

Ipse Dixit, But Not in Science: An Alternative Approach to Implicit Measures

Ottavia M. Epifania^{1,2},
Pasquale Anselmi³, Egidio Robusto³, Gianmarco Altoè^{2,3}



¹ University of Trento, Rovereto (IT)

² Psicostat (Scientific Community Beyond Borders)

³ University of Padova, Padova (IT)

Predoc Camp @ DiPSCo
University of Trento

April, 10th 2025

*L'idea che la soluzione venga dall'intuizione del matto genio per natura e non dal lavoro complicato e **collettivo** di centinaia, migliaia di scienziati, questa idea è un'idea falsa e sbagliata, che toglie valore all'Università [...]*

Matteo Bordone, 19 Febbraio 2025

Introduction

According to Greenwald & Banaji (1995), implicit attitudes are defined as:

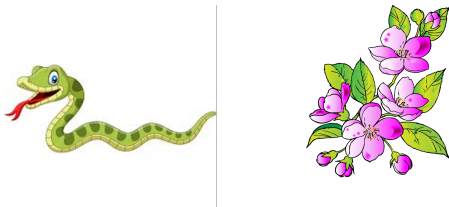
Introspectively unidentified – or inaccurately identified – traces of past experience that mediate favorable or unfavorable feelings, thoughts, or actions toward social objects

IMPLICIT = UNCONSCIOUS

Implicit attitudes are expressed through so-called **automatic associations**

Greenwald, A. G., & Banaji, M. R. (1995) Implicit Social Cognition: Attitudes, Self-Esteem, and Stereotypes. *Psychological Review*, 102-1, doi: 10.1037/0033-295X.102.1.4

Automatic associations



- Not controllable
- Triggered by “triggering” stimuli
- Not accessible through introspection
- Fast and almost immediate

Automatic associations and the unconscious

Studying automatic associations = Studying implicit attitudes



Gaining access to the unconscious and (finally!) being able to study it scientifically

WRONG!

Fazio & Olson (2003):

Being quick in associating snakes with negative adjectives or words does not imply that one is unaware of their negative attitudes toward snakes!

The only thing one is truly unaware of is that someone is measuring attitudes!

The attitude (the object of measurement) is not implicit, the measurement process itself is.

Implicitly measured constructs Vs. Unconscious constructs

Fazio, R. H., & Olson, M. A. (2003). Implicit measures in social cognition research: Their meaning and use. *Annual Review of Psychology*, 54(1), 297–327.

<https://doi.org/10.1146/annurev.psych.54.101601.145225>

Banaji & Greenwald (2013):

Theoretical definition of the term implicit as unconscious and unaware

Greenwald & Banaji (2017), Greenwald & Lai (2020):

Empirical definition of the term implicit as indirect

Banaji, M. R., & Greenwald, A. G. (2013). *Blindspot: Hidden biases of good people*. Delacorte Press,

Greenwald, A. G., & Banaji, M. R. (2017). The implicit revolution: Reconceiving the relation between conscious and unconscious. *American Psychologist*, 72(9), 861–871. <https://doi.org/10.1037/amp0000238>

Greenwald, A. G., & Lai, C. K. (2020). Implicit social cognition. *Annual Review of Psychology*, 71, 419–445. <https://doi.org/10.1146/annurev-psych-010419-050837>

Implicit = Indirect

Why...?

- Lack of scientific empirical evidence to support access to the unconscious
- Issues related to construct validity → What are we measuring?
Are we *sure* we are measuring what we think we are measuring?
- Definition without a supporting theory
- Results cannot be replicated...quite an issue

The Implicit Association Test

Coke
Good

Pepsi
Bad



Response key: E

Coke
Good

Pepsi
Bad



Response key: I

Coke
Good

Pepsi
Bad

terrible

Response key: I

Coke
Good

Pepsi
Bad

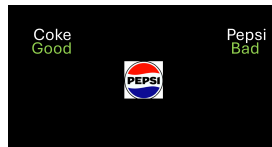
glory

Implicit Association Test

Greenwald et al. (1998):

Block	# Trials	Left Key (E)	Right Key (I)
1	20	Good	Bad
2	20	Coke	Pepsi
3	20	Coke + Good	Pepsi + Bad
4	40	Coke + Good	Pepsi + Bad
5	20	Pepsi	Coke
6	20	Pepsi + Good	Coke + Bad
7	40	Pepsi + Good	Coke + Bad

