

Ruoxuan Li*, Doriane Gras and Caterina Donati

Investigating locality constraints on *wh*-in situ in French: an experiment

<https://doi.org/10.1515/probus-2025-0002>

Received February 2, 2025; accepted March 10, 2025; published online April 8, 2025

Abstract: Most analyses of *wh*-questions posit the existence of a dependency between the scope position and the theta position of the *wh*-constituent, both in ex situ questions and in in situ questions, but they disagree on the nature of this dependency and in particular on whether it is the same or different in the two strategies. The aim of this article is to verify how things stand in French, which is a mixed language allowing both in situ and ex situ *wh*-phrases in normal information seeking questions. We tested with two acceptability judgment tasks and two sentence-picture matching tasks whether in situ questions in French display the same subject advantage that has been widely attested in ex situ A-bar dependencies across languages. We found as expected a subject advantage in French *wh*-ex situ questions but an object advantage in French *wh*-in situ questions. We explain this seemingly divergent pattern by a) assuming that the same kind of movement dependency is established in both strategies but b) that its landing site is different and hence interacts differently with the clitic left dislocation structure (CLLD) we introduced in our experimental items. In both strategies, a processing advantage is observed when the *wh*-dependency and the clitic left dislocation structure (CLLD) dependency are nested, while crossing dependencies result in a cost in both cases.

Keywords: *wh*-in situ questions; French; covert dependency; subject advantage

1 Introduction

Most analyses of *wh*-questions posit the existence of a dependency between the scope position and the theta position of the *wh*-constituent, but they disagree on the nature of this dependency and in particular on whether it is the same or different in ex situ

***Corresponding author: Ruoxuan Li**, Laboratoire de Linguistique Formelle, Paris, France,
E-mail: ruoxuan.li@etu.u-paris.fr. <https://orcid.org/0000-0002-0103-4697>

Doriane Gras, Laboratoire de Psychologie du Développement et de l'Éducation de l'Enfant, Paris, France,
E-mail: doriane.gras@cnrs.fr. <https://orcid.org/0000-0002-6198-1450>

Caterina Donati, Laboratoire de Linguistique Formelle, Paris, France, E-mail: caterina.donati@u-paris.fr.
<https://orcid.org/0000-0002-7173-1041>

and in situ questions. The aim of this article is to verify how things stand in a mixed language like French, that allows both in situ and ex situ *wh*-phrases in normal information seeking questions.

We tested empirically whether in situ questions in French display the same subject advantage that has been widely attested in *ex situ* A-bar dependencies across languages and found a surprising object advantage in *in situ* questions. We explain this unexpected result by assuming that the same kind of movement dependency is established in both strategies but its landing site differs and hence interacts differently with the clitic left dislocation structure (CLLD) we introduced in our experimental items for clarity reasons. This conclusion is consistent with previous findings on the relevance of the distinction between crossing and nested dependencies.

The paper is organized as follows. Section 1.1 summarizes the debate on *wh*-in situ in general and in particular in French. Section 1.2 turns to processing and presents knowns and unknowns of the processing of *wh*-questions, in situ and ex situ, and sets the bases for the study discussed in Section 2. Section 3 discusses the results of the study and Section 4 provides a conclusion.

1.1 *Wh*-in situ questions across languages and in French

Colloquial French can be described as a ‘mixed’ language, as it shows both in situ and ex situ *wh*-phrases in normal information seeking questions (e.g., Dryer 2013) in contrast to other languages that show either only *wh*-in situ items (e.g., Mandarin Chinese or Japanese: [1]) or only *wh*-ex situ items (e.g., English: [2]).

- (1) a. Mandarin Chinese
Mali mai-le shenme?
 Mary buy-PERF what
- b. Japanese
Mary-ga nani-o katta-no?
 Mary-NOM what-ACC bought-Q
 ‘What did Mary buy?’
 (Miyamoto and Takahashi 2002: 135)

- (2) *What did Mary buy?*

In many languages the absence of *wh*-movement correlates with a peculiarity concerning the nature of the *wh*-elements. In Japanese and in Mandarin Chinese, *wh*-words such as *dare* ‘who’ and *nani* ‘what’ in (3a) and *shenme* ‘what’ in (3b) do not

carry an interrogative force *per se* and can be used as indefinite pronouns when produced in a non-interrogative context.

(3) Japanese

- a. *Dare-mo-ga nani-ka-o tabe-te-iru*
everyone-NOM **something**-ACC eating-is
 ‘Everyone is eating something.’
 (Nishigauchi 1990: Ch. 4)

Mandarin Chinese

- b. *Ta yiwei wo xihuan shenme*
 He think I like what
 ‘He thinks that I like something.’
 (Li 1992: 125)

This is not the case in French, where seemingly identical non-ambiguous *wh*-elements can be either left in situ or moved at the left periphery in a variety of different questions strategies, illustrated in (4).

- (4) a. *Qui Tu as vu ?* (wh-ex situ)
 Who you have seen
 b. *Tu as vu qui ?* (wh-in situ)
 you have seen who
 c. *Qui est-ce que tu as vu ?* (wh-ex situ with question particle ESK)
 Who Q you have seen
 d. *Qui as-tu vu ?* (wh-ex situ with subject-verb inversion)
 who have you seen
 ‘Who have you seen?’

A number of differences, variously described in the literature, suggest that the two strategies, in situ and ex situ, are not simply two options of the same grammar in French. First of all, they partially belong to different registers or varieties. While ex situ strategies (illustrated above in [4a, c, d]) belong to registers of varying degrees of formality, French *wh*-in situ in (4b) only belongs to a register we might call Colloquial French, which is mostly oral. Even when focusing on this register, though, *wh*-in situ and *wh*-ex situ questions display a number of semantic and syntactic differences. *Wh*-in situ questions have been widely described as more constrained than fronted *wh*-questions (Boeckx 1999; Bošković 2000; Chang 1997; Cheng and Rooryck 2000; Mathieu 1999). Many researchers have argued that *wh*-in situ in French is only a root phenomenon and cannot be found (i) in embedded clauses introduced by the complementizer *que*, (ii) in embedded infinitival clauses introduced by *de* or *à*, (iii) with negation, (iv) with quantifiers, and (v) with adverbs. Other scholars disagree and describe embedded in situ *wh*-questions with negation, quantifiers, and adverbs

as attested in the relevant register (Adli 2006; Baunaz 2011; Starke 2001). This sharp disagreement may be due to the fact that the linguists who have worked on the topic do not describe the same dialect of French. We shall leave aside this important controversy as not immediately relevant to the research described here, which only focuses on direct positive interrogative clauses.

Researchers importantly agree that in general processing *wh*-questions involves constructing a non-local dependency between the scope position of the question operator and the position where the *wh*-phrase receives its thematic role. In *wh*-ex situ questions, an overt dependency is built between the moved *wh*-phrase and its original position (i.e., the gap: [5a]), while in *wh*-in situ questions, a covert dependency is formed: (5b).

- (5) a. *wh*-ex situ: [*wh*-word_i [.....t_i.....]]
 b. *wh*-in situ: [*wh*-scope_i [...*wh*-word_i...]]

There are three main approaches to covert *wh*-dependencies: *LF-movement*, *unselective binding*, and *heavy pied-piping*.

In a nutshell, the LF-movement hypothesis (stemming from Huang 1982) proposes that both in situ and ex situ questions involve *wh*-movement, but that it happens abstractly at Logical Form (LF) in in situ *wh*-questions, and concretely in overt syntax in ex situ *wh*-questions. In Minimalist terms, this can be rephrased saying that in ex situ questions it is the higher copy of the *wh*-element that is pronounced, while the lower copy is pronounced in in situ questions (see Chomsky 1995). As a consequence, *wh*-questions, whether in situ or ex situ, are globally expected to exhibit a shared sensitivity to the locality constraints associated with movement under this view.

The unselective binding hypothesis (Pesetsky 1987; Watanabe 1992) claims that the difference between ex situ and in situ *wh*-questions is more profound since the dependency associated with *wh*-in situ is an anaphoric dependency, not a movement one. In situ *wh*-phrases do not move (or internally merge), but instead receive their quantificational force by being bound by an operator in the clausal periphery. Crucially, under this view, in situ *wh*-questions are not expected to obey the same locality constraints as ex situ *wh*-questions, which involve movement.

The heavy pied-piping hypothesis (Richards 2000) claims that what moves covertly is not the single *wh*-element, but the entire clause containing it. *Wh*-phrases pied-pipe their clause up to their scope position at LF. As a consequence, they are expected to escape many locality constraints (and notably a number of islands) that restrict simple *wh*-movement.

Bayer and Cheng (2017) combine the first two hypotheses (LF-movement and unselective binding) as parametric options depending on how scope marking works in any given language. For languages with obligatory overt scope markers, like Japanese, they propose that *wh*-in situ is derived without movement through

unselective binding. No syntactic constraint associated with movement (e.g., islands) is observed. For languages with no overt scope marker or no obligatory scope marker, like Mandarin Chinese, *wh*-in situ is derived through LF-movement, and as a result, the dependency is constrained as a movement dependency.

The question we aim to address empirically here is: what about mixed languages like Colloquial French, where *wh*-in situ is only one strategy for questions alongside *wh*-ex situ?

Baunaz (2011: 75) describes in situ and ex situ structures in non-standard Colloquial French as syntactically similar, both involving movement and presupposing identical processing loads (following Adli 2006). Faure and Palasis (2021) also describe *wh*-in situ questions in Colloquial French as derived through covert *wh*-movement. They show, however, that movement is less restricted in *wh*-ex situ questions, as evidenced by lower sensitivity to weak islands and intervention. They attribute these differences to the presence of a different, lower, trigger in *ex situ* questions, called Exclusivity.

Both analyses remain purely theoretical in their method and scope.

1.2 Processing *wh*-questions

Studies about processing covert dependencies confirm the prediction of theoretical analyses, that posit the existence of a long-distance dependency in *wh*-in situ, which has a processing cost. Aoshima et al. (2004), Ueno and Kluender (2009) found in Japanese that longer *wh*-in situ dependencies implied more processing costs compared to shorter ones. Xiang et al. (2014, 2015) found that Chinese *wh*-in situ questions incurred more processing costs than their declarative counterparts. When involving one embedding, they were processed slower than with no embedding. Similar results were obtained with both simplex ('who') and complex ('which') in situ *wh*-questions in French (Pablos et al. 2018).

