



BEYOND DOMINANCE: GENDER AND INTERRUPTION IN ONLINE COMMUNICATION

Hanane Salif¹ⁱ,

Mina Afkir²

¹Faculty of Letters & Humanities Ain Chock,
Hassan II University of Casablanca,
Morocco

²Faculty of Letters & Humanities Ain Chock,
Hassan II University of Casablanca,
Morocco

Abstract:

Gendered patterns of conversational interruption have been a long-debated topic in sociolinguistics. It has been divided into two broad perspectives: some treat interruption chiefly as a marker of dominance, in which one speaker seizes the floor at another's expense, while others argue that interruption can serve cooperative or epistemic functions unrelated to control, with the same surface form, overlapping or simultaneous speech, serving markedly different interactional functions depending on context. This study examines gendered interruption in two Moroccan samples, comparing face-to-face focus-group discussion with parallel WhatsApp group messaging drawn from another peer group at a different institution. It applies a four-category coding framework, distinguishing competitive, corrective, and supportive interruption in face-to-face talk from a WhatsApp-specific sequential thread-bypass category, and asks whether the traditional dominance account adequately explains the patterns observed. A combined qualitative and quantitative approach was used, drawing on coded focus-group transcripts and a message-by-message coding of the WhatsApp corpus across both media. Findings show that female participants account for a higher share of interruptive and bypass events in the competitive and corrective functions across both settings, but not in the supportive function, where face-to-face overlap skews male while its WhatsApp analog skews female, the one function whose direction differs across media; none of these contrasts reach conventional statistical significance at this sample size and are reported descriptively throughout. Because the two media make different aspects of interaction observable, the study treats cross-medium comparison as a problem of operationalization rather than assuming that interruption constitutes an equivalent countable unit in both settings. Once this is accounted for, competitive floor-competition occurs markedly less often on WhatsApp than face-to-face, even though it still skews female, and corrective interruption, while female-led in both settings, occurs several times more often face-to-face than on WhatsApp, suggesting that direct floor competition specifically

ⁱ Correspondence: email salifhanane6@gmail.com

does not transfer to asynchronous messaging in the form it takes face-to-face. The study complicates the traditional dominance account without simply reversing it, and offers a medium-sensitive framework for future research on gendered interruption in offline and digital settings.

Keywords: interruption, gender, conversation analysis, computer-mediated communication, disrupted turn adjacency

1. Introduction

Conversation is organized by a systematic set of norms governing who speaks next and when (Sacks, Schegloff, and Jefferson, 1974), and interruption marks the point at which that system is overridden by force rather than by orderly allocation. An interruption is never only a structural event. It also raises the question of whose challenges are treated as legitimate and what kinds of knowledge count as authoritative within a given interactional community. Since Zimmerman and West's (1975) foundational study, interruption has become one of the most heavily researched topics in the sociolinguistics of gender, precisely because it sits at the junction between the micro-level machinery of conversation and the macro-level structures of social power that conversation reproduces, contests, and transforms. Despite this prominence, the concept remains contested, and researchers continue to disagree over its definition, measurement, and interpretation.

This paper examines interruptive events in two Moroccan university datasets, face-to-face focus-group discussion and parallel WhatsApp group messaging, drawn from separate peer groups at two different institutions with no participant overlap. It applies a coding framework distinguishing three interactional functions (competitive, corrective, and supportive), realized directly in face-to-face talk, from a WhatsApp-specific sequential thread-bypass mechanism that bundles those same three functions under a single medium-specific category. This framework is motivated by James and Clarke's (1993) critique that a single surface form can serve multiple discourse functions. It is designed to capture qualitative variation among interruptive events rather than treating all instances of simultaneous speech or floor competition as equivalent. The central argument is that the distribution and character of interruptive events in this corpus complicate the traditional dominance account without simply reversing it. Aggregate frequency and interactional function point in different directions, and only the latter supports a coherent theoretical claim.

A central methodological problem addressed in this study is that interruption is not equally observable across interactional media. In face-to-face conversation, interruption can be identified through the temporal overlap of speech and the displacement of an ongoing turn. In WhatsApp interaction, by contrast, participants do not have access to the same real-time unfolding floor: contributions appear as discrete messages, may be fragmented across several posts, and may respond to an earlier message only after several intervening exchanges. Consequently, the problem is not simply whether an online practice can be labeled an interruption, but whether interruption can be operationalized and counted in the same way once the organization of interaction itself has changed. The present study therefore treats the

comparison between the two datasets as a test of how interruption is measured across media, rather than assuming structural equivalence between face-to-face interruption and WhatsApp interaction.

Three questions organize the analysis that follows.

RQ1: How are gendered interruptive practices distributed across competitive, corrective, and supportive functions in face-to-face and WhatsApp interaction?

RQ2: To what extent can interruption be operationalized and quantitatively compared across face-to-face and WhatsApp interaction, given their different turn-taking and sequencing structures?

RQ3: What does this comparison reveal about the relationship between gender, interactional function, and medium?

Together, these questions link the paper's quantitative comparison to its interpretive account of what interruption accomplishes in each medium.

2. Literature Review

The foundational claim in the gender-and-interruption literature comes from Zimmerman and West's (1975) analysis of dyadic conversations, which found that men performed 96% of interruptions in mixed-gender pairs and interpreted this pattern as evidence of conversational dominance. West and Zimmerman (1983) extended this claim in a further study of interruptions between unacquainted cross-sex speakers, reinforcing the dominance account as the field's leading framework. This account was enormously influential but attracted sustained critique on several fronts. James and Clarke (1993) argued that studies varied widely in how they operationally defined interruption, and that equating interruption with dominance overlooked how the same surface form, overlapping or simultaneous speech, can serve radically different functions depending on sequential and relational context.