Implicit Association Test

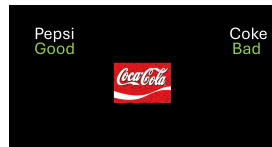


Greenwald et al. (1998):

Block	# Trials	Left Key (E)	Right Key (I)
1	20	Good	Bad
2	20	Coke	Pepsi
3	20	Coke + Good	Pepsi + Bad
4	40	Coke + Good	Pepsi + Bad
5	20	Pepsi	Coke
6	20	Pepsi + Good	Coke + Bad
7	40	Pepsi + Good	Coke + Bad

Condition Coke-Good/Pepsi-Bad

Implicit Association Test



Greenwald et al. (1998):

Block	# Trials	Left Key (E)	Right Key (I)
1	20	Good	Bad
2	20	Coke	Pepsi
3	20	Coke + Good	Pepsi + Bad
4	40	Coke + Good	Pepsi + Bad
5	20	Pepsi	Coke
6	20	Pepsi + Good	Coke + Bad
7	40	Pepsi + Good	Coke + Bad

Condition Coke-Good/Pepsi-Bad

Condition Pepsi-Bad/Coke-Good

D score

Greenwald et al. (2003)

$$D_{B6,B3} = \frac{M_{B6} - M_{B4}}{sd_{B6,B3}}$$

$$D_{B7,B4} = \frac{M_{B7} - M_{B4}}{sd_{B7,B4}}$$

$$D = \frac{D_{B6,B3} + D_{B7,B4}}{2}$$

Error responses? Fast responses?

Greenwald, A. G., Nosek, B. A., & Banaji, M. R. (2003). Understanding and using the implicit association test: I. An improved scoring algorithm. *Journal of Personality and Social Psychology*, 85(2), 197–216. <https://doi.org/10.1037/0022-3514.85.2.197>

Random Factors

Stimuli are fixed, respondents are random

However...

The stimuli can also represent a sample of a larger universe

Processing speed of positive and negative attributes

There is a universe of **positive attributes** as well as an universe of **negative attributes**

Only samples of **positive attributes** (e.g., good, nice, ...) and **negative attributes** (e.g., bad, evil, ...) are administered

Stimuli are fixed, respondents are random

However...

The stimuli can also represent a sample of a larger universe

Processing speed of positive and negative attributes

There is a universe of **positive attributes** as well as an universe of **negative attributes**

Only samples of **positive attributes** (e.g., good, nice, ...) and **negative attributes** (e.g., bad, evil, ...) are administered

So... there must be a sampling variability!

What if the sampling variability is not acknowledged

Generalizability

Generalizability is bounded to the specific set of stimuli used in the experiment

Results can be generalized if and only if the exact same set of stimuli is used

Robustness of the results

Random variability at the stimulus level might inflate the probability of committing Type I errors

Averaging across stimuli to obtain person-level scores results in biased estimates due to the noise in the data

Loss of information

All the variability is not considered as well as all the information that can be obtained from it

Every stimulus is assumed to be equally informative

Random Effects for Random Factors

Random effects and random factors

Linear combination of predictors in a Linear Model:

$$\eta = X\beta,$$

where β indicates the coefficients of the fixed intercept and slope(s), and X is the model-matrix.

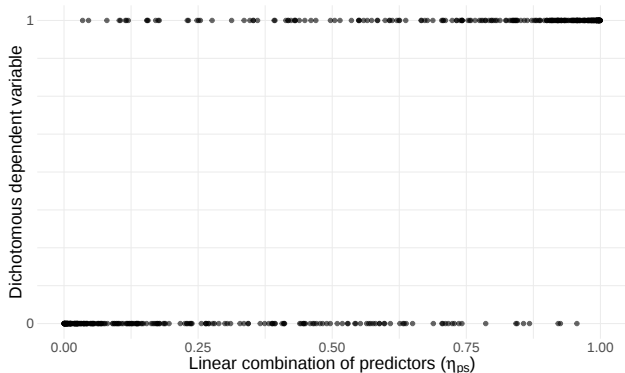
Linear combination of predictors in a Linear Mixed-Effects Model (LMM):