As for ex situ *wh*-questions, subject questions are in general easier to comprehend than object questions (see van Gompel 2013 for adults' typical populations, Guasti 2004 for a discussion about children; Deevy and Leonard 2004; Levy and Friedmann 2009; van der Lely and Battell 2003 for children with language impairments; Sheppard et al. 2015 among others for agrammatic adults). There are various accounts for this subject advantage.

For one, the subject question in (6a) is easier to parse over the object question in (6b) because of the shorter linear distance between the *wh*-word and its gap.

- (6) a. *wh*-subject: **Who** __ met the girl ?
- b. *wh*-object: **Who** did the girl meet __ ?

A more structural account consists of appealing to hierarchical rather than to linear distance and suggesting that more deeply embedded constituents require more computational effort to parse than less deeply embedded ones: the subject question in (7a) is easier to parse than the object question in (7b) because of the shorter structural distance between the *wh*-word and its gap. These two accounts make different predictions when associated with word order variation. In particular, the latter predicts a universal subject gap that has so far proved to be correct (see Lau and Tanaka 2021 for a systematic review).

- (7) a. *wh*-subject: [_{CP} **Who** [_{TP} [_{VP} __ met the girl]]]?
 b. *wh*-object: [_{CP} **Who** did [_{TP} [_{VP} the girl meet __]]]?

Among structural accounts, Featural Relativized Minimality (Friedmann and Rizzi 2009; Rizzi 1990) explains the subject advantage in yet a different, empirically accurate, way. According to Featural Relativized Minimality, the local relation between an extracted element and its trace is disrupted when it crosses an intervening element whose morphosyntactic featural specification matches the specification of the elements it separates. This approach naturally leads to a system that can capture degrees of processing difficulty: the relative acceptability of an intervention configuration will vary as a function of the total, partial, or zero featural overlap between the intervener and the target. In a nutshell, configurations involving less featural overlap are predicted to be easier than sentences involving a higher degree of overlap (Villata et al. 2016). Subject questions like (8a) are easy because no intervener with any overlap with the target holds in the relevant configuration.

- (8) a. *wh*-subject: [_{CP} **Who** [_{TP} [_{VP} __ met the girl]]]?
[+WH; +D] [+WH; +D]
b. bare *wh*-object: [_{CP} **Who** did [_{TP} [_{VP} the girl meet __]]]?
[+WH; +D] [+D] [+WH; +D]
c. complex *wh*-object: [_{CP} **Which friend** did [_{TP} [_{VP} the girl meet __]]]?
[+WH; +D; +NP] [+D; +NP] [+WH; +D; +NP]
d. *wh*-island: [_{CP} **which friend** did [_{TP} [_{VP} which girl meet __]]]?
[+WH; +D; +NP] [+WH; +D; +NP] [+WH; +D; +NP]

Bare *wh*-object questions like (8b) are more difficult because of the partial featural identity between the object and the intervening subject: both the subject and the *wh*-element are DPs and as such they share the same categorial [+D] feature. Since the intervention is only partial (concerning only one feature), the sentence can still be processed, although with some additional difficulty. But the processing difficulty associated with object questions is predicted to be stronger in *which*-object questions as (8c), where an additional [+NP] feature is shared with the non-*wh* lexical subject. The pattern of acquisition of *wh*-questions in children appears to conform to this

prediction, as among object questions, *wh*-questions with lexical restrictions (e.g., ‘which dog’) emerge significantly later in L1 acquisition than bare *wh*-questions (e.g., ‘who’) and are associated with difficulty until about age seven in monolingual development across languages (see, e.g., Avrutin 2000; Stromswold 1995; Tracy 1994).¹ Finally, (8d), displaying a total featural overlap configuration, corresponds to ungrammaticality (a *wh*-island).

Going back to the subject advantage, it has also been suggested that object questions are more difficult, in SVO languages at least, because after movement the arguments in the sentence are no longer in a canonical word order (Diessel and Tomasello 2005). This makes it more difficult to infer a meaning based on surface form. Subject questions are easier presumably because there are no overt word-order differences between the interrogative and declarative forms.

What do these different accounts predict as far as *wh*-in situ in a mixed language like French is concerned, and how does this interact with the various approaches to *wh*-in situ we have discussed briefly in the preceding section?

Starting from the canonical word order approach, this predicts **no subject advantage** in *wh*-in situ questions no matter the underlying derivation, since by definition in situ questions display canonical word order across the board.

On the other hand, if *wh*-in situ involves covert movement, all accounts predict *wh*-in situ questions to display a subject advantage as it obeys the same constraint as overt movement by hypothesis. If Featural Relativized Minimality is on the right track, moreover, we further expect object ‘which’ questions to be more difficult to comprehend than object ‘who’ questions. These effects are not expected if *wh*-in situ involves an unselective binding dependency, since a quantifier binds all unbound variables in its scope and a full definite NP in subject position is not expected to intervene in such dependency. Faure and Palasis (2021)’s description, which associates *wh*-ex situ and *wh*-in situ to different features, would be compatible with different featural Minimality effects for the two covert dependencies.

On the empirical side, little has been done to investigate whether *wh*-in situ questions display a subject advantage and whether this is modulated by the type of *wh*-element involved. A recent study on French (Bentea 2016) found a subject advantage only in ex situ questions. Object *which*-ex situ questions yielded lower response accuracy scores than subject *which*-ex situ questions, while *who*-questions

¹ But see Avrutin (2000, 2006) for an alternative account, suggesting that since ‘which’ phrases require linking to specific sets of discourse entities while *wh*-pronouns such as ‘who’ do not, computing mental representations for ‘which’ questions requires more computational resources than computing ‘who’ questions.

displayed no significant difference with respect to object ones, due to a ceiling effect. In contrast, higher accuracy scores were obtained with object in situ questions, irrespective of whether the *wh*-element was ‘who’ (*qui*) or ‘which’ (*quel*). French Sign Language (LSF) on the other hand have recently been shown to exhibit a subject advantage in in situ *wh*-questions (Hauser et al. 2023): authors found a significantly higher response accuracy in subject questions than in object questions and this effect was stronger in *wh*-questions with lexical restrictions (equivalent to *which*-questions) than in bare *wh*-questions (equivalent to *who*-questions).

The goal of our study is to verify if there is a subject-object asymmetry in the processing of French in situ *wh*-questions, as well as in ex situ *wh*-questions. If *wh*-in situ in French is derived through the same covert movement operation as the one deriving overt *wh*-questions, we expect to observe a subject advantage in the comprehension of *wh*-in situ questions as well as in that of *wh*-ex situ questions. We expect moreover to find a stronger effect in *which*-questions than in bare *wh*-questions, due to Featural Relativized Minimality effects. If Faure and Palasis (2021) are on the right track, though, and the movement path for *wh*-ex situ and for *wh*-in situ is not the same, we might also expect to see some difference between the two strategies.

We constructed two comprehension and acceptability experiments including in situ and ex situ questions, with the aim of verifying whether they display the same pattern.

2 Two experiments on *wh*-questions comprehension

Since the existence of a subject advantage in overt dependencies is a well-established generalization, *wh*-ex situ questions are used as controls in the experiments we set. The rationale is simple: if we manage to detect a subject-object asymmetry in our results for *wh*-ex situ questions, we can take in this result as a confirmation that our methodology is sensitive enough to detect the cost of this processing, and that we can rely on it to test whether the same cost is associated with processing *wh*-in situ. If on the other hand, no such asymmetry is to be observed, we would know our task is not sophisticated enough to detect the reflexes of a *wh*-dependency.

Notice however that our aim was not to compare the two strategies, but rather to verify whether some kind of subject-object asymmetry was to be found within each of them. What we needed was thus a complete paradigm of subject and object questions that were minimally different within the two strategies, but that were not necessarily paired across them.

2.1 The quest for the ideal paradigm

While preparing our experiment, we soon realized that constructing unambiguous *wh*-questions in French is not easy, as confirmed by several pilot studies involving about five native French speakers each. Our aim was to construct items that were unambiguous both for their status (in situ vs. ex situ) and for their interpretation (subject vs. object), and finally both with *qui* ‘who’ and with *quel* ‘which’. We faced two major problems in trying to build such a complex paradigm: (i) the lack of *quel*-questions in subject ex situ conditions and (ii) the ambiguity in the comprehension and in status of in situ subject questions.

Starting with ex situ questions, in fact, particle questions, i.e., questions with *est-ce que/qui*, which have the advantage of clearly marking the ex situ strategy, are defective: they are not available with *quel* ‘which’ subject questions:

- (9) **Quelle dame est-ce qui pousse le monsieur ?*
 Which lady Q pushes the gentleman
 Intended: Which lady is pushing the gentleman?

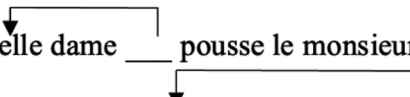
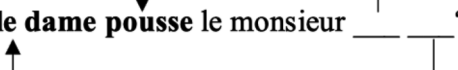
Since French does not allow *quel*-subject to be used with *est-ce que*, we tried to provide a complete paradigm by resorting to another strategy, namely interrogatives with inversion. However, as shown by (10a), which contrasts minimally with the corresponding object condition in (10b), subject questions with inversion are not grammatical either.²

2 The incompatibility of the verb with a suffixed subject with a *wh*-subject is confirmed by the *Grande grammaire du français* (Abeillé and Godard 2023). The suffixed subject verb is only compatible with a preverb subject (i). It is impossible with a *qui* (‘who’)-subject (ii), and difficult with a *quel* (‘which’)-object (iii).

- (i) *Qui Paul a-t-il rencontré?*
 Who Paul has-he met
 ‘Who did Paul meet?’
- (ii) **Qui vient-il?*
 Who comes-it
 ‘Who’s coming?’
- (iii) ?*Quelle solution a-t-elle été mise en place en attendant ?*
 Which solution has-she been put in place in-the-meantime
 ‘What solution has been put in place in the meantime?’
 (la-csf.org, 8 Aug, 2012)

- (10) a. **Quelle dame pousse-t-elle le monsieur ?* (subject ex situ questions)
 Which lady pushes-she the gentleman
 Intended: Which lady is pushing the gentleman?
- b. *Quel monsieur la dame pousse-t-elle ?* (object ex situ questions)
 Which gentleman the lady pushes-she
 ‘Which gentleman is the lady pushing?’

The participants to our pilots clearly preferred the non-inverted version, as in *Quelle dame pousse le monsieur?* ‘Which lady pushes the gentleman?’ This strategy of subject questions is however ambiguous both for its derivational status and for its meaning. As for its status, it could correspond both to an in situ and an ex situ derivation.³ As for its interpretation, it can be interpreted both as a subject question or an object question with inverted subject, as represented in (11).