Tannen (1994) distinguished cooperative overlaps, which express enthusiasm and collaborative floor-building, from competitive interruptions, arguing that high-involvement conversational styles treat simultaneous speech as a normal and positive feature of engaged interaction rather than as aggression. Goldberg (1990) extended this critique further, proposing that interruptive acts be classified not by whether they seize the floor but by their relational function, as relationally neutral, power-oriented, or rapport-oriented acts. This typology treats floor-claiming as one possible function among several rather than as interruption's defining feature. Across these critiques, the common thread is that dominance is too blunt an instrument to capture what interruption actually accomplishes in different sequential and relational contexts.

Even so, gender-linked interruption has not disappeared from view: Miller and Sutherland (2023) found that women members of the United States Congress were interrupted significantly more often than their male colleagues during committee hearings. A working paper by Tong (2025) reports a similar pattern in United States Supreme Court oral arguments, where interruptions directed at female advocates carried a more negative emotional tone than those directed at male advocates. Together, these critiques and these more recent findings suggest that the relationship between gender and interruption remains context-dependent rather than

settled, shaped by function and setting nearly five decades after Zimmerman and West's original study.

The concept of epistemic authority (Heritage, 2012a, 2012b; Heritage & Raymond, 2005) offers a theoretically richer framework than simple dominance for interpreting interruption. Epistemic authority is the socially distributed right to be treated as a knowledgeable and credible source of information about a given topic. In conversation, this right is not assigned by formal institutional role but is continuously claimed, contested, and negotiated through the sequential structure of interaction itself. A speaker who interrupts to correct another's factual claim is asserting the right to have their knowledge count as more reliable than the interrupted speaker's, a distinct interactional move from interrupting simply to seize the floor. This same negotiation of epistemic standing remains an active concern in current conversation-analytic research: Caronia (2024) shows that parents can undermine a teacher's epistemic authority during parent-teacher conferences by invoking outside experts, extending the concept well beyond the institutional contexts in which it was first developed.

Edelsky's (1981) concept of the collaborative floor complements this framework, describing a conversational mode in which participants contribute simultaneously toward a shared narrative, in contrast with the single floor, one speaker at a time, described by Sacks, Schegloff, and Jefferson (1974). In this collaborative mode, overlapping speech is generative rather than displacing, a pattern Edelsky documented across a range of same-gender and mixed-gender meetings, and one that Coates (1996) later extended in her account of how women jointly construct talk.

A further methodological question arises once interruption is compared across face-to-face and digital data: can a construct built for real-time, co-present talk be applied to asynchronous messaging at all? Herring's (1999) concept of disrupted turn adjacency describes how, in asynchronous computer-mediated communication, messages are composed without real-time coordination, so that a reply's position in a thread does not reliably indicate what it is responding to. A participant may reply to a message posted several exchanges earlier, producing a response that is non-adjacent in the linear sequence. Garcia and Jacobs (1999) extend this insight into the turn-taking system itself, showing that in quasi-synchronous computer-mediated communication, participants cannot observe a message being composed, so that every message effectively constitutes its own transition-relevance place.

Read together, these findings justify treating an out-of-sequence WhatsApp reply as a structural analog to face-to-face interruption rather than as ordinary asynchronous responding, since both involve a participant asserting a claim on the conversational floor in a way that departs from what the preceding talk projected, even though the two media accomplish this through different mechanisms. This methodological choice is further supported by Colom (2022), who found that WhatsApp's integration into participants' everyday communication routines gives it a level of ecological validity that purpose-built research platforms cannot easily replicate.

Research on interruption in Arabic and Moroccan discourse contexts remains comparatively sparse. Holes (2004) documents the diglossic character of Arabic sociolinguistic life, in which colloquial vernaculars such as Moroccan Darija operate under stylistic and pragmatic norms distinct from both the standard variety and the English-language contexts in

which most interruption research has been conducted. Caubet's (2004) collection of interviews with Maghrebi artists documents how Darija, often mixed with French and Berber, functions as a vehicle of irreverent humor, self-mockery, and subversive expressiveness in popular Maghrebi culture, though it does not itself address conversational turn-taking or interruption. Taken together, this literature suggests that the functional interpretation of overlapping speech in a Moroccan peer group may differ from the norms assumed in much of the interruption literature, which was developed primarily in Northern European and North American contexts, even though a dedicated empirical study of interruption or overlap norms in Moroccan Darija conversation has yet to be conducted.

3. Material and Methods

The face-to-face data were collected from three discussion groups, referred to here as Group 1, Group 2, and Group 3, drawn from the same university peer network. The groups differ in participants' relational history at the time of recording: Group 1 members had been close friends for several years, Group 2 members had a more moderate degree of acquaintance, and Group 3 members had known one another for only a few weeks, having recently formed as a peer group. This difference in relational history is treated as a contextual variable in the group-level analysis reported in Section 4.1.

The three groups are symmetrical in gender composition, each comprising two female and two male participants (twelve participants in total: six female, six male), aged eighteen to twenty-four and enrolled in the same English Studies program at a Moroccan public university. The WhatsApp corpus is drawn from an entirely separate population, a class-group chat at a different Moroccan higher education institution; no individual appears in both datasets, so the cross-medium comparisons in Section 4.5 contrast two media rather than tracking the same speakers across media. Gender for WhatsApp senders was assigned from display name, profile photo, and gendered self-reference in context; senders whose gender could not be determined with reasonable confidence were excluded from the analysis.