$$\eta = X\beta + Zd,$$

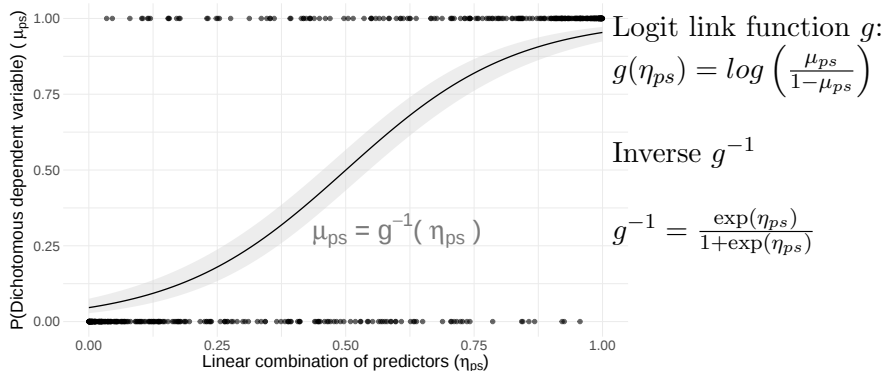
where Z is the matrix and d is the vector of the random effects (not parameters!)

Best Linear Unbiased Predictors

Generalized linear model (GLM) for dichotomous responses



Generalized linear model (GLM) for dichotomous responses



The Rasch model

$$P(x_{ps} = 1 | \theta_p, b_s) = \frac{\exp(\theta_p - b_s)}{1 + \exp(\theta_p - b_s)}$$

where:

θ_p : ability of respondent p (i.e., latent trait level of respondent p)

b_s : difficulty of stimulus s (i.e., “challenging” power of stimulus s)

Standard

$$P(x_{ps} = 1) = \frac{\exp(\theta_p - b_s)}{1 + \exp(\theta_p - b_s)}$$

GLM

$$P(x_{ps} = 1) = \frac{\exp(\theta_p + b_s)}{1 + \exp(\theta_p + b_s)}$$

The Rasch model

$$P(x_{ps} = 1 | \theta_p, b_s) = \frac{\exp(\theta_p - b_s)}{1 + \exp(\theta_p - b_s)}$$

where:

θ_p : ability of respondent p (i.e., latent trait level of respondent p)

b_s : difficulty of stimulus s (i.e., “challenging” power of stimulus s)

Standard

$$P(x_{ps} = 1) = \frac{\exp(\theta_p - b_s)}{1 + \exp(\theta_p - b_s)}$$

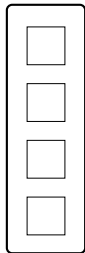
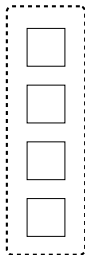
GLM