- (11) a. *Quelle dame* — *pousse le monsieur*

 b. *Quelle dame pousse le monsieur* — — ?

 ‘Which lady is the gentleman pushing?’

In (11b), the verb and *wh*-object move to the front of the subject or in other words, the subject is posited post verbally in sentences where the initial constituent (*quelle dame*) is extracted. This inversion, generally known as *stylistic inversion* (Kayne and Pollock 1978), is optional and occurs in interrogative (12a) or exclamative sentences (12b).

- (12) a. *Marie parlera quand ?* | *Quand parlera Marie ?*
 Marie will_speak when | when will_speak Marie
 ‘When will Mary speak?’
- b. *Quelle chance Paul a !* | *Quelle chance a Paul !*
 Which luck Paul has | which luck has Paul
 ‘What luck Paul has !’

To avoid both types of ambiguities, we first tried to embed the interrogative under a *believe*-verb, as in (13):

³ Interestingly, *Quelle dame pousse le monsieur?* is more likely to be a *wh*-in situ question given anti-locality analyses for subject extraction (Erlewine 2020), which assume that movement of a subject from Spec, TP to Spec, CP will be blocked if there are no intervening functional projections (XPs). But these accounts are far from being universally accepted in the field.

- (13)  Quelle dame tu crois qui pousse le monsieur ? (subject ex situ)
 which lady you believe who pushes the gentleman
 ‘Which lady do you believe is pushing the gentleman ?’

Note that French does not allow the complementizer *que* ‘that’ to be followed by a gap, but it is possible to use the form *qui* instead (Rizzi and Shlonsky 2007; Rooryck 2000; Taraldsen 2001). This allowed us to construct a complete set of items filling the 4 conditions.

- (14) a. *Tu crois que quelle dame pousse le monsieur ?* (subject in situ)
 You believe that which lady pushes the gentleman
 b. *Tu crois que le monsieur pousse quelle dame ?* (object in situ)
 You believe that the gentleman pushes which lady
 c. *Quelle dame tu crois qui pousse le monsieur ?* (subject ex situ)
 which lady you believe who pushes the gentleman
 d. *Quelle dame tu crois que le monsieur pousse ?* (object ex situ)
 which lady you believe that the gentleman pushes

Sentences in (14c-d) display no ambiguities in interpretation nor in status. The problem is that our participants disliked this strategy when applied to *qui*-questions. They rated badly most of the items in the *qui*-paradigm:

- (15) a. *?Tu crois que qui pousse le monsieur ?* (subject in situ questions)
 You believe that who pushes the gentleman
 b. *Tu crois que le monsieur pousse qui ?* (object in situ questions)
 You believe that the gentleman pushes who
 c. **Qui tu crois qui pousse le monsieur ?* (subject ex situ questions)
 who you believe who pushes the gentleman
 d. *?Qui tu crois que le monsieur pousse ?* (object ex situ questions)
 who you believe that the gentleman pushes

Our participants would rather use *Qui pousse le monsieur ?* ‘Who is pushing the gentleman?’ and *Par qui le monsieur est-il poussé ?* ‘By whom is the gentleman pushed?’ instead of (15a) and (15c). (15b) is slightly better than the others but less acceptable than *Qui est-ce que le monsieur pousse ?* and *Qui le monsieur pousse-t-il ?* ‘Who is the gentleman pushing?’.

The final decision was then to use clefted questions as in (16).

- (16) a. *C’est qui/quelle dame qui ___ pousse le monsieur ?* (subject ex situ)
 It’s who/which lady that pushes the gentleman
 b. *C’est qui/quelle dame que le monsieur pousse ___ ?* (object ex situ)
 It’s who/which lady that the gentleman pushes

These questions have the double advantage of being clearly *ex situ* and displaying no ambiguity in interpretation. The status and interpretation of clefted questions is far from clear, however. Some scholars have claimed that *wh*-elements in clefted questions do not move as far as to their scope position (Chang 1997), and they are thus not genuine instances of *wh-ex situ* questions. We acknowledge this complication but decided to stick to these items in lack of better alternatives. Indeed, it is uncontroversial that the *wh*-element is not in its *in situ* position either, and that the clefted question contains a gap.⁴ We shall consider this to be the minimal defining property of any *ex situ* strategy. We are of course aware that clefted questions are more complex than questions with no clefting. But recall that the aim of our experiment here is to verify if there is a subject-object asymmetry to be observed within each of the strategies of interrogatives, and not to compare them.

There is some normative stigma associated with the use of double *qui* in *c'est qui qui* since *qui donc*, *qui est-ce qui*, *qui c'est qui* or a single *qui* are deemed more elegant. This is one of the reasons why we decided to develop an oral version of the experiment, assuming that oral French is closer to Colloquial French and thus less prone to normative pressures (see below).

Turning to *in situ* conditions, we found that simple intonation questions are still potentially ambiguous in various ways. Subject questions with simple intonation can be either *in situ* or *ex situ*, and several interpretation ambiguities also arise. A sentence like *Qui lave le chat?* can be interpreted not only as a subject question 'Who is washing the cat?' but also an object question with subject inversion 'Who is the cat washing?' (as in [11]). We avoided these potential ambiguities by inserting a clitic left dislocation structure (CLLD), as in (17):

- (17) a. *Le monsieur, qui/quelle dame le pousse ?* (subject *in situ* questions)
 the gentleman who/which lady him pushes
 b. *Le monsieur, il pousse qui/quelle dame?* (object *in situ* questions)
 the gentleman he pushes who/which lady

In situ questions in (17a-b) are preceded by a left dislocated (i.e., topicalized) masculine NP, associated with a resumptive proclitic, and this addition makes them non-ambiguous and natural.

For the sake of uniformity, we decided to add this topicalization pattern across the board in all items. *Ex situ* questions (18a-b) also included topicalized masculine NPs besides being clefted to avoid the ambiguities discussed above.

⁴ See more about French cleft structures in Belletti (2008, 2009, 2015), Destruel (2013), among many others.

- (18) a. *Le monsieur, c'est qui/quelle dame qui le pousse?* (subject ex situ)
 the gentleman it's who/which lady that him pushes
 b. *Le monsieur, c'est qui/quelle dame qu'il pousse?* (object ex situ)
 the gentleman it's who/which lady that he pushes

The result we obtained is a homogenous set of questions, all containing a clitic left dislocation structure (CLLD), displaying no ambiguity in interpretation, and a distinction between ex situ and in situ which is again non ambiguous and capitalizes on clefting versus non-clefting.

We also eliminated verbs starting with a vowel (e.g., *appeler* 'call', *aimer* 'love', ...) and only used verbs starting with a consonant (e.g., *prévenir* 'inform', *désirer* 'desire'...) to avoid perceptual ambiguities in sequences such as *qui l'appelle* 'who calls him' and *qu'il appelle* 'who he calls':/i/ in the latter is a little longer than in the former but this difference is very subtle and difficult to distinguish. We then avoided female referents and kept only male referents to avoid any gender bias that might play a role wherever social gender is involved: men are more likely to be in a subject position and to be referred to via pronouns (Richy and Burnett 2020).

A potential issue with this data setting is that it is particularly natural in spoken Colloquial French but may be somewhat stigmatized in written French. This is why we decided to replicate the experiment in two modalities: one being a reading task, and the other a listening task. We expect any effect observed in the results to be relatively larger in the listening task compared to the reading task.

More generally, we expect to observe a subject advantage in cleft structures (as in [18]), because they are associated with overt *wh*-extraction (although with the proviso discussed above). If the experiment is well-designed and sufficiently sensitive, the effect should be stronger with *quel* 'which' questions than with *qui* 'who' questions. If the data pattern is as predicted in the ex situ condition, we will be able to consider the results in the in situ condition as evidence for their processing and derivation. If in situ questions involve covert movement to the same position as *wh*-ex situ, the hypothesis is that they will display the same subject advantage, and with the same modulation based on *wh*-material. If the dependency involved is of a different type (or targeting a different position; see above) the hypothesis is that no subject advantage should hold.

In what follows, we present the two experiments only once, because the items were just the same with the only difference concerning the modality of representation. The results will be discussed separately, due to some minor differences.

2.2 Materials

In total, our experiment included 48 groups of *wh*-questions: 24 sets of *who*-questions (*qui*-questions) and 24 sets of *which*-questions (*quel*-questions). Each set of *wh*-questions followed a 2×2 design: (A) in situ (19a, b)/ex situ (19c, d), and (B) subject (19a, c)/object (19b, d).

- (19) a. *Le monsieur, qui/quelle dame le pousse ?* (subject in situ)
 the gentleman who/which lady him pushes
 ‘The gentleman, who/which lady is pushing him?’
- b. *Le monsieur, il pousse qui/quelle dame?* (object in situ)
 the gentleman he pushes who/which lady
 ‘The gentleman, who/which lady is he pushing?’
- c. *Le monsieur, c’est qui/quelle dame qui le pousse?* (subject ex situ)
 the gentleman it’s who/which lady that him pushes
 ‘The gentleman, it’s who/which lady that is pushing him?’
- d. *Le monsieur, c’est qui/quelle dame qu’il pousse?* (object ex situ)
 the gentleman it’s who/which lady that he pushes
 ‘The gentleman, it’s who/which lady that he is pushing?’

The experimental task was twofold: participants had to rate each item, and then to reply to each question by selecting a character on a picture with three characters. Two of them were similar, either performing an action toward the third, distinct character standing between them, or experiencing the same action performed by the third character in the middle. An illustration is given in Figure 1. Details on the procedure are given in the next section.

Each participant saw only one item from each set and encountered each condition an equal number of times. In total, 48 items were shown to each participant. In addition to the test items, there were 14 *where*-questions (*où*-questions) included as fillers, as in (20). These filler questions always targeted the side characters, similar to the test items (see Figure 2). Four of the fillers were ungrammatical distractors (e.g., *Mouton le, il où est?* ‘sheep the, he where is?’) to ensure that participants were paying attention to the sentence items. These sentences were expected to be rated as 1 or 2. The data of participants who gave a high rating to these very bad sentences were excluded by the analysis.

- (20) fillers
La princesse avec les cheveux bleus, elle est où?
 ‘The princess with the blue hair, she is where?’



Figure 1: Picture illustrating the pushing event of (19): the woman on the right matches with the subject question in (19a, c) and the woman on the left matches with the object question in (19b, d).