The data comprise face-to-face focus-group discussions and parallel WhatsApp group messaging, collected from the two distinct populations described above. The coding framework distinguishes four categories of interruptive events, three defined by interactional function and one defined by medium. Competitive interruption occurs when a speaker takes a new turn or posts a message that displaces the current speaker's agenda before a natural completion point. Corrective interruption occurs when a speaker intervenes to introduce a factual correction, counter-example, or epistemic reframing of the current claim, and is defined by this epistemic function rather than by whether floor displacement occurs. Supportive overlap is cooperative simultaneous speech or messaging that completes, affirms, or elaborates the current contribution without displacing it. These three categories describe face-to-face interruptive events directly, either entry before a plausible completion point or genuinely simultaneous speech.

The fourth category, sequential thread-bypass, is specific to the WhatsApp data, where simultaneous speech is not structurally possible. A thread-bypass is a reply to a non-immediately preceding message that disrupts the linear sequential structure of the thread.

Following Herring (1999) and Garcia and Jacobs (1999), this out-of-sequence reply is treated as the medium-specific analog of interruption, where analog denotes a functionally comparable sequence-disruptive practice rather than a structurally equivalent event: the two are argued to serve a parallel interactional purpose, asserting a claim on the floor in a way that departs from what the preceding talk projected, without being produced by the same real-time mechanism (Section 4.5). Unlike the first three categories, thread-bypass is not itself a single interactional function: every bypass event was sub-coded as competitive, corrective, or supportive, using the same functional definitions applied to the face-to-face data. The four-category model should therefore be read as three functions realized directly in face-to-face talk, plus a fourth, medium-specific category that bundles those same three functions together under a single WhatsApp mechanism.

The coding procedure therefore treats measurement as a medium-sensitive problem rather than assuming that the same observable unit can be counted identically across the two datasets. In the face-to-face corpus, an interruptive event is anchored to the temporal organization of an unfolding turn: a next speaker enters before, or in overlap with, a projected completion point. In the WhatsApp corpus, no equivalent real-time overlap is observable; instead, the analysis identifies disruptions of sequential adjacency and then examines the interactional function performed by the bypass. This distinction is not a methodological footnote: the number of observable interruptive events is partly determined by how each medium packages interaction into analyzable units in the first place. The counts reported in Section 4 should therefore be understood not simply as behavioral frequencies, but as frequencies of events that are observable under the interactional conditions specific to each medium.

Thread-bypass events were identified through a message-by-message review of the fourteen WhatsApp segments, comprising 1,853 valid messages after excluding deleted, system-generated, and professor-authored messages. A message was coded as a bypass only when it could be matched, with reasonable confidence, to a specific earlier message that it answered, affirmed, or disputed. That earlier message was not the one immediately preceding it in the visible sequence. Each identified bypass was then sub-coded as competitive, corrective, or supportive using the same functional definitions applied to the face-to-face data, and sender gender was drawn directly from the corpus's anonymized sender coding. A secondary, more liberal coding pass additionally counted topic-revivals and repeated unanswered questions that lacked a single identifiable target message; both passes are reported in Section 4.4.

Because WhatsApp habitually divides a single contribution across several consecutive messages, a raw message count risks conflating platform-driven fragmentation with genuine behavioral difference. This threshold was derived from the corpus's own message-timing data rather than adopted by convention: across all fourteen segments, 94.7% of consecutive messages from the same sender arrived within ninety seconds of that sender's previous message, with the remaining messages separated by several minutes to multiple hours, a pattern consistent with a genuine return to an earlier point in the conversation rather than a continuation of a single fragmented turn.

Consecutive messages from the same sender arriving within this ninety-second window were therefore grouped into a single functional turn unless they clearly performed separate

interactional jobs; applying this rule to the WhatsApp corpus collapses its 1,853 raw messages into 1,395 functional turns, a 24.7 percent reduction. This threshold's effect on the turn count is not an artifact of the specific cutoff chosen: an independent, purely timestamp-based re-merge of the same message data run at sixty and one-hundred-twenty seconds yields reductions of 24.0% and 24.8% respectively, closely bracketing the reported 24.7% reduction at ninety seconds, so the turn-count reduction is stable across a reasonable range of merge windows. Cross-medium comparisons in Section 4.5 are computed against total word count rather than raw message or turn count, for the same reason.

Consent and disclosure followed different procedures for the two datasets, reflecting their different collection contexts. Face-to-face participants were informed before recording that their speech would be recorded for academic research and that they were free to withdraw at any time; the study's specific focus on gendered interaction was withheld until after recording, following a deferred-debriefing protocol consistent with British Association for Applied Linguistics guidelines (BAAL, 2021), and disclosed in full immediately afterward, at which point all participants were again offered the opportunity to withdraw their data; none did. The moderator introduced the discussion topic and then left the room for the remainder of each session, so the researcher was not a co-present participant in any face-to-face discussion.

The WhatsApp data are drawn from a closed, membership-restricted group of which the researcher was, and remains, a non-posting member; participants interacted without awareness that specific exchanges would later be selected for analysis. Consistent with Association of Internet Researchers guidance on semi-public and semi-private digital data (AoIR, 2019), no individual consent was sought for this closed-group archival material, and all quoted messages are reported only through anonymized, gender-coded labels. A single researcher carried out all coding for both datasets without a second coder; accordingly, the study does not report a conventional inter-rater reliability statistic. Intra-rater consistency was instead supported by explicit, consistently applied operational criteria for each coding category and by a final review pass re-examining all coded instances against those criteria.

