior on their branch length, relaxing the sampled ancestors model in (Chang et al., 2015). Specifically, the prior on the branch length above the sampled almost-ancestors is a half-normal distribution with $\sigma = 100$ years. The sampled almost-ancestors in our analysis are Old Irish, Old Norse, Old High German and Old English. For Latin, we use a slightly wider prior with $\sigma = 150$ years, and we include more recent versions of Latin (Section S4.1). In the future, the question of direct ancestry between sampled languages should be addressed using explicit models of both possibilities, as offered by the Sample Ancestors package (Gavryushkina et al., 2014) in BEAST 2. As the two packages are currently not compatible, we could not use this option for our case study.

The second type of calibrations are priors on internal nodes, i.e. splits in the language tree. We adopt these priors from Bouckaert et al. (2012). Table S1 lists all priors and Table S2 lists all sampling dates of ancient languages used for calibration.

S4.4 Results: topology

contactTrees can jointly infer phylogenetic trees and contact edges in one reconstruction. In the main article, we use a fixed tree topology. We focus on estimating node heights, contact edges and borrowings, for which we expect the most significant effects of the model. Here, we also report our results when

Cornish	300.0
Old Irish	1250.0
Old Norse	775.0
Old English	1000.0
Old High German	1050.0
Gothic	1650.0
Latin M	1000.0
Latin	2050.0

Table S2: Sampling dates of ancient languages in years before present.

reconstructing the topology. The following two pages show two densitree plots (Bouckaert and Heled, 2014) of the reconstructed language tree for the noCT model (Figure S6) and the CT model (Figure S7).

In what follows, we use the Chang et al. (2015) topology as a *reference tree*. Both inferred topologies deviate from the reference in the first splits between Celtic, Germanic and Romance. They show Celtic branching off first, followed by a Germanic-Romance split with a posterior support of 0.58 for noCT and 0.65 for CT. In the reference tree, the order is different, insofar as Germanic branches off first, followed by the Romance-Celtic split. However, the order of splits between the first three clades is highly uncertain in all models. Other differences are further down in the tree and only apply to one of the two models:

- The CT model places Catalan among the Gallo-Romance languages rather than Ibero-Romance.
- The noCT model places Rumanian in a clade with Sardinian, rather than a separate clade, branching from the Romance back-bone.
- The order of splits in the Frisian-Flemish-Dutch-Afrikaans clade is different from the reference tree in both reconstructions.
- The order of splits in the Norwegian-Icelandic-Faroese clade is different in the CT tree.

S4.5 Results: matches and mismatches compared to the annotated loanwords

This section compares the reconstructed borrowings with the loanwords listed in the IELex data set. Specifically, we list all potential false positives – reconstructed borrowings with a posterior support ≥ 0.8 not listed as a loanword in IELex (Figures S8 and S9) – and false negatives – reconstructed borrowings with a posterior support ≤ 0.2 listed as a loanword in IELex (Figure S10). Then, we summarise most common patterns and reasons for mismatches.