Figure 2: Picture corresponding to the filler *where*-question in (20): the princess on the right answers the question.

All items were randomized into four counterbalanced lists with the same fillers (see Appendix: <https://osf.io/y9x2h>). All items and fillers were recorded by a French native speaker in a soundproof room. Each sentence lasted approximately 3–4 s and was played at a volume of 20 dB.

2.3 Tasks and procedure

The test was hosted on PCibex (<https://farm.pcibex.net/>). It began with a consent page clarifying the purpose and procedure of the experiment,⁵ followed by a page gathering some basic information such as age, gender and native language and then proceeded to an instruction page. A short training phase, including four practice

⁵ The experiment complied with the GDPR (General Data Protection Regulation). The experiment was approved by the PGSL (Paris Graduate School of Linguistics) Ethical Committee.

items, with no time limit, was provided before each trial, allowing participants to familiarize themselves with the experiment. Answers from the training phase were not considered in the data analysis. The duration of each testing session was approximately 15 min.

Since the experiment was conducted online, we were unable to provide face-to-face instructions to participants, who were informed that they needed to complete the task in a well-connected and quiet environment. Participants were also free to withdraw at any time due to unstable internet connections, interruptions, or fatigue.

For each trial, participants either read (if performing the reading task) or listened (if performing the listening task) to the question stimulus and had then to perform two tasks:

- an acceptability rating task: they were asked to rate its acceptability on a 1–7 Likert scale. The time limit was 10 s.
- a forced choice task: each question was followed by a picture from which participants had to choose a character to answer, with a response time limit of 5 s.

Acceptability rating tasks are known to have relatively higher statistical power with smaller sample sizes compared to other informal offline methods (e.g., forced-choice, magnitude estimation, yes-no tasks) (Sprouse and Almeida 2017: 28). The task was meant to allow us to establish a scale for the comprehension levels of different sentences, enabling comparisons across conditions and supporting our factorial experimental design.

The forced-choice task is known to be highly sensitive in detecting differences between two contrasting conditions (Sprouse and Almeida 2017: 25), so the task was also meant to allow us to identify differences across conditions. But response accuracy was also meant as a sanity check, to verify whether participants were properly reading and understanding the experimental items. This is why we used the combination of these two tasks.

We expected the acceptability rating task to yield more significant results than the forced-choice task, as the former is more sensitive when testing the same syntactic phenomena (Sprouse and Almeida 2017: 28). We also expected the results of the forced-choice task to align with those of the acceptability judgment, with only minor differences. Specifically, higher response accuracy should correspond to higher ratings. Given our research hypothesis, if all *wh*-questions in French displayed a subject advantage, participants would be expected to answer subject questions more accurately than object questions and to rate subject questions higher than object ones.

2.4 Participants

A total of 68 participants completed our reading task. Most were recruited online through RISC (*Le relais d'information sur les sciences de la cognition*), while two participants joined after seeing the recruitment notice on campus. Among them, two were non-native speakers of French, and four had previously participated in the experiment as pilot testers (three of these four had low response accuracy: 54.76 %, 59.57 % and 62.50 %); their results were removed from the analyses. Three participants gave high ratings (6 or 7) to ungrammatical distractors. Although we removed their ratings, we retained their answers to the other questions since their response accuracy fell within the acceptable margin of error (i.e., greater than the mean minus two standard deviations, which is 67.49 %). The remaining valid participants were all native French speakers,⁶ consisting of 20 men and 42 women. Each participant received a €3 gift card as a reward. The average age was 28.27 years (range: 19 to 45).

A total of 62 participants consisting of 18 men and 44 women completed our listening task. Most were recruited online through RISC, Facebook, and SurveyCircle. Two participants were excluded due to low response accuracy (56.10 % and 60.42 %). Additionally, four participants gave high ratings (6 or 7) to ungrammatical distractors. Although we removed their ratings, we retained their answers to the *wh*-questions since their response accuracy fell within the acceptable margin of error (i.e., greater than the mean minus two standard deviations, which is 72.19 %). Each participant received a €3 gift card as a reward. The average age was 30.58 years (range: 18–66).

2.5 Results

2.5.1 The reading task

Results for *who*-questions (*qui*-questions) and *which*-questions (*quel*-questions) were analyzed separately. Ten sentence items (within a given list) with mean response accuracy smaller than the overall mean by more than two standard deviations (SD) (i.e., smaller than 64.56 %) were removed. The results were analyzed using R Studio. Response accuracy was analyzed using generalized linear mixed-effects models with

⁶ We did not dig into the dialectal background of the participants because our test items were all simple clauses widely used in Colloquial French. This was confirmed by our pilots. Debates regarding the constraints associated with French *wh*-in situ questions are mostly centered around embedded clauses or involving negation and adverbs, none of which are included in the present study.

glmer() from the *lme4* package (Bates et al. 2015), while acceptability judgments were analyzed using cumulative link models with *clm()* from the *ordinal* package (Christensen 2019). Log-transformed response times $\log(\text{RT})$ were analyzed using linear mixed-effects models with *lmer()* from the *lmerTest* package (Kuznetsova et al. 2017).⁷

Random effects included intercepts for subjects and items. The fixed factors were question type (subject question vs. object question) and the position of the *wh*-element (in situ vs. ex situ). All effects with a *p*-value less than 0.05 were considered significant.

Starting with mean response accuracy, Figure 3 shows a significant interaction between the two main factors (subject/object and in situ/ex situ) for *qui*-questions ($\beta = 0.6029$, $z = -2.498$, $p < 0.05$), reflecting a significant subject-object asymmetry in the ex situ condition only ($\beta = 0.2990$, $z = 3.352$, $p < 0.001$). This indicates that there

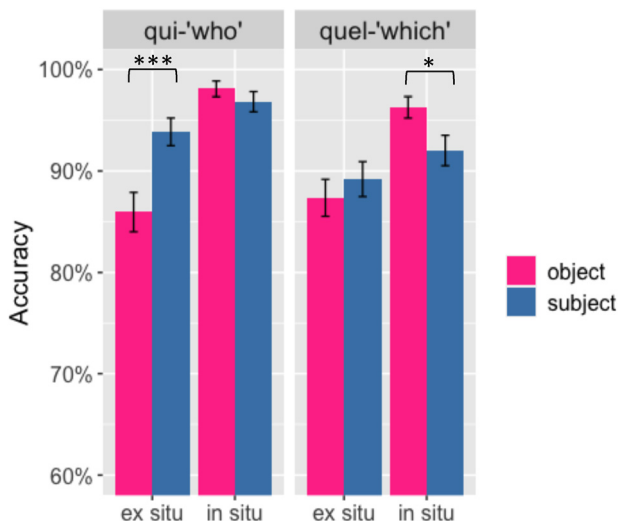


Figure 3: Mean accuracy rates (%) for *qui*-questions (left panel) and *quel*-questions (right panel) for in situ (right in each panel) and ex situ (left) questions across subject (blue) and object (pink) conditions. Error bars represent standard errors of the mean and stars are used to indicate statistically significant results.

⁷ They were written as (all data files and scripts can be found on the OSF public repository: <https://osf.io/7yru6/>):

- (i) `glmer(accuracy ~ SUOB*EXIN + (1|Subject) + (1|Item), data = Data, family= binomial)`
- (ii) `clm(judgment ~ SUOB*EXIN, data = Data)`
- (iii) `lmer(log(RT) ~ SUOB*EXIN + (1|Subject) + (1|Item), data = Data)`

were more correct answers for subject questions than for object questions (i.e., subject advantage). For *quel*-questions, there is also a significant interaction between the two main factors ($\beta = 0.5210$, $z = -2.175$, $p < 0.05$). In this case, a subject-object asymmetry appears only in the in situ condition ($\beta = 0.4017$, $z = -1.997$, $p < 0.05$), with object questions answered more accurately than subject questions (i.e., object advantage).

The mean ratings of *qui/quel*-questions are presented in Figure 4.

Figure 4 shows a slight but not significant interaction between subject/object and in situ/ex situ conditions for *qui*-questions ($\beta = 0.19392$, $z = -1.647$, $p < 0.1$). A subject-object asymmetry is present, but only in the in situ condition ($\beta = 0.1380$, $z = -2.867$, $p < 0.01$). This interaction is significant for *quel*-questions ($\beta = 0.195204$, $z = -4.060$, $p < 0.001$), reflecting a subject-object asymmetry also in the in situ condition ($\beta = 0.1405$, $z = -5.608$, $p < 0.001$). The mean ratings for object questions are significantly higher than those for subject questions (i.e., object advantage).

The mean response time to *qui/quel*-questions are presented in Figure 5.

In Figure 5, there is no interaction between subject/object and in situ/ex situ conditions in either case (*qui*: $\beta = 2.548e-02$, $t = 0.316$, ns; *quel*: $\beta = 0.023468$, $t = 0.257$, ns). We observe no subject-object asymmetry in *qui*-questions ($\beta = 1.808e-02$,

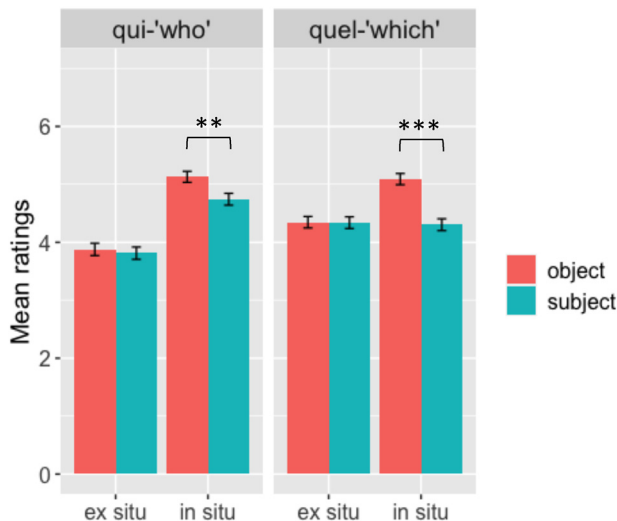


Figure 4: Mean ratings for *qui*-questions (left panel) and *quel*-questions (right panel) in situ (right in each panel) and ex situ (left) questions across subject (blue) and object (orange) conditions. Error bars represent standard errors of the mean and stars are used to indicate statistically significant results.