4. Results and Discussion

4.1 Face-to-Face Interruptive Events

Table 1 presents the distribution of face-to-face interruptive events by function and gender. Across the three discussion groups, female participants accounted for 80 of the 157 coded interruptive events (51.0%). This higher female share was not uniform across functions. It was largest in the corrective function (60.0% female, 40.0% male), smaller in the competitive function (52.3% female), and reversed in the supportive function, where male participants accounted for 70.0% of events. Overall, gender differences in face-to-face interruption depend heavily on which interactional function is at stake, though as reported below, none of these differences reach conventional statistical significance at this sample size.

Table 1: Face-to-Face Interruptive Events by Function and Gender

Category	Female	Male	Total	Female %	Male %
Competitive interruptions	56	51	107	52.3%	47.7%
Corrective interruptions	18	12	30	60.0%	40.0%
Supportive overlaps	6	14	20	30.0%	70.0%
Total	80	77	157	51.0%	49.0%

The group-level breakdown in Table 2 reveals that this aggregate pattern conceals substantial variation across the three discussion groups. Group 1 generated sixteen interruptive events, evenly split between genders (50.0% female). Group 2 generated thirty-two events with a male share of 65.6%, reflecting the greater activity of particular male speakers despite the group's balanced gender composition. Group 3 generated the most events of any group, 109 in total, with female participants producing sixty-one of these, or 56.0%, the group driving most of the aggregate higher female share in the corpus. This group-level variation indicates that the aggregate higher female share reported above is concentrated in one especially dense pairing rather than distributed evenly across the sample, and that Group 2 alone runs counter to the aggregate direction.

Group 3's outsized contribution may partly reflect the group's shorter relational history: unlike Group 1, whose members had been friends for several years, and Group 2, whose members had a more moderate acquaintance, Group 3's participants had known one another for only a few weeks at the time of recording. In a newly formed peer group, interactional roles and floor-access patterns may not yet be evenly distributed among speakers, and the concentration of activity in a small number of highly active individuals, such as S1 in Example 1, may reflect this early-stage social positioning rather than a stable, group-wide behavioral pattern. This limitation should be kept in mind when interpreting Group 3's disproportionate contribution to the corpus-wide higher female share reported in Table 1.

Table 2: Face-to-Face Interruptive Events by Discussion Group and Gender

Group	Female	Male	Total	Female %
1	8	8	16	50.0%
2	11	21	32	34.4%
3	61	48	109	56.0%
Total	80	77	157	51.0%

None of the gender contrasts reported in Tables 1 and 2 reach conventional statistical significance. Exact binomial tests against a chance baseline of 50% give $p = .70$ for competitive interruption, $p = .36$ for corrective interruption, and $p = .12$ for supportive overlap (Table 1), and $p = 1.00$, $p = .11$, and $p = .25$ for Groups 1, 2, and 3 respectively (Table 2); a chi-square test of function by gender association is likewise non-significant ($\chi^2(2) = 4.58$, $p = .10$), as is a chi-square test of group by gender association ($\chi^2(2) = 4.62$, $p = .10$). At this sample size ($N = 157$ face-to-face events), the reported percentages are consistent with chance. They are read descriptively throughout, as a higher or lower share of events, rather than as evidence of a robust gendered effect. The qualitative readings developed in Sections 4.2 through 4.5 characterize the interactional function of the specific examples analyzed there; they are not offered as evidence

that these functions are performed at a statistically robust, population-level rate by one gender. The frequency comparisons introduced in Section 4.5 compare face-to-face and WhatsApp rates directly rather than comparing genders, and so are unaffected by this limitation; only claims about which gender shows the higher share within a function should be read as descriptive rather than established.

Competitive interruption in this corpus is not a uniform floor-seizure event. Its rhetorical force and face-management strategy vary considerably across speakers and groups. In Group 3, the corpus's most interactionally dense pairing, S1 (F) produces the highest individual interruption count. She frequently deploys direct negation and narrative counter-example to challenge the current speaker's implicit premise rather than simply displacing his turn, as the following example illustrates.

Example 1: (Competitive interruption, Group 3)

S4 (M) [extended turn]	[...dismissive framing: this man will spend his whole life earning 70 dirhams a day...]	
S1 (F) [Interruption]	Non mais ghayt3allem! Ghaymchi ykhdem, ghayhezz l7jar, ghaytsekhr lihom, chwiya ghayt3allem kifach ydreb blmla7, chwiya ghayt3allem kifach ybni 7ayet, chwiya bchwiya bchwiya ghaywli m3allem, fhemti? <i>/nɔ̃ mɛ ʔajtʃallem ʔajmʃi ʔəxdem ʔajhəzz lhʒar ʔajtsəxr lihum fwiʒa ʔajtʃallem kifaf ʔədrəb bəlmɫah fwiʒa ʔajtʃallem kifaf ʔəbni həjət fwiʒa bfwija bfwija ʔajwəlli mʃallem fhəmti/</i>	<i>No, but he will learn! He will go work, carry stones, run errands for them, slowly learn how to apply plaster, slowly learn how to build a wall, slowly by slowly he becomes a master craftsman, you know?</i>

The interruption opens with a direct negation, 'Non mais ghayt3allem!' ('No, but he will learn!'), immediately countering S4 (M)'s fatalistic framing of the laborer's prospects. S1 (F) then reuses S4's own scenario, the construction laborer earning a fixed wage, but redirects its trajectory through a cascading sequence, 'chwiya bchwiya bchwiya' ('slowly by slowly'), that traces a path from unskilled labor to master craftsman (m3allem). This is a form of what Goodwin and Goodwin (1987) call format tying. This entry borrows the grammatical and topical structure of the current speaker's proposition while inverting its content and implied outcome. The closing tag question, 'fhemti?' ('you know?'), seeks alignment even as the interruption contests the substance of S4's claim.