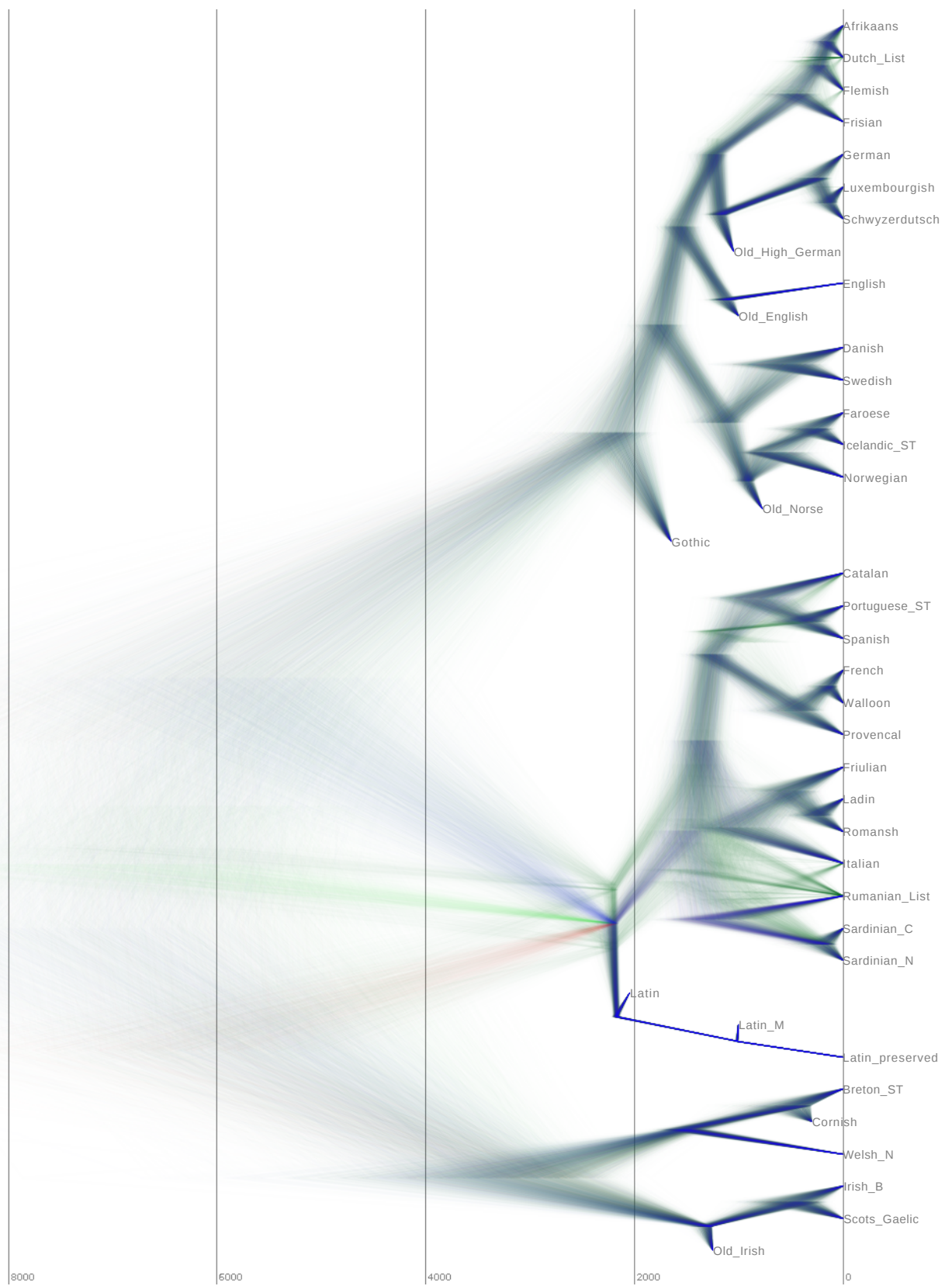


Figure S6: A densitree plot of the reconstructed Celtic-Germanic-Romance phylogeny under the noCT model.

