

BRIEF-2 Report Write-Up Cheat Sheet

Reviewed by the BastionGPT Clinical Advisory Board · July 2026

The BRIEF-2 (Behavior Rating Inventory of Executive Function, Second Edition, PAR, 2015; updated 2026) is a multi-informant rating scale of everyday executive function for ages 5–18: parent and teacher forms plus a self-report from age 11, nine clinical scales under three indexes (BRI, ERI, CRI) and a Global Executive Composite. T-scores (mean 50, SD 10); higher means more reported difficulty. Two validated norm sets since 2026: name the forms, raters, and norm set every time.

1 – Measures, forms, informants, and norms which edition and raters produced these scores?

Do: name the edition, each form and who completed it, the norm set (age-and-sex or combined-sex), scoring platform, dates.

Pitfall: "a BRIEF was completed"; the BRIEF, BRIEF-2, and adult BRIEF2A even use different index architectures.

2 – Validity scales first before any score, for every rater

Do: report Inconsistency, Negativity, and Infrequency per rater, then the action taken on any flag.

Pitfall: clean validity scales read as proof the profile is accurate; they screen response style, not truth.

3 – Index and composite results BRI · ERI · CRI, then the GEC convention

Do: give T-scores with percentiles; band language in prose; state the GEC convention when indexes diverge.

Pitfall: leading with the GEC across divergent indexes; the summary score describes no ability the student shows.

4 – Scale-level drivers the scales doing the work, nothing more

Do: name the clinical scales driving each elevated index, one functional example each.

Pitfall: the nine-scale score dump; the scales share variance, so pick the drivers.

5 – Cross-informant comparison expect modest agreement (mean $r = .28$)

Do: report each rater separately; write divergence as setting and self-awareness information.

Pitfall: averaging raters into one number, or crowning one rater accurate.

6 – Integration with performance testing complementary levels, median $r = .19$

Do: write ratings as everyday goal pursuit and tests as structured efficiency; interpret divergence.

Pitfall: "ecologically valid" as a trump card, or intact test scores used to dismiss the ratings.

7 – Functional-impairment translation what the elevation looks like in class

Do: convert elevations into observable behavior; adverse effect and substantial limitation live in the narrative.

Pitfall: "elevated Working Memory scale" offered as the impairment; the scale is the measure, not the behavior.

8 – Interpretive summary and recommendations answer the referral question

Do: consistent-with language; state the not-standalone and no-single-measure limits; tie recommendations to elevated scales.

Pitfall: the elevation-equals-ADHD leap, or a T-score cutoff written as if some authority set one.

Before you sign

- ☐ Edition, forms, raters, and norm set named
- ☐ Validity scales reported for every rater
- ☐ T-scores carry percentiles; one band vocabulary
- ☐ Divergence interpreted, not averaged
- ☐ Ratings and tests kept complementary
- ☐ Each recommendation traces to an elevated scale

Drafting with AI

Paste your index and scale summary: raters, T-scores with percentiles, validity results. Not the items, not the software narrative.

"Draft a BRIEF-2 results section from this summary: parent CRI T 73 (Working Memory 74, Plan/Organize 71), teacher CRI T 70, self CRI T 61; ERI average for all raters; validity acceptable; referral question: ..."

Never paste client-identifying information into consumer AI tools. BastionGPT is HIPAA-compliant with a signed BAA on every plan.

300–700 words
results section

~10 min per form
5-min screeners

Ages 5–18
parent · teacher · self (11+)

IEP teams · evaluators
who reads it

BRIEF-2 Results Section Template

Copy or print for your practice · Adapt fields to your organization's, payer's and jurisdiction's documentation requirements.

Student (initials): _____ Age / grade: _____ Rating dates: _____

Edition / norm set: _____ Scoring platform: _____ Referral question: _____

Measures and informants

Each form administered; each rater's role, relationship, setting, and date completed

Validity scales

Inconsistency, Negativity, Infrequency per rater; action taken on any flag

Index and composite results

BRI, ERI, CRI, GEC per rater: T-score with percentile; band language in prose; GEC convention stated

Scale-level drivers

Clinical scales driving each elevated index; one functional example each

Cross-informant comparison

Convergence, divergence, and what the divergence means; no averaging, no crowned rater

Integration with performance testing

Complementary levels of analysis; convergence and divergence interpreted, not explained away

Interpretive summary

Answer the referral question; consistent-with language; what the ratings do NOT establish

Recommendations linkage

Each recommendation tied to a specific elevated scale; monitoring plan

GEC convention stated: _____ Re-rating due (raters): _____

Evaluator signature / credentials: _____ Date signed: _____

Supervisor signature, where required: _____ Date: _____