4.2 Corrective Interruption as an Epistemic Authority Claim

Corrective interruption shows the largest female share of any face-to-face category, at 60.0% female. In this corpus, it can be read as the function most closely tied to negotiating epistemic authority. Producing a corrective interruption can be understood as asserting not merely the right to speak. Still, the right to be treated as more knowledgeable or more empirically grounded than the speaker being corrected. The following example illustrates how this claim may be constructed in practice.

Example 2: (Corrective interruption, Group 2)

S4 (M) [Floor Defense]	[...claims job opportunities in Morocco exist only in Casablanca and Rabat...]	
S2 (F) [Corrective Overlap]	C'est la capitale. C'est la capitale... <i>/sɛ la kapital sɛ la kapital.../</i>	<i>It's the capital [city]. It's the capital...</i>

The correction here is compressed rather than elaborated. S4 (M) has just asserted that opportunities exist only in Casablanca and Rabat without explaining why; S2 (F)'s two-word phrase, repeated for emphasis, 'C'est la capitale' ('It's the capital'), supplies the missing causal account in a single stroke: these cities are chosen for their administrative and economic centrality, not by accident. This is a compact instance of what Pomerantz (1984) calls giving a source or basis, or telling how I know, a discursive strategy in which speakers construct the credibility of their claims through explicit reference to the evidential basis of their knowledge, here achieved through repetition rather than expansion.

4.3 Supportive Overlap and the Collaborative Floor

Supportive overlap is the one face-to-face function that runs counter to the corpus-wide higher female share, skewing male at 70.0% and female at 30.0%. This pattern, like the others reported above, is not statistically distinguishable from chance at this sample size, but it runs counter to the assumption, associated with Coates' (1996) collaborative floor model, that supportive overlapping is primarily a female interactional strategy; among the events coded here, it is more often male participants who perform the collaborative work of completing, extending, or affirming another speaker's still-unfolding turn. Male supportive overlaps in this corpus range from brief affirmative backchannels to full collaborative completions of a co-participant's unfinished sentence, as the following example shows.

Example 3: (Supportive overlap, Group 2)

S2 (F) [trailing turn]	3ndhom... ma-3reftsh kifash gh-nshre7. <i>/ʃandhum... ma ʃraftʃ kifaf ʃanʃrah/</i>	They have... I don't know how to explain.
S4 (M) [Collaborative Completion]	Des opportunités. <i>/dezɔpɔʁtynite/</i>	<i>Opportunities.</i>

Here S4 (M) supplies the single word S2 (F) is visibly searching for, completing her syntactic project rather than displacing it. This is a textbook instance of what Coates (1996) describes as jointly constructing talk, extending Edelsky's (1981) collaborative floor: two speakers contribute to a single utterance without either individually owning the resulting claim. That the male participant performs this collaborative completion, rather than the female participant as Coates' original account would predict, is precisely what makes supportive overlap the function most resistant to a straightforward gendered reading in this corpus.

4.4 WhatsApp Sequential Thread-Bypasses

A message-by-message re-coding of the fourteen WhatsApp segments identified 29 sequential thread-bypass events meeting the coding criteria set out in Section 3. Table 3 presents their

distribution by function and gender. Female participants account for 19 of these, or 65.5%. Unlike the dissertation-chapter figures this re-coding supersedes, the higher female share holds across all three functions rather than running the other way in the competitive category: competitive bypasses are 75.0% female, corrective bypasses 63.6% female, and supportive bypasses 64.3% female. A secondary, more liberal coding pass, which additionally counted topic-revivals and repeated unanswered questions lacking a single identifiable target message, identified 52 events in total and preserved the same pattern (competitive 75.0% female, corrective 57.9% female, supportive 61.9% female), indicating that the female-led pattern is not an artifact of the coding threshold chosen.

Table 3: WhatsApp Sequential Thread-Bypasses by Function and Gender

Category	Female	Male	Total	Female %	Male %
Competitive (topic-shift) bypasses	3	1	4	75.0%	25.0%
Corrective bypasses	7	4	11	63.6%	36.4%
Supportive bypasses	9	5	14	64.3%	35.7%
Total	19	10	29	65.5%	34.5%

One illustrative case shows this pattern clearly. Early in one WhatsApp segment, two male participants offer conflicting, hedged claims about whether a student who misses the retake exam is marked absent or keeps their original grade. The question goes unresolved for roughly one hundred intervening messages, through an unrelated stretch of conversation, until a female participant returns to it with a detailed, authoritative clarification: attendance at the retake is mandatory, she explains, outlining the specific exception for students who have already compensated on a prior grade. The clarification is accepted without further dispute. As in the face-to-face corrective examples above, the epistemic authority being claimed here appears not tied to any formal institutional role, but rather to the participant's demonstrated willingness and ability to resolve a piece of circulating misinformation with precise, actionable detail.

Code-switching in the WhatsApp corpus is not incidental to this pattern; it is one of its mechanisms, and unlike the interruption and bypass counts above, this pattern is itself statistically robust. Of the 1,853 coded messages, 249 (13.4%) involve a language switch. Of these, 143 were coded as serving an epistemic function, marking or reinforcing a knowledge claim, rather than signaling formality, humor, or affect: 96 were produced by female senders and 39 by male senders, with the remaining 8 produced by senders whose gender could not be determined with reasonable confidence and who are excluded from the comparison that follows, consistent with the criterion set out in Section 3. Among the 135 gender-identified epistemic switches, the split is highly significant (exact binomial test, $p < .001$); formality-marking switches show the same skew (51 female to 24 male of 75 gender-identified instances, $p = .002$; a further 7 instances from gender-undetermined senders are again excluded).