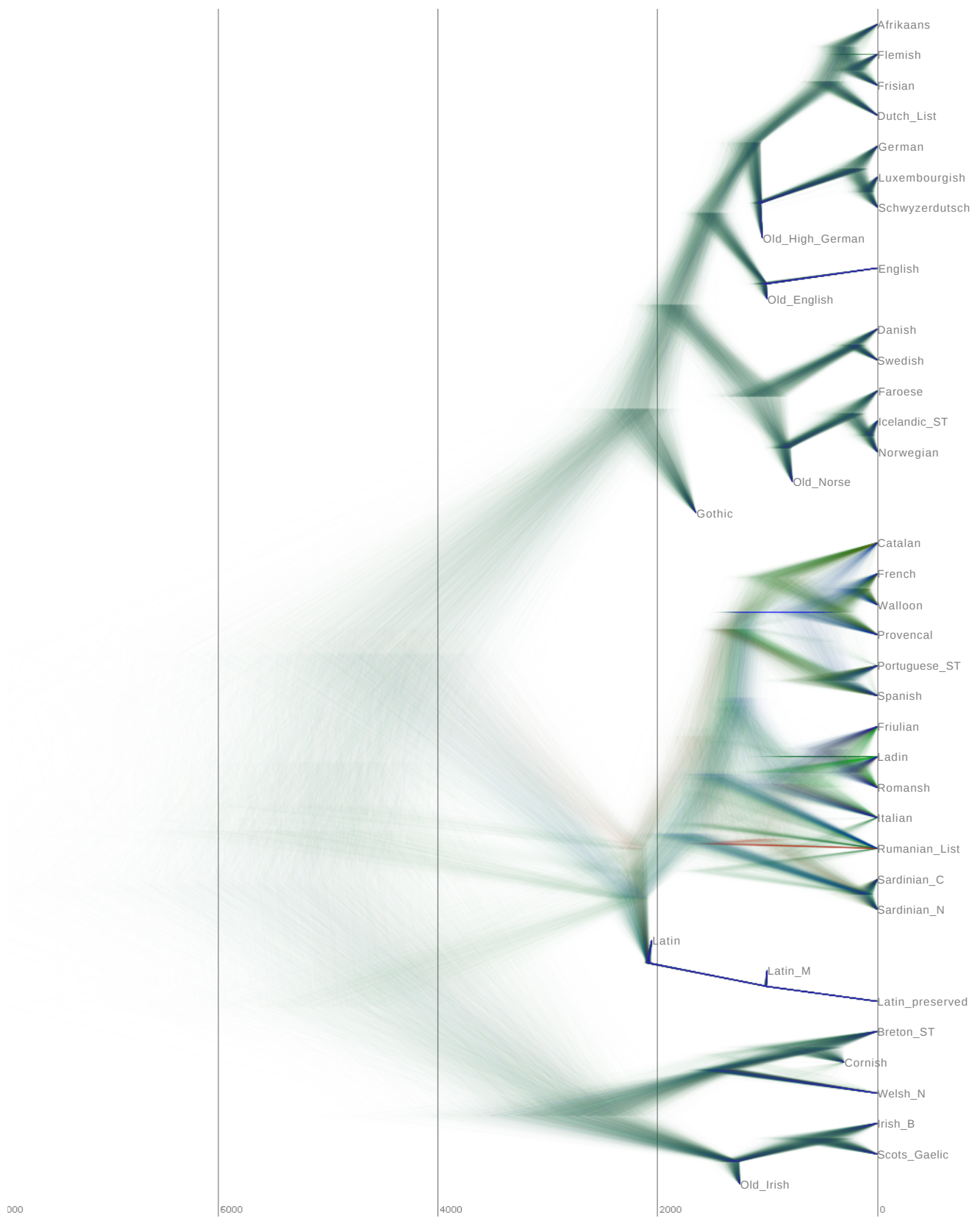


Figure S7: A densitree plot of the reconstructed Celtic-Germanic-Romance phylogeny under the CT model.