$$P(x_{ps} = 1) = \frac{\exp(\theta_p + b_s)}{1 + \exp(\theta_p + b_s)}$$

p_1

p_2

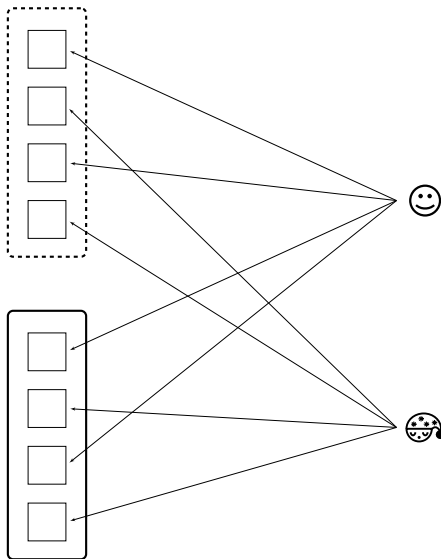
p_3

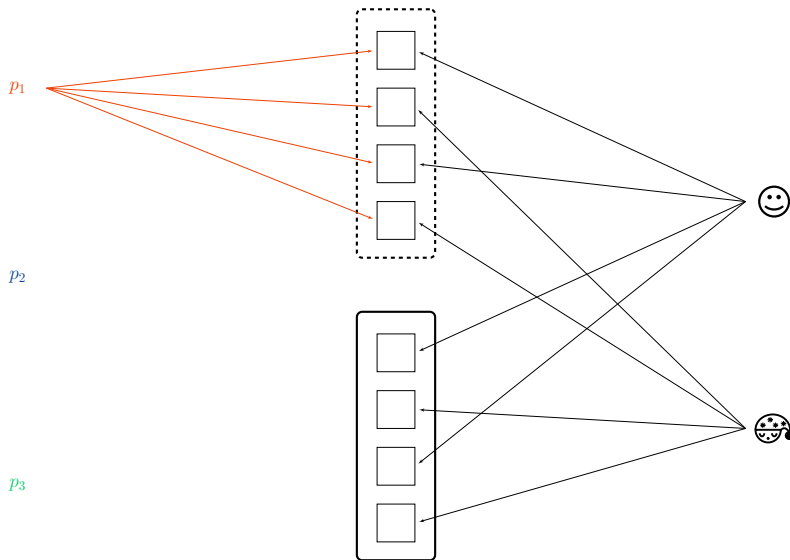


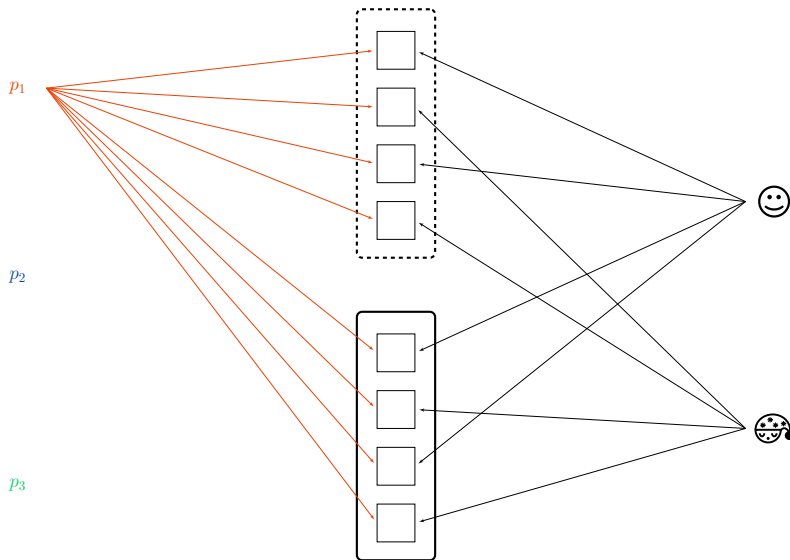
p_1

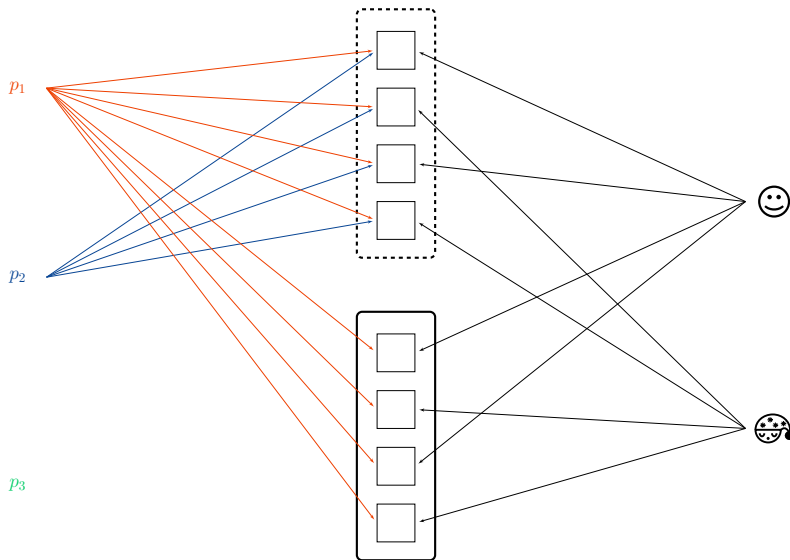
p_2

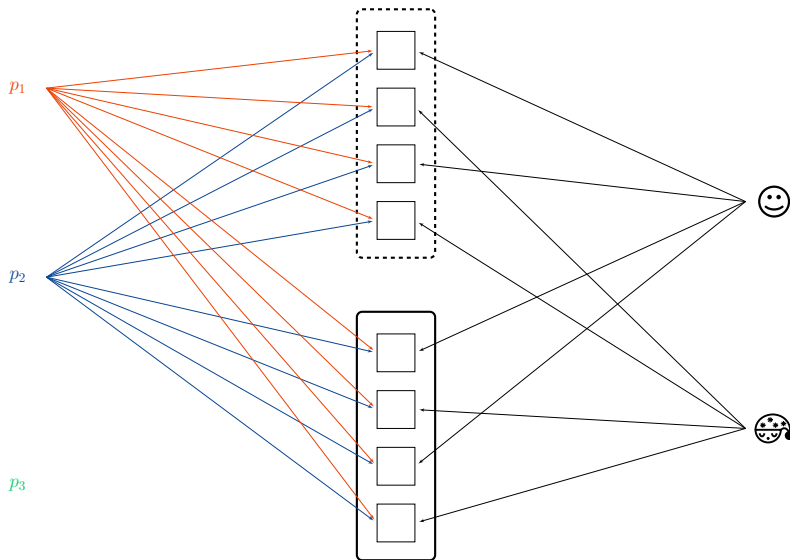
p_3

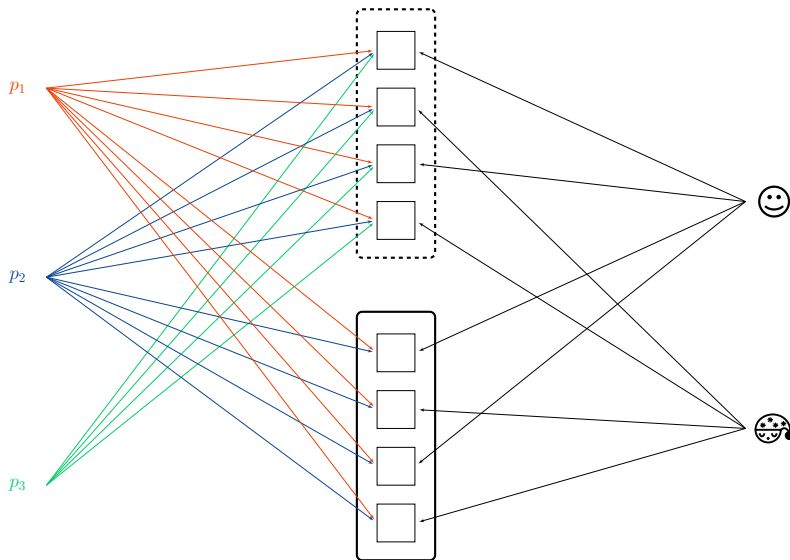


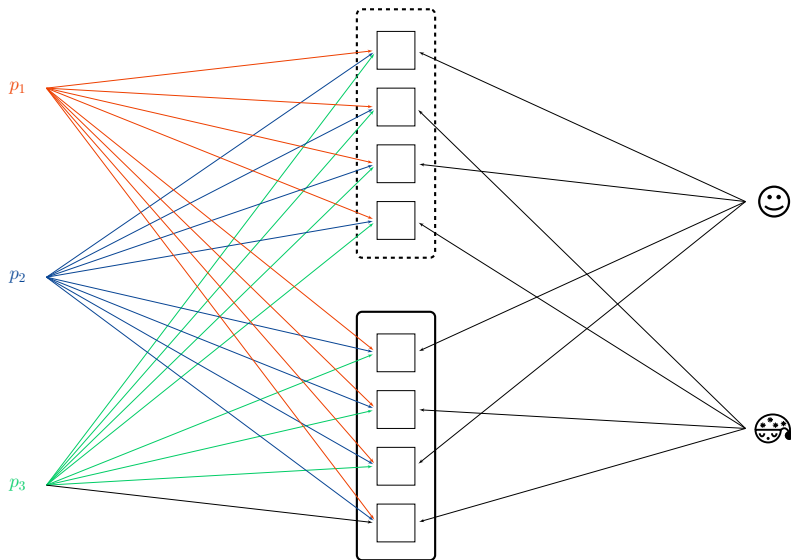












The expected response y for the observation $i = 1, \dots, I$ for respondent $p = 1, \dots, P$ on stimulus $s = 1, \dots, S$ in condition $c = 1, \dots, C$:

Model 1:

$$\begin{aligned} y_i &= \text{logit}^{-1}(\alpha + \beta_c X_c + \alpha_{p[i]} + \alpha_{s[i]}) \\ \alpha_p &\sim \mathcal{N}(0, \sigma_p^2), \\ \alpha_s &\sim \mathcal{N}(0, \sigma_s^2). \end{aligned}$$

Model 2:

$$\begin{aligned} y_i &= \text{logit}^{-1}(\alpha + \beta_c X_c + \alpha_{p[i]} + \beta_{s[i]} c_i) \\ \alpha_p &\sim \mathcal{N}(0, \sigma_p^2), \\ \beta_s &\sim \mathcal{MVN}(0, \Sigma_{sc}). \end{aligned}$$

Model 3:

$$\begin{aligned} y_i &= \text{logit}^{-1}(\alpha + \beta_c X_c + \alpha_{s[i]} + \beta_{p[i]} c_i) \\ \alpha_s &\sim \mathcal{N}(0, \sigma_s^2), \\ \beta_p &\sim \mathcal{MVN}(0, \Sigma_{pc}). \end{aligned}$$