These switches predominantly pair Darija with English rather than French. This is a distinct claim from the corrective-interruption gender split reported in Section 4.1, which is not statistically significant: rather than shoring up that specific contrast, the code-switching finding independently establishes that female participants in this corpus more often deploy an overt linguistic marker of epistemic authority when a knowledge claim is at stake. Read alongside the

qualitative analysis in Section 4.2, it supports treating epistemic-authority marking, in general, as a gendered practice in this dataset (Heritage, 2012a, 2012b; Heritage & Raymond, 2005), without implying that corrective interruption specifically is itself a statistically robust gendered effect.

As with the face-to-face data, none of the WhatsApp gender contrasts in Table 3 reach significance (competitive $p = .63$, corrective $p = .55$, supportive $p = .42$, overall $p = .14$; chi-square $p = .91$), a limitation of the smaller bypass sample ($N = 29$) that should temper the strength of any claim about a female-led pattern in this dataset.

4.5 Measuring Interruption Across Media

The comparison demonstrates why interruption cannot be treated as a medium-neutral countable unit. Table 4 places the face-to-face and WhatsApp results side by side as reported, without adjustment for how each medium packages interaction into countable units. Read this way, the picture is considerably less uniform than a single aggregate higher female share. The corrective function shows a female-skewed pattern in both settings, at 60.0% face-to-face and 63.6% on WhatsApp. The competitive function also skews female in both settings, more strongly on WhatsApp (75.0%) than face-to-face (52.3%). The supportive function is where the two media diverge most sharply: it is female-led on WhatsApp (64.3%) but male-led face-to-face (70.0% male, 30.0% female), the only function in the dataset whose direction differs across media. Read without adjustment, one might conclude that face-to-face and WhatsApp interaction carry a broadly consistent female association for competitive and corrective interruption, but opposite associations for supportive work specifically. This comparison, however, treats interruption and thread-bypass as though they were the same measurement taken twice, once in each medium, which is precisely the assumption that requires justification rather than being taken for granted.

Table 4: Face-to-Face and WhatsApp Interruptive/Bypass Events by Function: Unadjusted Comparison

Function	Face-to-Face Female %	WhatsApp Female %	Comparability (see Table 5)
Competitive	52.3%	75.0%	Not directly comparable
Corrective	60.0%	63.6%	Conditional
Supportive	30.0%	64.3%	Conditional

Three features of the comparison make a face-value reading unsafe. First, WhatsApp habitually divides a single contribution across several consecutive messages, so a raw count of bypass events, or of any WhatsApp unit, risks conflating platform-driven message fragmentation with genuine behavioral difference; applying the ninety-second same-sender merge rule described in Section 3 to this corpus collapses 1,853 raw messages into 1,395 functional turns, a 24.7% reduction. Second, and more fundamentally, face-to-face interruption is a real-time floor-competition event, observable and interruptible only because participants can monitor an utterance as it unfolds; WhatsApp bypass is a sequence-disruption event, occurring precisely because no such real-time monitoring is available (Herring, 1999; Garcia and Jacobs, 1999). The two are not the same underlying mechanism relabeled for a new medium, and a percentage-of-total comparison implicitly treats them as though they were. Third, and independent of measurement, the two corpora are drawn from different populations at different institutions

(Section 3): the face-to-face and WhatsApp comparisons that follow are therefore consistent with a medium effect but do not isolate one, since institutional or group-level differences between the two student populations cannot be ruled out as an alternative or contributing explanation.

Table 5 addresses both problems: it reports event rates per 1,000 words of message or speech text rather than percentages of differently sized totals, and it states, function by function, whether direct comparison is warranted. The face-to-face word count used for this normalization, 10,408 words of transcribed speech across the three focus groups, was recomputed from the corrected merged transcripts described in the accompanying interruption-coding audit and supersedes an earlier figure drawn from a corpus snapshot that predated those corrections.

The corrected comparison changes what the divergence in Table 4 supports, though not uniformly. Once expressed as a rate, corrective interruption occurs at 2.88 events per 1,000 words face-to-face, roughly two to three and a half times the rate of its WhatsApp analog (0.83 to 1.43, depending on the coding threshold), a real gap rather than the near-parity the raw percentages alone might suggest; both are marked conditionally comparable in Table 5, and the persistent gap is read here as evidence that correction is more readily produced within the real-time organization of face-to-face conversation than through an asynchronous typed reply, not as evidence that the two are different interactional moves.

Supportive overlap is closer to parity in raw rate (1.92 face-to-face against 1.05 to 1.58 on WhatsApp) than corrective is. Still, the two media diverge sharply in who performs it: face-to-face supportive overlap is majority-male in this corpus (70.0%), while WhatsApp supportive bypass remains majority-female (64.3%), the one function whose gender association differs in direction across media even where its frequency does not.

Competitive events tell the clearest story of non-comparability. Face-to-face competitive interruption occurs at 10.28 events per 1,000 words, a substantially higher rate than its WhatsApp counterpart (0.30 to 0.90, depending on the coding threshold described in Section 3; see Table 5 for the full comparison). Because face-to-face competitive interruption depends on floor competition that has no WhatsApp equivalent, this function is marked as not directly comparable in Table 5. The rate gap is read here not as a reversal of gender association, but as evidence that direct floor competition specifically does not transfer to asynchronous messaging in anything like its face-to-face form; when a functional equivalent does occur on WhatsApp, as the topic-shift bypass, it remains female-led, just markedly less frequent.