Afrikaans		Faroese	
0.93: REG (rightside)		0.99: BUKUR (belly)	
0.90: ? (round)		0.97: nø:s (nose)	
0.87: HOU (hold)		0.95: 'həusfrəu 'həusfrəu KONA (wife)	
Breton_ST		0.91: 'mɔani (moon)	
0.86: mām: (mother)		0.91: bain 'lɛg:ur (leg)	
0.82: 'askɔrn (bone)		0.91: VISS ʊm: (if)	
0.81: dəj, dəjs (day)		0.86: stɔv DUST (dust)	
		0.81: 'meavur 'hys.bændi (husband)	
Catalan		Flemish	
0.98: 'marə (mother)		0.93: REGTSCH (rightside)	
0.83: grə (seed)		0.90: ? (round)	
0.83: 'parə (father)		0.88: RIVIER (river)	
		0.87: VASTHOUDEN GRYPEN (hold)	
Cornish		French	
0.86: mam (mother)		0.97: mɛR (mother)	
0.82: askorn (bone)			
0.81: dyth wyth geydh (day)		Frisian	
Danish		0.93: ? (rightside)	
1.00: 'vasqʰə (wash)		0.90: ? (round)	
1.00: 'gʰkəsə (scratch)		0.88: hɔ:də, hɔ:rə (hold)	
1.00: dʰkagʰə (pull)			
0.99: 'slɔŋə (snake)		Friulian	
0.99: bu (belly)		0.98: mari (mother)	
0.99: sdʰʌgʰ (stick)		0.93: stɔmi (breast)	
0.99: 'jɛ:jə (hunt)		0.87: pari (father)	
0.99: 'dʰkægʰə (squeeze)		0.86: suiā (wipe)	
0.99: dʰɔŋ (heavy)		0.86: scolā (flow)	
0.99: 'wʌðən (rotten)		0.84: femine (wife)	
0.98: 'sdʰoɐ (big)		0.84: vint aiar, arie (wind)	
0.98: 'dʰɛŋ gʰə (think)		0.84: bosc (woods)	
0.98: sɛð (seed)		0.80: cognossi savê (know)	
0.98: 'gʰɛmbə (fight)			
0.98: gʰlɔw (split)		German	
0.97: 'sdø:ðə (push)		1.00: rɛçt (rightside)	
0.97: huð (skin)		1.00: 'haltŋ (hold)	
0.96: 'sdʰegʰə (stab)		1.00: kla_in (small)	
0.95: ben (leg)		0.99: 'ra_ibŋ (rub)	
0.95: 'gʰo:nə (wife)		0.99: 've:niç (few)	
0.94: 'hɛ:lə (tail)		0.99: fra_u 'gatin (wife)	
0.94: sdʰɔw (dust)		0.98: 've:ən (blow)	
0.94: 'bjɛw (mountain)		0.98: 'a_usbrɛçŋ (vomit)	
0.94: bɛ'sgʰid (dirty)		0.98: hals 'nakŋ (neck)	
0.94: bʰɛð (wide)		0.97: 'fly:gl (wing)	
0.93: gladh (smooth)		0.96: glat (smooth)	
0.92: 'fʰægdʰə (fear)		0.96: 'e:əman 'e:əman 'gatə (husband)	
0.92: smal (narrow)		0.94: gə'dɛrm 'a_ingəva_idə (guts)	
0.90: 'dʰɛlə (count)		0.94: ftumpf (dull)	
0.90: 'fly:ðə (flow)		0.92: flɛçt (bad)	
0.88: oɐm (worm)		0.84: fra_u va_ip (woman)	
0.86: so (lake)		0.81: 'fto:sŋ (push)	
0.85: sgʰa:b (sharp)			
0.81: 'vənə sɔj (turn)		Italian	
0.81: sbʰi:sə (eat)		0.98: 'madre (mother)	
0.80: man (man)		0.95: galled'dʒare (float)	
Dutch_List		0.85: ro'tella dʒi'nɔkkjo (knee)	
0.93: rɛxts, 'rɛxtər (rightside)		0.85: 'sabbja 'rena (sand)	
0.87: 'haudə(n) (hold)		0.81: 'dʒusto (right)	
		0.80: per'sona (person)	
English		0.80: 'padre (father)	
0.99: rait (rightside)		0.80: dʒet'tare (throw)	
0.98: bləv (blow)			
0.95: bɛli (belly)		Latin_M	
0.91: fæt (fat)		0.87: calidus (warm)	
0.89: skrætʃ (scratch)		0.84: ficatum, fecatum (liver)	
0.89: hi: (he)			
0.88: tri: (tree)			
0.87: mæn (man)			
0.86: wɛt (wet)			
0.84: həuld (hold)			
0.80: bæk (back)			
0.80: pul (pull)			