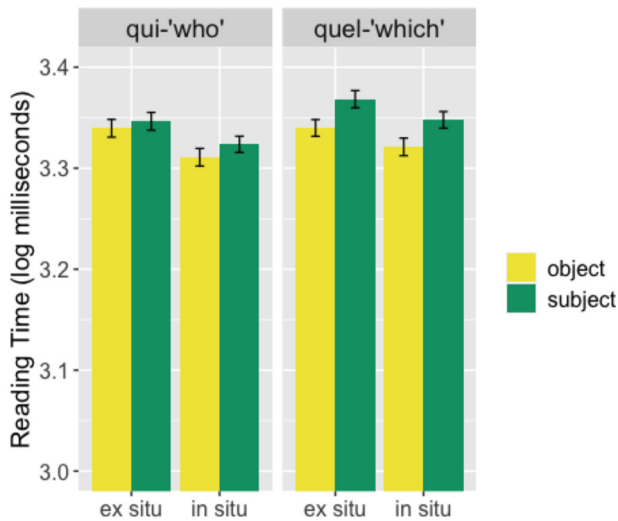


Figure 5: Mean response time (log milliseconds) for *qui*-questions (left panel) and *quel*-questions (right panel) for in situ (right in each panel) and ex situ (left) questions across subject (green) and object (yellow) conditions. Error bars represent standard errors of the mean and stars are used to indicate statistically significant results.

$t = -0.028$, ns) or in *quel*-questions ($\beta = 0.016499$, $t = 1.355$, ns). For *qui*-questions, there is a slight but not significant difference between in situ and ex situ conditions ($\beta = 1.790e-02$, $t = -1.734$, $p < 0.1$). This difference is not present in *quel*-questions ($\beta = 0.016526$, $t = -1.012$, ns).

The mean reading time to *qui/quel*-questions is presented in Figure 6.

Figure 6 shows an interaction between subject/object and in situ/ex situ conditions for *qui*-questions ($\beta = 0.01708$, $t = 4.211$, $p < 0.001$), while this interaction is not significant for *quel*-questions ($\beta = 0.01730$, $t = 1.696$, ns). For *qui*-questions, subject questions are read faster than object questions in ex situ conditions ($\beta = 0.009898$, $t = -4.545$, $p < 0.001$), indicating a subject advantage. There is no significant difference in in situ conditions ($\beta = 0.01370$, $t = 1.853$, ns). For both types of questions, the mean reading time for ex situ questions is significantly longer than for in situ questions (*qui*: $\beta = 0.01209$, $t = -6.360$, $p < 0.001$; *quel*: $\beta = 0.01221$, $t = -4.707$, $p < 0.001$). There is no significant subject-object asymmetry within any condition.

2.5.2 The listening task

Results for *whom*-questions (*qui*-questions) and *which*-questions (*quel*-questions) were analyzed separately. Seventeen items with mean response accuracy smaller

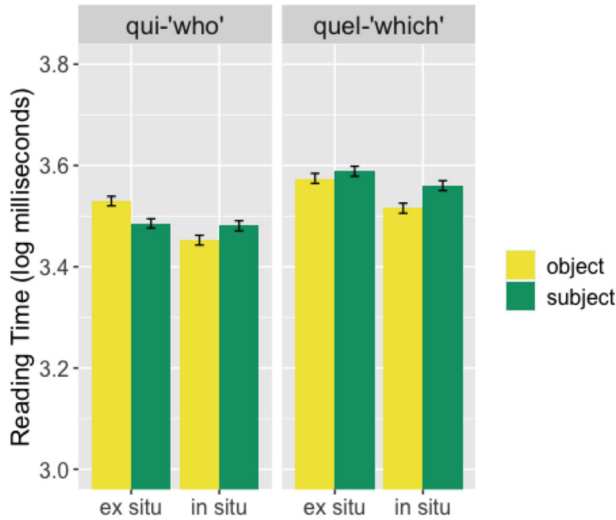


Figure 6: Mean reading time (log milliseconds) for *qui*-questions (left panel) and *quel*-questions (right panel) for in situ (right in each panel) and ex situ (left) questions across subject (green) and object (yellow) conditions. Error bars represent standard errors of the mean and stars are used to indicate statistically significant results.

than the overall mean by more than two standard deviations (SD) (i.e., smaller than 64.16 %) were removed. The models used are the same as those in Section 2.5.1 of the reading task.

Figure 7 represents the mean response accuracy. There is an interaction between the two main factors (subject/object and in situ/ex situ) in both cases (*qui*: $\beta = 0.9074$, $z = -2.913$, $p < 0.001$; *quel*: $\beta = 0.5946$, $z = -2.615$, $p < 0.001$). The difference between subject and object questions in the ex situ condition is significant. It shows a subject advantage in both types of *wh*-questions (*qui*: $\beta = 0.5256$, $z = 3.366$, $p < 0.001$; *quel*: $\beta = 0.3268$, $z = 4.375$, $p < 0.001$). This means that there are more correct answers in subject questions than in object questions (i.e., subject advantage). No significant asymmetry shows in in situ conditions.

The mean ratings of *qui/quel*-questions are presented in Figure 8.

Figure 8 shows almost no interaction between subject/object and in situ/ex situ conditions for *qui*-questions ($\beta = 0.21599$, $z = -0.662$, ns) and for *quel*-questions ($\beta = 0.21790$, $z = -1.817$, $p < 0.1$). However, there is an object advantage in *quel*-questions, but only in in situ conditions ($\beta = 0.1564$, $z = -3.111$, $p < 0.01$). The mean ratings for object questions are significantly higher than those for subject questions (object advantage).

The mean response time to *qui/quel*-questions are presented in Figure 9.

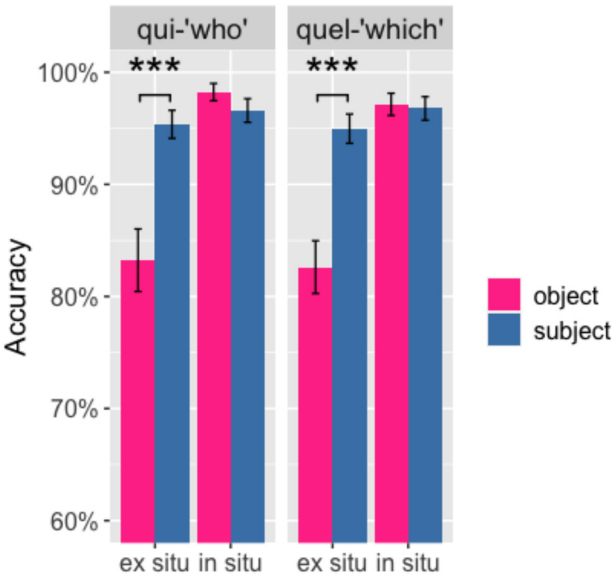


Figure 7: Mean accuracy rates (%) for *qui*-questions (left panel) and *quel*-questions (right panel) for in situ (right in each panel) and ex situ (left) questions across subject (blue) and object (pink) conditions. Error bars represent standard errors of the mean and stars are used to indicate statistically significant results.

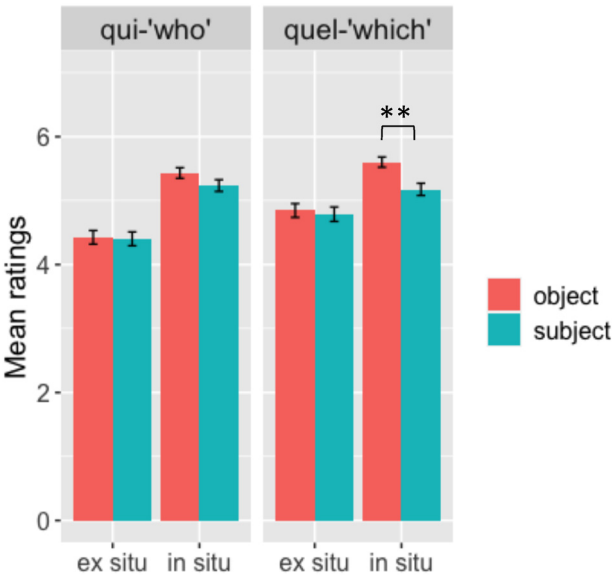


Figure 8: Mean ratings for *qui*-questions (left panel) and *quel*-questions (right panel) for in situ (right in each panel) and ex situ (left) questions across subject (blue) and object (orange) conditions. Error bars represent standard errors of the mean and stars are used to indicate statistically significant results.

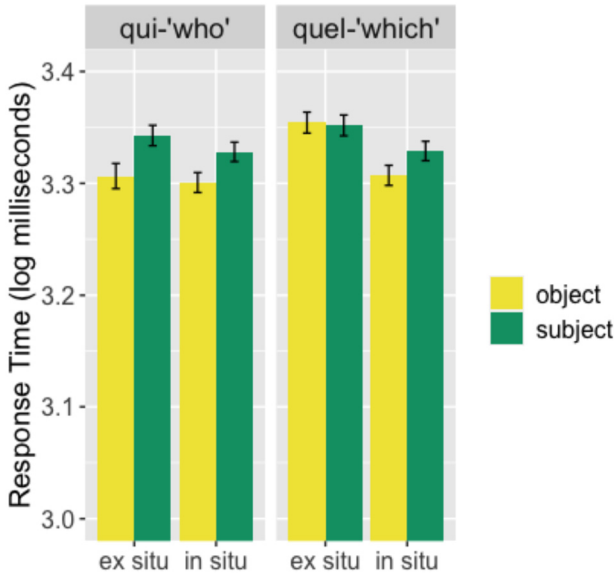


Figure 9: Mean response time (log milliseconds) for *qui*-questions (left panel) and *quel*-questions (right panel) for in situ (right in each panel) and ex situ (left) questions across subject (green) and object (yellow) conditions. Error bars represent standard errors of the mean and stars are used to indicate statistically significant results.

In Figure 9, there is no interaction between subject/object and in situ/ex situ conditions in either case (*qui*: $\beta = 0.025017$, $t = 0.048$, ns; *quel*: $\beta = 0.024767$, $t = 1.067$, ns). This indicates that there is no subject-object asymmetry in *qui*-questions ($\beta = 0.018871$, $t = 1.593$, ns), nor in *quel*-questions ($\beta = 0.017719$, $t = -0.156$, ns). For *quel*-questions, there is a significant difference between in situ and ex situ conditions ($\beta = 0.017708$, $t = -2.596$, $p < 0.05$). Ex situ questions have longer response times, indicating that they are relatively more difficult to process. This difference is not present in *qui*-questions ($\beta = 0.018843$, $t = -0.537$, ns).

The mean listening time to *qui/quel*-questions is presented in Figure 10.