Read together, Tables 4 and 5 support a more precise and more heterogeneous version of the paper's central claim than either table supports alone. The unadjusted comparison in Table 4 shows a higher female share in the competitive and corrective functions that holds, in the same direction, across both media. In comparison, supportive work shows higher shares by different genders depending on the medium. The volume-corrected comparison in Table 5 adds that even where the direction of the gender association holds, as with corrective interruption, its frequency does not: face-to-face correction occurs at several times the rate of its WhatsApp analog. The one function tied most closely to real-time floor competition, competitive interruption, all but disappears on WhatsApp in rate terms rather than simply changing hands, while the one function tied most closely to cooperative floor-building, supportive overlap, keeps a broadly similar rate across media but changes hands entirely.

Table 5: Rate-Normalized Comparison by Function and Medium

Function	Comparability status	F2F rate (events / 1,000 words)	WhatsApp rate (events / 1,000 words)
Competitive	Not directly comparable: real-time floor-seizure (F2F) vs. agenda-redirection with no floor to seize (WhatsApp)	10.28	0.30
Corrective	Conditional: epistemic challenge, plausibly the same move regardless of channel	2.88	0.83
Supportive	Conditional: cooperative affirmation/elaboration, plausibly the same move regardless of channel	1.92	1.05
Total	N/A	15.08	2.18

The traditional dominance account is complicated by this pattern in more than one way: not by a single reversal in who dominates, but by three functions that diverge from a uniform cross-medium account in three different respects: frequency, direction, and mechanism. Of these, frequency is a direct rate comparison between media rather than a gender contrast, and mechanism is a conceptual distinction between floor-competition and sequence-disruption; only direction, which gender shows the higher share within a function, rests on the contrasts shown in Section 4.1 not to be statistically distinguishable from chance at this sample size, and is read accordingly as suggestive rather than established.

5. Conclusion

These findings do not support the traditional dominance account in its original form, nor do they support a simple reversal of it, and the corrected data complicate the picture further than the raw counts alone. Female participants in this corpus account for a higher share of interruptive and bypass events in both settings for the competitive and corrective functions, but this direction does not hold for supportive overlap, which skews male face-to-face and female on WhatsApp, the one function whose gender association differs in direction across media. None of these contrasts reach conventional statistical significance at this sample size (Section 4.1), and are treated throughout as descriptive rather than as evidence of a robust effect.

The higher female share is most consistent in direction in the corrective function, where interruption can be read as an epistemic authority claim (Heritage & Raymond, 2005) rather than as the undifferentiated floor-seizure that Zimmerman and West (1975) treated as the defining feature of male-enacted interruption. Still, even here the two media are not equivalent in frequency; correction occurs several times more often face-to-face than its WhatsApp analog, suggesting that epistemic correction, while structurally similar across media, is more readily produced within the real-time organization of conversation than through an asynchronous typed reply. The function most closely tied to the original dominance account's concern with unilateral floor control, competitive interruption, does not reverse gender association on WhatsApp; instead, its WhatsApp analog, topic-shift bypass, occurs at a small fraction of its face-to-face rate, suggesting that direct floor competition is the component of the dominance

account least able to survive the move to a medium where no real-time floor exists to compete for.

More broadly, the comparison demonstrates that interruption is not a medium-neutral countable unit. What can be identified and counted as an interruptive event depends partly on the temporal and sequential organization through which interaction is produced. Face-to-face conversation makes simultaneous floor competition directly observable, whereas WhatsApp makes sequence disruption observable but not simultaneous competition in the same form. The resulting differences in event frequency should therefore not be interpreted solely as differences in participants' interruption behavior; they also reflect differences in the interactional conditions under which interruption can become analytically visible in the first place.

Creative Commons License Statement

This research work is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License. To view a copy of this license, visit <https://creativecommons.org/licenses/by-nc-nd/4.0>. To view the complete legal code, visit <https://creativecommons.org/licenses/by-nc-nd/4.0/legalcode.en>. Under the terms of this license, members of the community may copy, distribute, and transmit the article, provided that proper, prominent, and unambiguous attribution is given to the authors, and the material is not used for commercial purposes or modified in any way. Reuse is only allowed under the terms of the Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License.

Conflict of Interest Statement

The authors declare there are no conflicts of interest.

About the Authors

Hanane Salif is a doctoral researcher in Sociolinguistics at the Faculty of Letters & Humanities Ain Chock, Hassan II University of Casablanca, specializing in online discourse and digital communication. She holds a Master's degree in Gender Studies and researches how participation, authority, and identity are negotiated through digital discourse.

ORCID: <https://orcid.org/0000-0002-0005-416X>

Mina Afkir is a Full Professor of Linguistics at the Faculty of Letters & Humanities Ain Chock, Hassan II University of Casablanca, where she directs the Language and Linguistics Studies (LLS) Research Laboratory. Her research interests include Moroccan Arabic, diglossia, and language in education.