Figure S8: False positives (part 1)

Luxembourgish	Sardinian_N
1.00: hien (he)	0.98: PESANTE (heavy)
1.00: halen (hold)	0.98: PROKE PROKE (because)
1.00: Kapp (head)	0.97: KURPIRE (hit)
1.00: riets (rightside)	0.89: ISTARE (stand)
1.00: Bësch (woods)	0.87: ISTUKKITHTHARE (stab)
1.00: riicht (straight)	
1.00: kleng (small)	Schwyzerdutsch
0.99: reiwen (rub)	1.00: rächts (rightside)
0.99: wéineg (few)	1.00: wit (wide)
0.99: dréinen (turn)	1.00: loufe (walk)
0.98: Fra (wife)	1.00: chlii (small)
0.97: Papp (father)	0.99: ribe (rub)
0.97: Flillek (wing)	0.99: weni (few)
0.96: glat (smooth)	0.99: trääie (turn)
0.94: stomp (dull)	0.98: frou (wife)
0.92: schlecht (bad)	0.97: flügu (wing)
0.90: schméissen (throw)	0.96: glatt (smooth)
0.89: Fra (woman)	0.94: schtumpf (dull)
0.87: deelen (split)	0.92: schlächt (bad)
0.85: Mamm (mother)	0.89: wärfe (throw)
0.82: drécken (push)	0.89: frou (woman)
	0.81: schtoosse (push)
Norwegian	Spanish
1.00: GLATT JEVN (smooth)	0.98: 'maðre (mother)
1.00: STIKKE (stab)	0.92: a'rena (sand)
1.00: TENKE (think)	0.88: 'laryo (long)
1.00: GA PA JAKT (hunt)	0.87: es'tar de pje (stand)
1.00: TREKKE (pull)	0.86: ro'ði'la (knee)
1.00: STOKK (stick)	0.86: afi'laðo a'yüðo (sharp)
1.00: SMAL (narrow)	0.84: 'boske (woods)
1.00: VASKE (wash)	0.83: 'paðre (father)
1.00: BEN (leg)	0.83: 'kwerno hasta (horn)
1.00: FORDI FORDI (because)	
1.00: LUKTE (smell)	Swedish
0.99: STOV (dust)	0.99: 'ja:ga (hunt)
0.99: SKARP (sharp)	0.99: gru:v (big)
0.99: VARM (warm)	0.98: 'røtøn (rotten)
0.99: KLEMME (squeeze)	0.98: bæ:k MAGE (belly)
0.99: ROTTEN (rotten)	0.98: 'tenka (think)
0.98: DERE (you)	0.98: stur svo:r (heavy)
0.98: nese (nose) (nose)	0.98: se:d (seed)
0.94: buk MAVE (belly)	0.97: 'cempa 'stri:da (fight)
0.92: DARLIG (bad)	0.95: 'splitra (split)
0.88: hud (skin)	0.95: 'stika GENOMBORRA (stab)
0.88: elv (river)	0.93: 'bærj (mountain)
0.84: jord (earth)	0.91: 'fa:rho:ga 'frækta (fear)
0.83: RETT (straight)	0.90: fæt svans (tail)
0.81: RYGG (back)	0.89: 'fju:ta 'stø:ta (push)
0.80: VEI (road)	0.84: støft (dust)
	0.84: fjø: (lake)
Old_English	0.82: be:n (leg)
0.87: hē (he)	
0.85: trēow (tree)	Walloon
0.83: wæt fūht (wet)	0.97: MERE (mother)
0.82: hielt (hold)	
Portuguese_ST	Welsh_N
0.98: mēi (mother)	0.93: hwi (they)
0.92: e'reje (sand)	0.91: CLUST (ear)
0.87: iʃ'tar ɛj pɛ porse di pɛ (stand)	0.91: LLAW (hand)
0.85: ɛgu'sadu ɛfi'adu (sharp)	0.89: YCHYDIG (few)
0.82: pai (father)	0.88: COES (leg)
0.81: ? (knee)	0.86: mam (mother)
0.81: 'lōyu COMPRIDO (long)	0.86: OER (cold)
	0.84: ADAIN, ADEN (wing)
Provençal	0.82: ASGWRN (bone)
0.97: 'majre (mother)	0.81: di:ð, di:ð (day)
0.80: 'pajre (father)	0.80: CYWIR (right)
Romansh	
0.93: sain pèz (breast)	
0.82: ba:p (father)	
0.80: enconuscher saver (know)	
0.80: mamma (mother)	