Figure 10 shows no interaction between subject/object and in situ/ex situ conditions for *qui*-questions ($\beta = 0.011159$, $t = 0.206$, ns) and *quel*-questions ($\beta = 0.010794$, $t = 0.395$, ns), indicating no subject-object asymmetry in *qui*-questions ($\beta = 0.007896$, $t = 0.851$, ns) or *quel*-questions ($\beta = 0.007612$, $t = -0.284$, ns). Additionally, there is no difference between in situ and ex situ conditions (for *qui*: $\beta = 0.007886$, $t = -0.525$, ns; for *quel*: $\beta = 0.007624$, $t = 0.633$, ns).

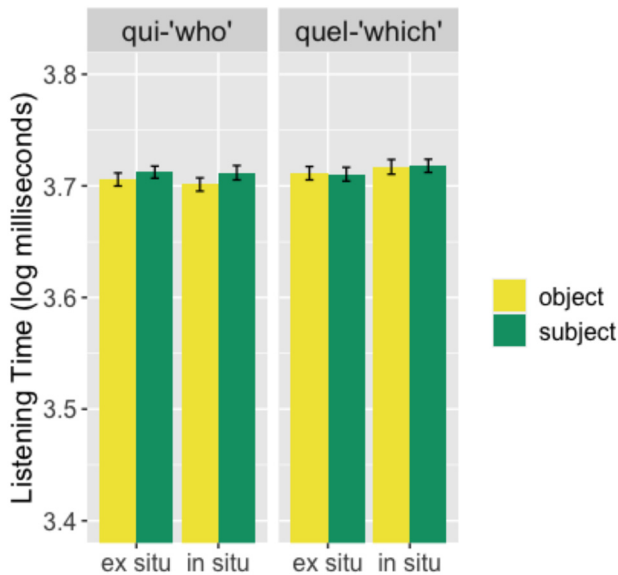


Figure 10: Mean listening time (log milliseconds) for *qui*-questions (left panel) and *quel*-questions (right panel) for in situ (right in each panel) and ex situ (left) questions across subject (green) and object (yellow) conditions. Error bars represent standard errors of the mean and stars are used to indicate statistically significant results.

2.5.3 Summary

In both experiments, we did not obtain an exact one-to-one correspondence between response accuracy and acceptability ratings; however, they tended to converge rather than diverge, despite differing in significance.

For instance, in the reading task, response accuracy demonstrated a subject-object asymmetry in both ex situ and in situ conditions ($*** p \leq 0.001$ and $* p \leq 0.05$, respectively, as shown in Figure 3), while mean ratings showed the same asymmetry only in the in situ condition ($** p \leq 0.01$ and $*** p \leq 0.001$, respectively, in Figure 4). In the listening task, response accuracy exhibited a subject-object asymmetry only in the ex situ condition ($*** p \leq 0.001$ and $*** p \leq 0.001$ in both instances, as shown in Figure 7), whereas mean ratings displayed an asymmetry only in the in situ condition ($** p \leq 0.01$ in Figure 8). This discrepancy may be due to the tasks not being sensitive enough to reveal significant differences in each condition. The direction that this asymmetry took was systematically the same across the reading task and the listening task: it always pointed towards a subject advantage in ex situ questions and towards an object advantage in in situ questions.

In addition, the size of the effect (i.e., subject-object asymmetry) in *quel*-questions was larger than in *qui*-questions, but only in *in situ* conditions (see Figures 3 and 4 for the reading task, and Figure 8 for the listening task). In *ex situ* conditions, *qui*-questions exhibited a larger effect (see Figures 3 and 6 for the reading task).

In addition to responses and ratings, we measured response time and reading/listening time, but did not observe any overall significant effects, except for a subject advantage indicated by shorter reading times when processing *ex situ qui*-questions in the reading task. Our design and tasks were probably not sensitive enough to detect time differences, as other tasks (e.g., eye-tracking, self-paced reading, maze tasks) do.

3 Discussion

The results of our two experiments are identical except for the modality of administration (written vs. oral). There was also a difference in effect size between the two tasks, again depending on *ex situ* and *in situ* conditions. The effect in the listening task was more significant than in the reading task in *ex situ* conditions, whereas in *in situ* conditions, it was the reading task where the effect was more significant. We believe that these inconsistencies are related to the relatively lower accuracy and ratings of *ex situ* questions in both tasks. *Ex situ* questions were less natural than *in situ* questions in our design, which affected participants' comprehension and led to results that deviated from our expectations.

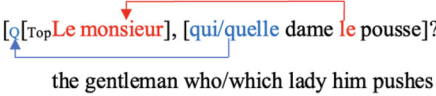
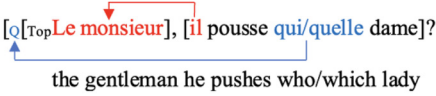
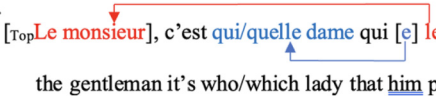
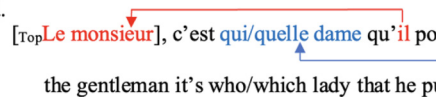
The first important, although partial, result we obtained, is that we found an asymmetry between *qui*-questions and *quel*-questions, which is an important indication of a relativized minimality effect (as discussed in Section 1.2). Specifically, the subject-object asymmetry in *quel*-questions is more significant than in *qui*-questions, and *quel*-questions are overall more difficult to process than *qui*-questions. This suggests that understanding *wh*-*in situ* questions indeed requires the movement of the *wh*-word and this is disrupted by intervening elements that share morphosyntactic features with it. The lexical *wh*-word *quel* 'which' shares more features with a lexical intervener ([+D], [+NP]) than the bare *wh*-word *qui* 'who' (only [+D]): see above the discussion in Section 1.2. As a result, *quel*-questions are harder to comprehend than *qui*-questions. This observation validates the movement analysis of *wh*-questions, since unselective binding does not predict this selective effect. Interestingly, this difference was significant only

for *wh*-in situ, while we were not able to detect any difference in the *wh*-ex situ conditions. Notice however that this relativized minimality effect is well-attested and has been found in *wh*-movement in French (cf. recently Bentea 2016). There might be many reasons for why we did not find it here, including the fact that our measures are basically offline and shallow.

In both tasks, a subject advantage in ex situ questions emerged in response accuracy, while an object advantage in in situ questions was displayed by mean ratings. Why ex situ questions display a *subject advantage* while in situ questions display an *object advantage* in our results? As we discussed in Section 2, our initial expectation was to find the same subject advantage in both strategies, as a reflex of the same underlying dependency between the scope position of the *wh*-element and its scope position. The results we found for ex situ *wh*-questions were in line with these expectations: as predicted, subject questions were comprehended significantly more accurately than object questions. But what about the apparently opposite result we obtained with in situ questions? This seems at first sight completely at odds with the hypothesis that the same underlying derivation holds for both strategies of questions.

Remember, however, that in our quest for a complete and unambiguous paradigm we decided to insert in all the questions in the experiments a clitic left dislocation structure. We believe that the asymmetry we found in our results, opposing a subject advantage in *wh*-movement questions and an object advantage in *wh*-in situ questions, can be explained in a unitary way if we assume that a movement dependency holds in both *wh*-questions strategies but that the landing site is different. As discussed in Section 1, Faure and Palasis (2021) argue on the basis of descriptive observations on the distribution of the two strategies in Colloquial French, that *wh*-ex situ is associated with a lower position, that they label Exclusivity.

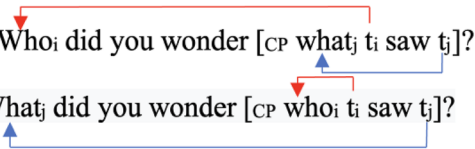
Although analyses of topic-pronoun dependencies differ (we can grossly divide them between “movement analyses”, see Kayne 1975, 1989, 1995; Rizzi 1997; Sportiche 1989, 1990; and “base generation analyses”, see Anagnostopoulou 1997; De Cat 2007; Demirdache 1992; Raposo 1998), researchers agree that a dependency needs to be established when interpreting pronouns associated with a topicalized item. Rizzi (1997) and the literature stemming from this seminal paper have established that the Topic position is higher than the position where the *wh*-elements move overtly. But what about *wh*-in situ? If its (covert) landing site in the tree is higher than the landing site of overt *wh*-movement, as argued by Faure and Palasis (2021), it is possible that this position is higher than Topic. If this is the case, we can assume that the two dependencies (topic, in red below, and *wh*-, in blue below) interact differently. Consider (21).

- (21) a.  (*wh*-subject in situ questions)
the gentleman who/which lady him pushes
- b.  (*wh*-object in situ questions)
the gentleman he pushes who/which lady
- c.  (*wh*-subject ex situ questions)
the gentleman it's who/which lady that him pushes
- d.  (*wh*-object ex situ questions)
the gentleman it's who/which lady that he pushes

As can be seen in (21), *wh*-dependencies and topic-pronoun dependencies interact differently in the different conditions if the movement associated with *wh*-ex situ and *wh*-in situ differ. In ex situ subject questions, illustrated in (21c), the landing site of *wh*-movement is uncontroversially on the right of the Topic P. We shall assume that this lower position is associated with the Exclusivity/contrastive feature discovered by Faure and Palasis (2021). As a result, in ex situ subject questions the topic-pronoun dependency and the *wh*-dependency have an inclusive relation, namely the *wh*-dependency is fully included in the topic-pronoun dependency, as illustrated in (21c). In ex situ object questions, illustrated in (21d), the topic-pronoun dependency and the *wh*-dependency have instead an intersective relation, i.e., they cross each other. In other words, within the ex situ strategy, subject questions display nested dependencies while object questions exhibit crossing dependencies.

In in situ *wh*-questions, on the other hand, the *wh*-movement does not land in this lower position because no Exclusivity feature is present, and the landing site, arguably, is higher (corresponding to the scope position). If we adopt the hypothesis that in situ *wh*-questions involve covert movement to a higher position, we can interpret our results as a similar pattern. The *wh*-dependency *intersects* with the topic-pronoun dependency in subject questions (21a), while it *includes* the topic-pronoun dependency in object questions. This entails that subject ex situ questions and object in situ questions do have a strong property in common after all: they both display nested dependencies.

Now, it is well known from the sentence processing literature that nested dependencies are relatively easier to process compared to crossing dependencies (Fodor 1978; Frazier and Fodor 1978; Pickering and Barry 1991; Rochemont and Culicover 1990). This is illustrated for example by the contrast in (22).