ORCID: <https://orcid.org/0000-0002-7515-7702>

References

AoIR (Association of Internet Researchers), 2019. Internet Research: Ethical Guidelines 3.0. Association of Internet Researchers. Retrieved from <https://aoir.org/reports/ethics3.pdf>

- BAAL (British Association for Applied Linguistics), 2021. Recommendations on Good Practice in Applied Linguistics, 4th edition. BAAL, UK. Retrieved from <https://aila.info/baals-recommendations-on-good-practice-in-applied-linguistics-2016-aila/>
- Caronia L, 2024. Epistemic and Deontic Authority in Parent-Teacher Conference: Referring to the Expert as a Discursive Practice to (Jointly) Undermine the Teacher's Expertise. *Journal of Teacher Education* 75(4): 397-411. <https://doi.org/10.1177/00224871231153088>
- Caubet D, 2004. *Les Mots du Bled*. L'Harmattan, Paris, France. Retrieved from https://books.google.ro/books/about/Les_mots_du_bled.html?id=EXvA85oIvaoC&redir_esc=y
- Coates J, 1996. *Women Talk: Conversation between Women Friends*. Blackwell, Oxford, UK. Retrieved from https://books.google.ro/books/about/Women_Talk.html?id=V8_zHs8JPRUC&redir_esc=y
- Colom A, 2022. Using WhatsApp for Focus Group Discussions: Ecological Validity, Inclusion, and Deliberation. *Qualitative Research* 22(3): 452-467. Retrieved from <https://doi.org/10.1177/1468794120986074>
- Edelsky C, 1981. Who's Got the Floor? *Language in Society* 10(3): 383-421. Retrieved from <https://www.cambridge.org/core/journals/language-in-society/article/whos-got-the-floor/655A684D40C1F0310A46A7BC909059B5>
- Garcia AC, Jacobs JB, 1999. The Eyes of the Beholder: Understanding the Turn-Taking System in Quasi-Synchronous Computer-Mediated Communication. *Research on Language and Social Interaction* 32(4): 337-367. Retrieved from <https://www.learntechlib.org/p/89433/>
- Goldberg JA, 1990. Interrupting the Discourse on Interruptions: An Analysis in Terms of Relationally Neutral, Power- and Rapport-Oriented Acts. *Journal of Pragmatics* 14(6): 883-903. [https://doi.org/10.1016/0378-2166\(90\)90045-F](https://doi.org/10.1016/0378-2166(90)90045-F)
- Goodwin MH, Goodwin C, 1987. Children's Arguing. In: Philips SU, Steele S, Tanz C (eds) *Language, Gender, and Sex in Comparative Perspective*. Cambridge University Press, Cambridge, UK, pp 200-248. Retrieved from <https://psycnet.apa.org/record/1987-98360-008>
- Heritage J, 2012a. Epistemics in Action: Action Formation and Territories of Knowledge. *Research on Language and Social Interaction* 45(1): 1-29. <https://doi.org/10.1080/08351813.2012.646684>
- Heritage J, 2012b. The Epistemic Engine: Sequence Organization and Territories of Knowledge. *Research on Language and Social Interaction* 45(1): 30-52. <https://doi.org/10.1080/08351813.2012.646685>
- Heritage J, Raymond G, 2005. The Terms of Agreement: Indexing Epistemic Authority and Subordination in Talk-in-Interaction. *Social Psychology Quarterly* 68(1): 15-38. <https://doi.org/10.1177/019027250506800103>
- Herring SC, 1999. Interactional Coherence in CMC. *Journal of Computer-Mediated Communication* 4(4). <https://doi.org/10.1111/j.1083-6101.1999.tb00106.x>
- Holes C, 2004. *Modern Arabic: Structures, Functions, and Varieties* (rev. ed). Georgetown University Press, Washington, DC, USA. Retrieved from

<https://archive.org/details/georgetown-classics-in-arabic-languages-and-linguistics-series-clive-holes-roger>

- James D, Clarke S, 1993. Women, Men, and Interruptions: A Critical Review. In: Tannen D (ed) Gender and Conversational Interaction. Oxford University Press, New York, NY, USA, pp 231-280. <https://psycnet.apa.org/record/1993-98939-009>
- Miller MG, Sutherland JL, 2023. The Effect of Gender on Interruptions at Congressional Hearings. American Political Science Review 117(1): 103-121. Retrieved from <https://www.cambridge.org/core/journals/american-political-science-review/article/effect-of-gender-on-interruptions-at-congressional-hearings/DDD33F28A1C8ED9C162E1793D8126243>
- Pomerantz A, 1984. Giving a Source or Basis: The Practice in Conversation of Telling 'How I Know.' Journal of Pragmatics 8(5-6): 607-625. [https://doi.org/10.1016/0378-2166\(84\)90002-X](https://doi.org/10.1016/0378-2166(84)90002-X)
- Sacks H, Schegloff EA, Jefferson G, 1974. A Simplest Systematics for the Organization of Turn-Taking for Conversation. Language 50(4): 696-735. <https://doi.org/10.2307/412243>
- Tannen D, 1994. Gender and Discourse. Oxford University Press, New York, NY, USA. <https://doi.org/10.2307/2655317>
- Tong Y, 2025. Heard or Halted? Gender, Interruptions, and Emotional Tone in U.S. Supreme Court Oral Arguments. Working paper, Georgetown University. Retrieved from <https://arxiv.org/abs/2512.05832>
- West C, Zimmerman DH, 1983. Small Insults: A Study of Interruptions in Cross-Sex Conversations between Unacquainted Persons. In: Thorne B, Kramarae C, Henley N (eds) Language, Gender and Society. Newbury House, Rowley, MA, USA, pp 102-117. Retrieved from https://books.google.ro/books/about/Language_Gender_and_Society.html?id=xiliAAAAIAAJ&redir_esc=y
- Zimmerman DH, West C, 1975. Sex Roles, Interruptions, and Silences in Conversation. In: Thorne B, Henley N (eds) Language and Sex: Difference and Dominance. Newbury House, Rowley, MA, USA, pp 105-129. Retrieved from https://books.google.ro/books/about/Language_and_Sex.html?id=SWywAAAAIAAJ&redir_esc=y