The expected response y for the observation $i = 1, \dots, I$ for respondent $p = 1, \dots, P$ on stimulus $s = 1, \dots, S$ in condition $c = 1, \dots, C$:

Model 1:

$$y_i = \text{logit}^{-1}(\alpha + \beta_c X_c + \alpha_{p[i]} + \alpha_{s[i]})$$
$$\alpha_p \sim \mathcal{N}(0, \sigma_p^2),$$
$$\alpha_s \sim \mathcal{N}(0, \sigma_s^2).$$

Model 2:

$$y_i = \text{logit}^{-1}(\alpha + \beta_c X_c + \alpha_{p[i]} + \beta_{s[i]} c_i)$$
$$\alpha_p \sim \mathcal{N}(0, \sigma_p^2),$$
$$\beta_s \sim \mathcal{MVN}(0, \Sigma_{sc}).$$

Model 3:

$$y_i = \text{logit}^{-1}(\alpha + \beta_c X_c + \alpha_{s[i]} + \beta_{p[i]} c_i)$$
$$\alpha_s \sim \mathcal{N}(0, \sigma_s^2),$$
$$\beta_p \sim \mathcal{MVN}(0, \Sigma_{pc}).$$

Fixed Effects

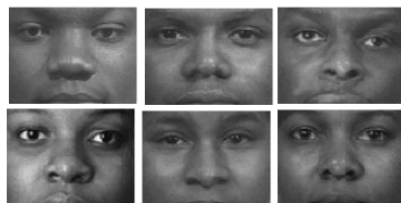
Random structure

12 Object stimuli

White people faces



Black people faces



16 Attribute stimuli

Positive attributes

Good, laughter, pleasure, glory, peace,
 happy, joy, love

Negative attributes

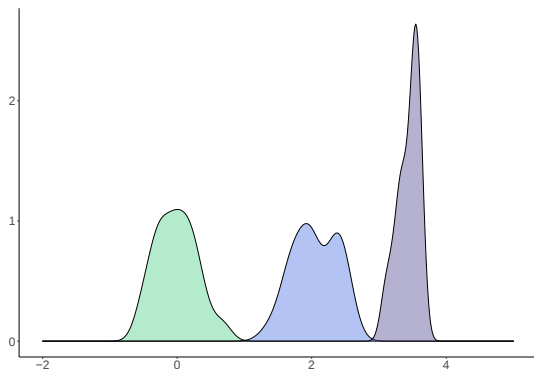
Evil, bad, horrible, terrible, nasty,
 pain, failure, hate

Participants: 62 ($F = 48.39\%$, Age = 24.92 ± 2.11 years)

Model 2 is the least wrong model

$$y_i = \text{logit}^{-1}(\alpha + \beta_c X_c + \alpha_{p[i]} + \beta_{s[i]} c_i)$$

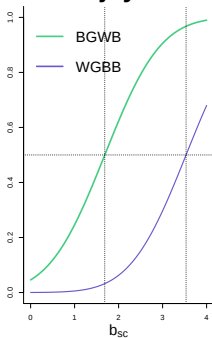
θ_p b_{BGWB} b_{WGBB}



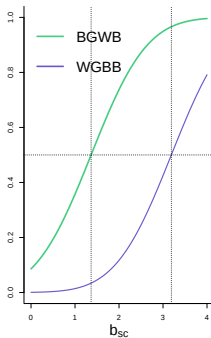
Condition-specific easiness

HIGHLY CONTRIBUTING STIMULI

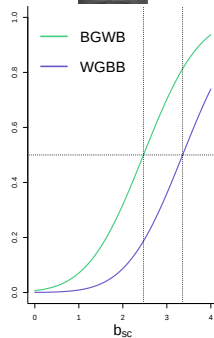
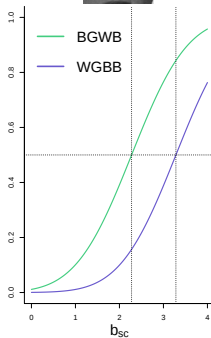
joy



evil



LOWLY CONTRIBUTING STIMULI



Discussion

Further Information

Epifania, O. M., Anselmi, P., & Robusto, E. (2024). A guided tutorial on linear mixed-effects models for the analysis of accuracies and response times in experiments with fully crossed design. *Psychological Methods*. doi: <https://doi.org/10.1037/met0000708>



Thank you!

ottavia.epifania@unitn.it

