

Abstract

The Irrelevant Sound Effect (ISE) is a well-established phenomenon where background sound disrupts serial recall. However, debate persists regarding whether this disruption arises solely from acoustic variation (Changing-State Hypothesis) or if semantic content exerts additional interference (Interference-by-Content). This study investigated whether Speech is uniquely disruptive compared to non-semantic changing-state sound (Music). Sixty-five undergraduate students ($N = 65$) completed a serial recall task of eight random digits under three within-subjects conditions: Silence, Music (*Eine Kleine Nachtmusik*), and Speech (Harvard Sentences). A One-Way Repeated Measures ANOVA revealed a significant main effect of sound condition on recall performance ($p < .001$). Bonferroni-corrected comparisons demonstrated that while both sounds significantly impaired recall compared to Silence, **Speech was significantly more disruptive than Music** ($p = .002$). These findings challenge a strict acoustic-only account of the ISE and support the view that semantic content creates specific, additional interference in short-term memory. Practically, this suggests that intelligible Speech is more detrimental to cognitive performance than instrumental noise in work environments.

I. Introduction

Context: The Irrelevant Sound Effect (ISE) refers to the finding that task-irrelevant background sound disrupts serial recall performance compared to Silence [1].

The Debate:

- **Changing-State Hypothesis:** Argues that disruption is driven solely by acoustic variation (e.g., changes in pitch/timbre), not meaning [1]. This predicts that **Speech** and **Music** should be equally disruptive if their acoustic variability is comparable.
- **Interference-by-Content:** Suggests that the semantic content (meaning) of Speech creates specific, additional interference with the rehearsal of to-be-remembered items [2].

The Current Study: We compared **Silence**, **Music** (*Eine Kleine Nachtmusik*), and **Speech** (Harvard Sentences) to determine if semantic content exerts a unique disruptive effect.

Hypotheses:

- **H1:** Both Music and Speech will significantly disrupt recall compared to Silence (Replicating the ISE).
- **H2:** If semantic content creates additional interference, **Speech** will be significantly more disruptive than **Music**.

II. Method

- **Participants:** $N = 65$ undergraduate students recruited via opportunity sampling.
- **Design:** A one-way repeated measures design was used.
 - **IV:** Sound Condition (Silence, Music, Speech).
 - **DV:** Serial Recall Score (Number of digits correctly recalled in order out of 8).
 - **Note:** Each condition comprised 10 trials; block order was counterbalanced across participants.
- **Materials:**
 - **Task:** Visual presentation of 8 random digits (1–9).
 - **Sounds:** *Eine Kleine Nachtmusik* (Music) and Harvard Sentences (Speech).

Procedure:

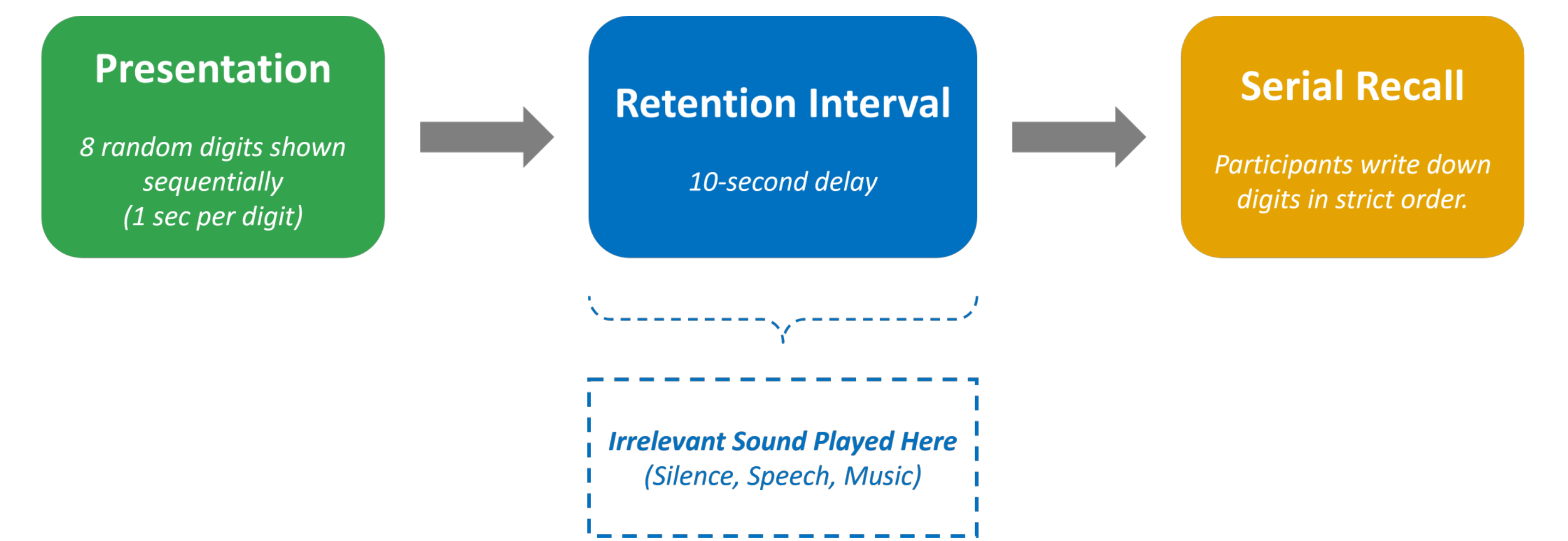


Figure 1. Schematic representation of a single experimental trial.

III. Results

Assumption Checks & Data Treatment

Preliminary analysis indicated recall scores in all conditions were negatively skewed ($z < -1.96$) and non-normal (Shapiro-Wilk, $p < .05$), indicating a ceiling effect. Standard transformations were inappropriate for negative skew. Given the large sample size ($N = 65$), the One-Way Repeated Measures ANOVA was considered robust to these violations [3] and analysis proceeded on raw data.

Descriptive Statistics

Mean serial recall scores were highest in the Silence condition and lowest in the Speech condition (**Table 1**).

Condition	Mean (M)	Std. Deviation
Silence	6.69	1.57
Music	6.42	1.53
Speech	6.08	1.57

Table 1. Mean serial recall scores (max 8) and standard deviations ($N = 65$).

Main Analysis

Mauchly's test indicated a violation of sphericity ($\chi^2(2) = 6.87, p = .032$); degrees of freedom were corrected using Greenhouse-Geisser estimates ($\epsilon = .906$). A One-Way Repeated Measures ANOVA revealed a significant main effect of sound condition on serial recall performance, $F(1.81, 116.00) = 17.55, p < .001, \eta_p^2 = .215$ (indicating a large effect).

Pairwise Comparisons

Bonferroni-corrected comparisons indicated that recall was significantly higher in **Silence** compared to both **Music** ($p = .014$) and **Speech** ($p < .001$). Crucially, recall in the **Speech** condition was significantly lower than in the **Music** condition ($p = .002$, Mean Diff = -0.33), suggesting that Speech was more disruptive than Music.