- (22) a. * Who_i did you wonder [CP what_j t_i saw t_j]?
 b. What_j did you wonder [CP who_i t_i saw t_j]?


The pattern of data that we found can be explained along these same lines if we assume, as in (21), that all questions involve the same kind of *wh*-dependencies, only differing in their landing site (and associated feature): subject ex situ questions, and object in situ questions, involving nested dependencies, are better comprehended and/or accepted than object ex situ and subject in situ questions, which involve crossing dependencies. The only specificity of our data concerns the status of the contrast at stake: while (22) display a grammaticality contrast, because the dependencies interacting there are of the same type involving the exact same (WH) feature, our data display a simple processing asymmetry, because the dependencies interacting here involve two different A-bar features (TOP and Q), hence only partially interacting.⁸

As for the grammatical basis for this processing effect, Pesetsky (1982) reformulated this bias as a syntactic constraint on A'-dependencies:

- (23) Path Containment Condition (PCC) Pesetsky (1982).⁹
 If two paths overlap, one must contain the other.¹⁰

Crucially, in the conditions that appear to be easier in our results one path contains the other, while the two paths of topic and *wh*-movement intersect in the conditions that are disfavored.

Notice that a very similar interaction between topics and *wh*-movement questions have been observed in French in recent experimental work by Barbosa and De Cat (2019), and explained along the same lines: they observe a disruption in acceptability each time a topic dependency and a *wh*-dependency cross (as in 24b),

⁸ The same difference holds between grammatical and processing effects of featural relativized minimality phenomena as discussed above (Section 1).

⁹ Where a path is, roughly, the set of nodes that connect the lexical projections dominating the foot and the head of a chain.

¹⁰ The PCC can also be reformulated as an intervention effect (Relativized Minimality, Rizzi 1990), if complete chains do not act as interveners the way chain links do.

while no such disruption holds when the same dependencies are in a nested configuration (as in 24a).

- (24) a. Voici [les médailles] que, [les athlètes], ils sont fiers d'avoir remportées [t]
 PRES. the medals that the athletes they are proud to=have.INF won
 'Here are the medals that the athletes are proud to have won'
- b. * Voici [les athlètes] qui, [les médailles d'or], [t] les ont remportées
 PRES. the athletes that, the medals of gold, they have won
 'Here are the athletes who have won the gold medal'

(adapted from Barbosa and De Cat 2019: 28)

Interestingly, however, they observe that this restriction holds only in those configurations in which the Topic appears inside the scope of the *wh*-phrases, as in (24), and not in cases where the Topic precedes the *wh*-constituent, as in (25).

- (25) a. Et Jean, quel livre a-t-il acheté [t]
 and Jean, which book has=he bought
- b. Et Jean, tu sais quel livre il a acheté [t]
 and Jean, you know which book he as bought

(adapted from Barbosa and De Cat 2019: 43)

Barbosa and De Cat (2019) discuss in great length on how to derive the contrast between (24) and (25) from syntactic constraints: "On the one hand, the fact that CLLD is not subject to the no-crossing constraint is consistent with non-movement analyses of the construction (Anagnostopoulou 1997; De Cat 2007; Demirdache 1992; Raposo 1998). But if the clitic left dislocation structure (CLLD) doesn't involve movement, one cannot appeal to a constraint on movement to rule out the cases" like (24). They go on discussing the puzzle, and tentatively conclude that the asymmetry might be due to a difference in the processing cost of integrating topics and *wh*-phrases.

Be as it may, what seems to be the case is that we observed the kind of generalized effect that would be expected if a syntactic constraint is at play here that Barbosa and De Cat (2019) failed to capture in their experiment. In our results, we observe a degradation associated to crossing dependencies, both when topic precedes WH and when WH precedes Topic.

Whether this is due to a difference in the data set (we also included *wh*-in situ items, while they only considered *wh*-ex situ questions) or to a difference in the methodology (they only recorded acceptability data, while we also measured comprehension accuracy) is an open issue.

4 Conclusions

This paper reports and discusses the results of two online experiments measuring the comprehension and acceptability of two *wh*-questions strategies available in French, the *in situ* strategy, where the *wh*-constituent is not moved to its scope position, and the *ex situ* strategy, where the *wh*-constituent does move to the edge of the clause. For reasons related to the need to avoid any ambiguity in interpretation and to be able to dispose of a complete paradigm for all the desired conditions (in situ vs. ex situ $\times$ subject vs. object $\times$ qui vs. quel), we systematically enriched the *wh*-question with a clitic left dislocation structure (as in *Le monsieur, qui le pousse?*).

Our results are superficially divergent for *wh*-ex situ and *wh*-in situ. While we observe, as predicted (given what is known about the processing of A-bar dependencies) a subject advantage in the former, we find an object advantage in the latter. We discuss these surprising results, arguing that they eventually indirectly confirm that both in situ and ex situ questions involve the same type of covert dependency. This dependency systematically interacts with the Topic-pronoun dependency also present in our items, yielding different results because the landing site associated to *wh*-in situ and *wh*-ex situ differ, as argued by Faure and Palasis (2021). As a matter of fact, the superficial dissymmetry between in situ and ex situ questions that we found provides a strong confirmation that covert and overt dependencies are similarly constrained.

Further investigations are needed in typologically unrelated languages that display *wh*-in situ. One priority is to examine *wh*-in situ questions in Mandarin Chinese, with or without an interaction with topicalization. If French and Chinese *wh*-in situ operate similarly in terms of covert dependency establishment (cf. Huang 1982), a subject advantage is expected when Topic is absent, and an object advantage when it is present. Additionally, it is necessary to explore how *wh*-fronting interacts with topic in Mandarin Chinese to determine whether fronted questions in this language are associated with the same lower position. Notably, Cheung (2008) argues that *wh*-fronting marks contrastive focus in Chinese, a claim closely related to the Exclusivity feature that Faure and Palasis (2021) associate with *wh*-fronting in French.

Another important direction is to extend the study to Japanese, which has been argued not to exhibit such a dependency in *wh*-in situ (Watanabe 1992). If this hypothesis is correct, no interaction with topicalization is expected.

Finally, future research should consider employing online eye-tracking methods, as suggested by findings from Hauser et al. (2021) and Fornasiero et al. (2024). Hauser et al. (2021) demonstrated that offline methods were not sufficiently sensitive to detect asymmetries between overt and covert dependencies in relative

clauses, while Fornasiero et al. (2024) observed a subject advantage in bilinguals using eye-tracking. These results highlight the potential of online methods to reveal dependency-related asymmetries that may be obscured in offline approaches, supporting their adoption in future investigations.

Acknowledgments: We would like to thank all the participants in the study, all the reviewers and the editor for their valuable comments. This research was supported by China Scholarship Council and Laboratoire de Linguistique Formelle (LLF). Authors' contributions: CD and RL conceived and designed the study; all authors prepared the experimental materials. RL implemented the tasks; RL and DG analyzed the data; all authors discussed the results; RL and CD drafted the manuscript; DG revised the first draft. All authors approved the manuscript before submission.

References

- Adli, Aria. 2006. French wh-in-situ questions and syntactic optionality: Evidence from three data types. *Zeitschrift für Sprachwissenschaft* 25(2). 163–203.
- Anagnostopoulou, Elena. 1997. Clitic left dislocation and contrastive left dislocation. In Elena Anagnostopoulou, Henk van Riemsdijk & Frans Zwarts (eds.), *Materials on left dislocation*, 151–192. Amsterdam: John Benjamins Publishing Company.
- Abeillé, Anne & Danièle Godard (dir.). 2023. *Grande Grammaire du Français*. Arles, Actes Sud, Paris: Imprimerie nationale éditions.
- Aoshima, Sachiko, Colin Phillips & Weinberg Amy. 2004. Processing filler-gap dependencies in a head-final language. *Journal of Memory and Language* 51(1). 23–54.
- Avrutin, Sergey. 2000. Comprehension of discourse-linked and non-discourse-linked questions by children and Broca's aphasics. In Yosef Grodzinsky, Lewis P. Shapiro & David Swinney (eds.), *Language and the brain: Representation and processing*, 295–313. San Diego, CA: Academic Press.
- Avrutin, Sergey. 2006. Weak syntax. In Yosef Grodzinsky & Katrin Amunts (eds.), *Broca's region*, 49–62. Oxford: Oxford University Press.
- Barbosa, Pilar & Cécile De Cat. 2019. Intervention effects in wh-chains: The combined effect of syntax and processing. *Glossa: a Journal of General Linguistics* 4(1). 127. 1–26.
- Bates, Douglas, Martin Mächler, Ben Bolker & Steve Walker. 2015. Fitting linear mixed-effects models using lme4. *Journal of Statistical Software* 67(1). 1–48.
- Baunaz, Lena. 2011. *The Grammar of French quantification*. Studies in Natural Language and Linguistic Theory 83. Dordrecht: Springer Netherlands.
- Bayer, Josef & Lisa Lai-Shen Cheng. 2017. Wh-in-Situ. In Martin Everaert & Henk C. van Riemsdijk (eds.), *The wiley blackwell companion to syntax*, 2nd edn., 1–44. John Wiley & Sons, Inc.
- Belletti, Adriana. 2008. The CP of clefts. *Rivista di Grammatica Generativa* 33. 191–204.
- Belletti, Adriana. 2009. Answering strategies: New information subjects and the nature of clefts. In Adriana Belletti (ed.), *Structures and strategies*, 242–265. London: Routledge.