Figure S9: False positives (part 2)

Afrikaans	Latin_M
0.15: PLUIM (feather)	0.12: troppus (many)
0.14: PERSOON (person)	
Breton_ST	Luxembourgish
0.17: 'kōntā (count)	0.10: katzen (vomit)
0.09: 'tōŋ (dull)	
0.07: ja'seal (hunt)	Old_English
0.05: bu'zelen (guts)	0.01: munt (mountain)
0.05: 'büntā (push)	0.01: stræt (road)
0.03: 'pōner, 'pūner (heavy)	
0.02: 'pluɛn (feather)	Old_High_German
0.01: 'sō:ʒal (think)	0.01: kurz (short)
0.01: stri:s (narrow)	
Catalan	Old_Irish
0.18: əs'kenə (back)	0.01: anmandae (animal)
0.16: əni'mal (animal)	0.00: clúmh (feather)
Cornish	Portuguese_ST
0.17: akontya, comtya nivera (count)	0.17: 'brẽku (white)
0.07: chassya (hunt)	
0.06: clini (ice)	Provençal
0.06: sewya (sew)	0.19: TOUMBA (fall)
0.02: pluven (feather)	0.19: SALE, ALO (dirty)
0.02: bad (bad)	0.18: FLOUTA, FLOUQUEJA (float)
0.01: cewsel cows (say)	0.15: GARRO (leg)
0.01: poder podrethek (rotten)	0.11: ESQUINO (back)
0.01: forh (road)	0.09: BLANC (white)
	0.08: RUSCO (bark)
	0.06: GRATA (scratch)
Dutch_List	Rumanian_List
0.05: kort (rope)	0.03: a pluti (float)
English	0.01: zə'padə (snow)
0.14: skai (sky)	0.01: tociť (dull)
0.14: big (big)	0.01: nisip (sand)
0.14: hit (hit)	0.01: a omo'ri (kill)
0.11: dig (dig)	
0.05: d3:ti (dirty)	Schwyzerdutsch
	0.04: chopf (head)
Flemish	Scots_Gaelic
0.14: PERSOON (person)	0.07: cūnnt (count)
	0.02: flùr (flower)
French	0.01: pearsa (person)
0.20: sal (dirty)	0.01: pahta (stick)
0.20: tōbe (fall)	0.01: ainmhidh (animal)
0.19: flōte (float)	0.01: pàiste (child)
0.09: blā (white)	0.01: ròpa (rope)
0.09: grate (scratch)	0.01: speuran (sky)
	0.00: rathad (road)
Friulian	Swedish
0.08: sgrafâ (scratch)	0.14: SMUTSIG (dirty)
German	0.04: da:m (woman)
0.04: kɔpf (head)	0.04: 'spe:la (play)
0.02: 'o:tsɛa:n (sea)	
Irish_B	Walloon
0.01: PEARSA, -ANN, -ANNA (person)	0.19: TOUMER (fall)
0.01: 'bvatvə (stick)	0.19: FLOTER (float)
0.01: 'pva:et'i (child)	0.10: BLANC (white)
0.01: SPEIR (sky)	0.09: GRETER (scratch)
0.01: ko:rʏŋʏ (horn)	
Italian	Welsh_N
0.18: spak'kare (split)	0.19: ko:x (red)
0.14: man'dzare (eat)	0.08: 'plent.in (child)
0.12: grat'tare (scratch)	0.02: PLUEN (feather)
	0.01: PWDR, PYDREDIG (rotten)
	0.01: FFORDD (road)
Ladin	
0.05: SGRATTER, GRATTER (scratch)	
Latin	
0.00: semita (road)	