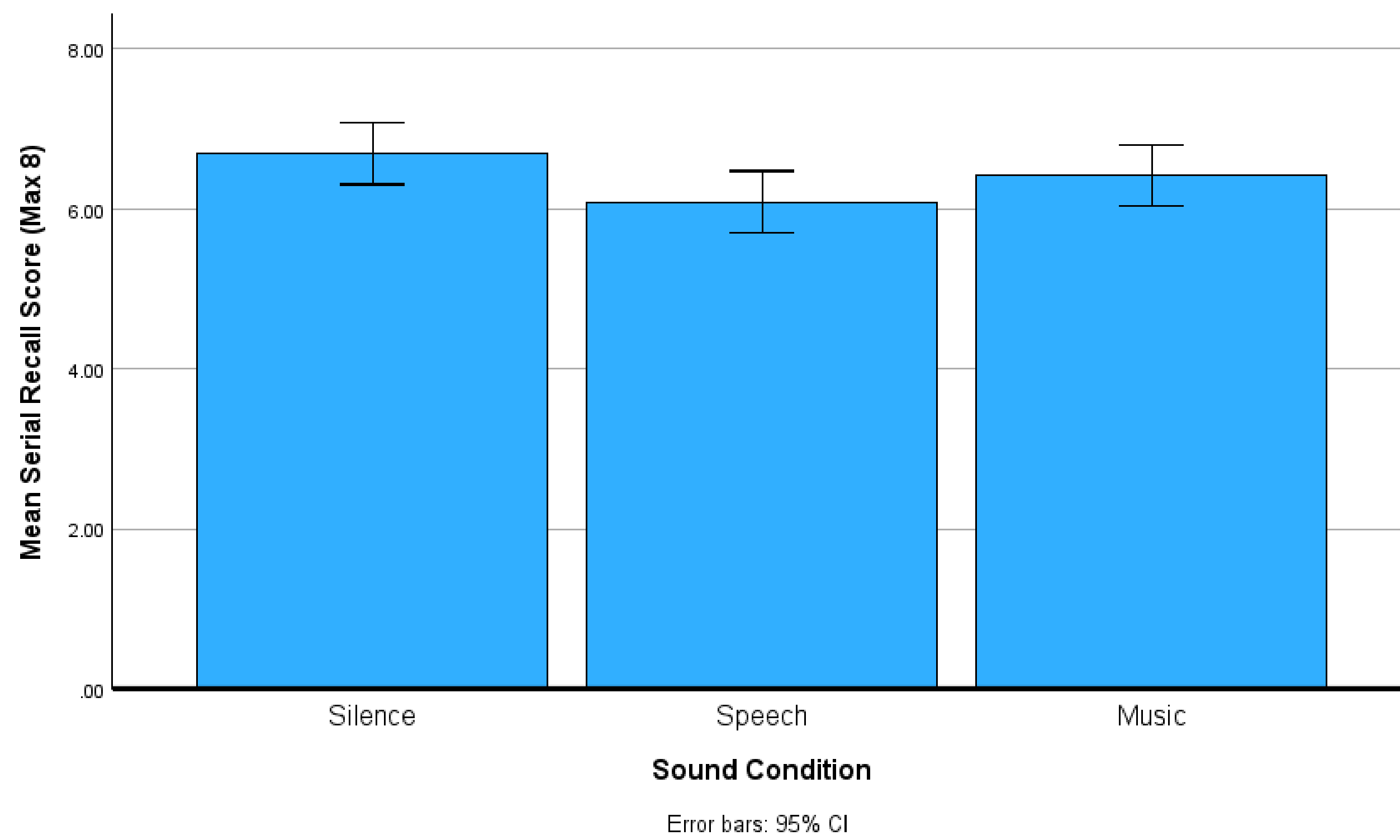


Figure 2. Mean serial recall performance across sound conditions. Error bars represent 95% Confidence Intervals.

IV. Discussion

Interpretation of Findings

The results supported **H1**: irrelevant sound significantly disrupted serial recall compared to Silence, replicating the robust nature of the ISE [1]. Numerically, Music performance fell between Silence and Speech, indicating that while Music is disruptive, it is significantly less harmful than Speech.

Crucially, the data also supported **H2**: **Speech** was significantly more disruptive than **Music**. This finding challenges a strict acoustic-only account (Changing-State Hypothesis), which predicts equal disruption if acoustic variation is comparable [1]. Instead, the greater disruption caused by Speech supports interference-by-content accounts [2], suggesting that the semantic nature of Speech creates specific, additional interference with the rehearsal of to-be-remembered items.

Limitations

A significant limitation was the presence of a **ceiling effect** (negative skew), indicated by high mean scores ($M > 6$) across conditions. This suggests the task was insufficiently demanding for this sample, which compressed the scores and may have masked the true magnitude of the difference between the Speech and Music conditions.

V. Conclusion

While all changing-state sound is disruptive, this study confirms that Speech is **”special”** in its capacity to degrade short-term memory. This suggests that the processing of meaning cannot be fully disengaged even when irrelevant to the task.

Wider Application: In safety-critical settings or open-plan offices, strictly reducing volume is insufficient; effective noise abatement policy must specifically target the intelligibility of background Speech to maximize cognitive performance [1].

References

- [1] R. Hughes and D. M. Jones, “The intrusiveness of sound: Laboratory findings and their implications for noise abatement,” *Noise and Health*, vol. 4, no. 13, pp. 51–70, 2001.
- [2] N. Perham, H. Hodgetts, and S. Banbury, “Mental arithmetic and non-speech office noise: An exploration of interference-by-content,” *Noise and Health*, vol. 15, no. 62, pp. 73–78, 2013.
- [3] A. Field, *Discovering statistics using IBM SPSS statistics*. Sage, 4th ed., 2013.