- Belletti, Adriana. 2015. The focus map of clefts: Extraposition and Predication. In Uri Shlonsky. (ed), *Beyond functional sequence: The cartography of syntactic structures*, Vol. 10, 42–59. New York: Oxford University Press.
- Bentea, Anamaria. 2016. *Intervention effects in language acquisition: The comprehension of A-bar dependencies in French and Romanian*. Geneva: University of Geneva doctoral thesis.
- Boeckx, Cedric. 1999. Decomposing French questions. *University of Pennsylvania Working Papers in Linguistics* 6(1). 69–80.
- Bošković, Željko. 2000. Sometimes in [Spec, CP], sometimes in situ. In Roger Martin, David Michaels & Juan Uriagereka (eds.), *Step by step: Essays on minimalist syntax in honor of Howard Lasnik*, 53–87. Cambridge, Mass: MIT Press.
- Chang, Lisa. 1997. *Wh-in-situ phenomena in French*. Vancouver, BC: University of British Columbia MA thesis.
- Cheng, Lisa & Johan Rooryck. 2000. Licensing wh-in-situ. *Syntax* 3. 1–19.
- Cheung, Candice Chi-Hang. 2008. *Wh-fronting in Chinese*. Los Angeles: University of Southern California PhD dissertation.
- Chomsky, Noam. 1995. *The minimalist program*. Cambridge, Mass: MIT press.
- Christensen, Rune H. B. 2019. Ordinal – regression models for ordinal data. *R Package Version* 2019. 12–10.
- Deevy, Patricia & Laurence B. Leonard. 2004. The comprehension of Wh-questions in children with specific language impairment. *Journal of Speech, Language, and Hearing Research* 47(4). 802–815.
- De Cat, Cécile. 2007. *French dislocation: Interpretation, syntax, acquisition*, (Oxford Studies in Theoretical Linguistics 17). Oxford: Oxford University Press.
- Demirdache, Hamida. 1992. *Resumptive chains and restrictive relative clauses, appositives and dislocation structures*. Cambridge, Mass: MITWPL.
- Destrue, Emilie. 2013. *The French C'est-cleft: Empirical studies of its meaning and use*. Austin, Texas: University of Texas PhD dissertation.
- Diessel, Holger & Michael Tomasello. 2005. A new look at the acquisition of relative clauses. *Language* 81. 1–25.
- Dryer, Matthew S. 2013. Position of interrogative phrases in content questions. In Matthew S. Dryer & Martin Haspelmath (eds.), *The world atlas of language structures online* (<http://wals.info/chapter/93>). Leipzig: Max Planck Institute for Evolutionary Anthropology.
- Erlwine, Michael Yoshitaka. 2020. Anti-locality and subject extraction. *Glossa: a Journal of General Linguistics* 5(1). 84. 1–38.
- Faure, Richard & Katerina Palasis. 2021. Exclusivity! Wh-fronting is not optional wh-movement in Colloquial French. *Natural Language & Linguistic Theory* 39(1). 57–95.
- Fodor, Janet Dean. 1978. Parsing strategies and constraints on transformations. *Linguistic Inquiry* 9(3). 427–473.
- Fornasiero, Elena, Charlotte Hauser & Chiara Branchini. 2024. The subject advantage in LIS internally headed relative clauses: An eye-tracking study. *Bilingualism: Language and Cognition*. Published online 2024. 1–14. <https://doi.org/10.1017/s1366728924000415>.
- Frazier, Lyn & Janet Dean Fodor. 1978. The sausage machine: A new two-stage model of the parser. *Cognition* 6. 291–325.
- Friedmann, Naama Adriana Belletti & Luigi Rizzi. 2009. Relativized relatives: Types of intervention in the acquisition of A-bar dependencies. *Lingua* 119(1). 67–88.
- Guasti, Maria Teresa. 2004. *Language acquisition: The growth of grammar*. Cambridge, Mass: MIT Press.
- Hauser, Charlotte, Giorgia Zorzi, Valentina Aristodemo, Beatrice Giustolisi, Doriane Gras, Rita Sala, Jordina Sánchez Amat, Carlo Cecchetto & Caterina Donati. 2021. Asymmetries in relative clause comprehension in three European sign languages. *Glossa: a Journal of General Linguistics* 6(1). 72.

- Hauser, Charlotte, Valentina Aristodemo & Caterina Donati. 2023. A subject advantage in covert dependencies: The case of *wh*-questions comprehension in French Sign Language. *Syntax* 26. 280–310.
- Huang, C.-T. James. 1982. Move WH in a language without movement. *The Linguistic Review* 1. 369–416.
- Kayne, Richard. 1975. *French syntax. The transformational cycle*. Cambridge, Mass: MIT Press.
- Kayne, Richard & Jean-Yves Pollock. 1978. Stylistic inversion, successive cyclicity, and move NP in French. *Linguistic Inquiry* 9(4). 595–621.
- Kayne, Richard. 1989. Null subjects and clitic climbing. In Osvaldo Jaeggli & Ken Safir (eds.), *The null subject parameter*, 239–261. Dordrecht: Kluwer.
- Kayne, Richard. 1995. *The Antisymmetry of syntax*. Cambridge, Mass: MIT Press.
- Kuznetsova, Alexandra, Per B. Brockhoff & Rune H. B. Christensen. 2017. lmerTest package: Tests in linear mixed effects models. *Journal of Statistical Software* 82(13). 1–26.
- Lau, Elaine & Nozomi Tanaka. 2021. The subject advantage in relative clauses: A review. *Glossa: a Journal of General Linguistics* 6(1). 1–34.
- Levy, Hagar & Naama Friedmann. 2009. Treatment of syntactic movement in syntactic SLI: A case study. *First Language* 29(1). 15–49.
- Li, Yen-Hui Audrey. 1992. Indefinite wh in Mandarin Chinese. *Journal of East Asian Linguistics* 1(2). 125–155.
- Mathieu, Eric. 1999. WH in situ and the intervention effect. In Corinne Iten & Ad Neeleman (eds.), *UCL working papers in linguistics*, 441–472. London: University College.
- Miyamoto, Edson T. & Shoichi Takahashi. 2002. The processing of wh-phrases in Japanese. *Scientific approaches to language* 1. 133–172.
- Nishigauchi, Taisuke. 1990. *Quantification in the theory of grammar*. Dordrecht: Kluwer.
- Pablos, Leticia, Doetjes Jenny & Lisa L.-S. Cheng. 2018. Backward dependencies and in-situ wh-questions as test cases on how to approach experimental linguistics research that pursues theoretical linguistics questions. *Frontiers in Psychology* 8. <https://doi.org/10.3389/fpsyg.2017.02237>.
- Pesetsky, David. 1982. *Paths and categories*. Cambridge, Mass: MIT PhD dissertation.
- Pesetsky, David. 1987. Wh-in-situ: Movement and unselective binding. In Eric Reuland & Alice G. B. ter Meulen (eds.), *The Representation of (In)Definiteness*, 98–129. Cambridge, Mass: MIT Press.
- Pickering, Martin & Guy Barry. 1991. Sentence processing without empty categories. *Language & Cognitive Processes* 6(3). 229–259.
- Raposo, Eduardo. 1998. Definite/zero alternations in Portuguese: Towards a unification of topic constructions. In Bernard Tranel, Armin Schwegler & Myriam Uribe-Etxebarria (eds.), *Romance linguistics: Theoretical perspectives*, 197–212. Amsterdam: John Benjamins.
- Richards, Norvin. 2000. An island effect in Japanese. *Journal of East Asian Linguistics* 9. 187–205.
- Richy, Célia & Heather Burnett. 2020. Jean does the dishes while marie fixes the car: A qualitative and quantitative study of social gender in French syntax articles. *Journal of French Language Studies* 30(1). 47–72.
- Rizzi, Luigi. 1990. *Relativized minimality*. Cambridge, Mass: MIT Press.
- Rizzi, Luigi. 1997. The fine structure of the left periphery. In Liliane Haegeman (ed.), *Elements of grammar: Handbook of generative grammar*, 281–337. Dordrecht: Kluwer.
- Rizzi, Luigi. & Ur Shlonsky. 2007. Strategies of subject extraction. In Uli Sauerland & Hans-Martin Gärtner (eds.), *Interfaces + recursion = language? Chomsky's Minimalism and the view from syntax semantics*, 115–160. Berlin: Walter de Gruyter.
- Rochemont, Michael Shaun & Peter William Culicover. 1990. *English focus constructions and the theory of grammar* (Cambridge Studies in Linguistics 52). Cambridge: Cambridge University Press.

- Rooryck, Johan. 2000. A unified analysis of French interrogative and complementizer *qui/que*. In *Configurations of sentential complementation*, 223–246. London: Routledge.
- Sheppard, Susan M., Matthew Walenski, Tricia Love & Lewis P. Shapiro. 2015. The auditory comprehension of *wh*-questions in aphasia: Support for the intervener hypothesis. *Journal of Speech, Language, and Hearing Research* 58(3), 781–797.
- Sportiche, Dominique. 1989. Le Mouvement Syntaxique. Contraintes et Paramètres. *Langages* 95, 35–80.
- Sportiche, Dominique. 1990. *Movement, agreement and case*. Los Angeles, CA: University of California, Los Angeles manuscript.
- Sprouse, Jon & Diogo Almeida. 2017. Design sensitivity and statistical power in acceptability judgment experiments. *Glossa: A Journal of General Linguistics* 2(1), 14, 1–32.
- Starke, Michal. 2001. *Move dissolves into merge: A theory of locality*. Geneva: University of Geneva doctoral thesis.
- Stromswold, Karin. 1995. The acquisition of subject and object *wh*-questions. *Language Acquisition* 4(1-2), 5–48.
- Taraldsen, Knut Tarald. 2001. Subject extraction, the distribution of expletives and stylistic inversion. In Aafke Hulk & Jean-Yves Pollock (eds.), *Subject inversion in Romance and the theory of universal grammar*, Vol. 6, 163–182. New York: Oxford University Press.
- Tracy, Rosemarie. 1994. Raising questions: Formal and functional aspects of the acquisition of *wh*-questions in German. In Rosemarie Tracy & Elisabeth Lattey (eds.), *How tolerant is universal grammar?* 1–34. Tübingen, Germany: Max Niemeyer Verla.
- Ueno, Michiko & Robert Kluender. 2009. On the processing of Japanese *wh*-questions: An ERP study. *Brain Research* 1290, 63–90.
- van der Lely, Heather K. J. & John Battell. 2003. *Wh*-movement in children with grammatical SLI: A test of the RDDR hypothesis. *Language* 79(1), 153–181.
- van Gompel, Roger P. G. 2013. *Sentence processing*. Hove, the United Kingdom: Psychology Press.
- Villata, Sandra, Luigi Rizzi & Julie Franck. 2016. Intervention effects and relativized minimality: New experimental evidence from graded judgments. *Lingua* 179, 76–96.
- Watanabe, Akira. 1992. Subjacency and S-structure movement of *wh*-in-situ. *Journal of East Asian Linguistics* 1, 255–291.
- Xiang, Ming, Brian Dillon, Matthew Wagers, Fengqin Liu & Taomei Guo. 2014. Processing covert dependencies: An SAT study on Mandarin *wh*-in-situ questions. *Journal of East Asian Linguistics* 23, 207–232.
- Xiang, Ming, Shuai Wang & Yi Cui. 2015. Constructing covert dependencies – the case of Mandarin *wh*-in-situ dependency. *Journal of Memory and Language* 84, 139–166.