Figure S10: False negatives

S4.6 Results: contact edges

In the case study, we reconstructed 32 contact edges. We think different aspects of each contact edge are of interest, including the involved donor and receiver languages, the proposed date of the contact event, the borrowings reconstructed at the contact event and the presence/absence of cognates in the involved and related languages. We visualised all these aspects for each contact edge in a separate plot on the last 32 pages of this document. Each plot shows the language tree in grey and a yellow arrow for the contact edge, placed at the median height of all posterior samples. The teal scatter lines show the uncertainty about the donor, the orange lines the uncertainty about the receiver. The teal/orange numbers give the posterior support for each possible donor/receiver branch. The meaning classes (rows) most likely borrowed at this edge are listed below the tree in decreasing order of posterior support. In each language (columns), the forms coded for a meaning class are shown in colour-coded squares, with the corresponding lexemes written out below.

References

- R. R. Bouckaert and J. Heled. DensiTree 2: Seeing trees through the forest. *BioRxiv*, page 012401, 2014.
- R. R. Bouckaert, P. Lemey, M. Dunn, S. J. Greenhill, A. V. Alekseyenko, A. J. Drummond, R. D. Gray, M. A. Suchard, and Q. D. Atkinson. Mapping the origins and expansion of the Indo-European language family. *Science*, 337(6097):957–960, 2012. doi: 10.1126/science.1219669.
- R. R. Bouckaert, C. Bowern, and Q. D. Atkinson. The origin and expansion of Pama-Nyungan languages across Australia. *Nature Ecology & Evolution*, 2(4):741–749, 2018. doi: 10.1038/s41559-018-0489-3.
- W. Chang, D. Hall, C. Cathcart, and A. Garrett. Ancestry-constrained phylogenetic analysis supports the Indo-European steppe hypothesis. *Language*, 91(1):194–244, 2015. doi: 10.1353/lan.2015.0005.
- A. Gavryushkina, D. Welch, T. Stadler, and A. J. Drummond. Bayesian inference of sampled ancestor trees for epidemiology and fossil calibration. *PLoS Computational Biology*, 10(12):e1003919, 2014.
- R. D. Gray and Q. D. Atkinson. Language-tree divergence times support the Anatolian theory of Indo-European origin. *Nature*, 426(6965):435–439, 2003. doi: 10.1038/nature02029.
- M. Pagel, Q. D. Atkinson, and A. Meade. Frequency of word-use predicts rates of lexical evolution throughout Indo-European history. *Nature*, 449(7163):717–720, 2007.
- T. Stadler, D. Kühnert, S. Bonhoeffer, and A. J. Drummond. Birth-death skyline plot reveals temporal changes of epidemic spread in HIV and hepatitis C virus (HCV). *Proceedings of the National Academy of Sciences*, 110(1):228–233, 2013.
- C. Tuffley and M. Steel. Modeling the covarion hypothesis of nucleotide substitution. *Mathematical biosciences*, 147(1):63–91, 1998.
- T. G. Vaughan, D. Welch, A. J. Drummond, P. J. Biggs, T. George, and N. P. French. Inferring ancestral recombination graphs from bacterial genomic data. *Genetics*, 205(2):857–870, 2017.
- Z. Yang. Maximum likelihood phylogenetic estimation from DNA sequences with variable rates over sites: approximate methods. *Journal of Molecular evolution*, 39(3):306–314, 1994.

