INTERNATIONAL JOURNAL OF
EDUCATION, PSYCHOLOGY
AND COUNSELLING
(IJEPC)www.ijepec.comTRANSLATION AND VALIDATION OF NARROW
PERSONALITY TRAITS FROM IPIP AMONG MALAYSIAN
CIVIL SERVANTSMuhammad Faez Zakaria¹, Harris Shah Abd Hamid², Nurul Shahnaz Ahmad Mahdzan^{3*}¹ Institute Advance Studies, University of Malaya (UM), Malaysia
Email: muhfeez@gmail.com.my² Faculty of Management, Education & Humanities, University College MAIWP International (UCMI), Malaysia,
Email: drharris@ucmi.edu.my³ Faculty of Business and Economic, University of Malaya (UM), Malaysia
Email: n_shahnaz@um.edu.my

* Corresponding Author

Article Info:**Article history:**

Received date: 01.05.2024

Revised date: 20.05.2024

Accepted date: 05.06.2024

Published date: 13.06.2024

To cite this document:

Zakaria, M. F., Abd Hamid, H. S., & Mahdzan, N. S. (2024). Translation And Validation of Narrow Personality Traits from IPIP Among Malaysian Civil Servants. *International Journal of Education, Psychology and Counseling*, 9 (54), 271-289.

DOI: 10.35631/IJEPC.954020**This work is licensed under** [CC BY 4.0](https://creativecommons.org/licenses/by/4.0/)**Abstract:**

Numerous psychological instruments can be used to assess personality, including those from the IPIP website. Portions from the IPIP measures had been translated and validated in Malaysia with unclear validity evidence. Additionally, narrow traits (from the broad five factors) had not been examined in the Malaysian civil servant's context. As part of a research project examining over-indebtedness among civil servants, this study attempted to validate 11 narrow traits (89 items). Back translation into Bahasa Melayu involved adapting the items to be culturally appropriate. Through expert validation, pre-testing, and a survey of Malaysian civil servants, it was determined that the narrow traits demonstrated good internal consistency, albeit with some items removed. Exploratory Factor Analysis based on data from 134 civil servants and Composite Reliability as well as Discriminant Validity (SmartPLS Analysis) based on data from 490 civil servants provided evidence that the 11 narrow traits have satisfactory structural validity. However, the traits tend to be structurally separated based on wording direction (positive vs negative). It can be concluded that the Bahasa Melayu version is appropriate for measuring the personality of Malaysian civil servants.

Keywords:

IPIP, Narrow Traits, Validation, Malaysian Civil Servants

Introduction

Psychological testing and evaluation in psychology help researchers delve deeper into people's issues and problems, allowing relevant experts to develop appropriate treatment plans and interventions. Assessment aids in collecting information about an individual by observing behaviour under specific conditions and against specific simulations. The psychological assessment uses objective methods like standardised tests and less objective methods like observations and interviews. Objective tests have been frequently used to investigate abilities, but it was not until the writings of Cattell and Warburton that these instruments were also methodically created to research personality (Santacreu, (2024).

Psychological testing is the systematic and scientific process of evaluating or assessing an individual's behaviour. In other words, the psychological assessment significantly contributes valuable information to understanding individual characteristics and attributes by accumulating, incorporating, and interpreting data about an individual (Groth-Marnat, 2009).

According to the American Psychological Association, each person has a unique personality that is reflected in their distinctive ways of thinking, feeling, and behaving. This helps to distinguish the individual differences between people (BW Robert, 2022). Allport and other 'personality experts' claimed that understanding personality traits is the most significant way to explain differences between people. Personality traits represent fundamental dimensions individuals vary (Matthew, Deary, & Whiteman, 2003). Personality traits become more stable with age, peaking during early life and leveling off in young adulthood. This trend aligns with increased maturity and valuing individual differences is crucial (Bleidorn, 2022). Personality studies focus on two broad areas: the first is an understanding of individual differences in personality traits, such as togetherness or irritability. The second area is to comprehend how the various components of a person work together to form a whole.

Personality traits imply consistency and stability. The IPIP Big-Five factors' (B5F) items were chosen to measure the Big Five dimensions and facets with as few items as feasible while maintaining appropriate levels of validity and reliability. The Big Five model identifies five major personality traits: Extraversion, Agreeableness, Conscientiousness, Neuroticism, and Openness. To measure someone's personality, a series of questionnaires are presented with a broad range of items. However, due to the complexity of human personality, thorough inventories are necessary (JO Ogunsemi, 2022). The Big Five's enormous scope is one of its limitations. Facets, or narrow traits nested inside domains, frequently account for more result variance than the Big Five domains and give higher specificity to personality-outcome correlations (Paunonen & Ashton, 2001). Growing research demonstrates that nuance-level models are more specific and predictive than facets or domain-based models (Elleman et al., 2020). The instrument's scales showed strong internal reliability ($M = 0.83$) and convergent validity with Goldberg's (1992) adjectives and Costa and McCrae's (1992) NEO Five-Factor Inventory (NEO-FFI; John & Srivastava, 1999). To fully understand the potential of the Big Five model, practical instruments need to be developed in various languages.

The aim of this research is to produce a Bahasa Melayu version of the IPIP B5F. Thus, the specific objectives are (1) to evaluate the Bahasa Melayu translation with subject matter experts, and (2) to evaluate the psychometric properties of the BM version. The researchers assessed the psychometric properties of the IPIP B5F items in Bahasa Melayu using two main

criteria: (a) the degree of correspondence between Bahasa Melayu translation's factor structure and the original English translations; and (b) the internal consistency of the IPIP B5F items.

Methodology

Instruments

The researcher referred to questions adapted and modified from the International Personality Item Pool (IPIP), which was also used by Johnson (2014), among others, to assess 11 narrow traits of the Five Factor Model using an 89-item public domain collection which involved Openness: Ideas (10 items) and Actions (8 items), Conscientiousness: Deliberation (6 items), Competence (7 items), Extraversion: Assertiveness (9 items), Excitement (7 items), Agreeableness: Trust (6 items), Compliance (8 items) and Neuroticism: Self-Consciousness (8 items), Vulnerability (10 items), Impulsiveness (10 items). The International Personality Item Pool's official website, accessible at <http://ipip.ori.org>, boasts an impressive collection of over 3,000 items and more than 250 scales derived from these items. This comprehensive resource is widely recognized for its diverse applications, notably in the context of the Big Five personality traits (Goldberg, 2001). This item is derived from 100 unipolar Big-Five factor indicators and employs a four-point Likert scale. The responses are aggregated to form a single composite score, providing a quantitative assessment of a character or personality trait.

Ethical considerations, such as the permission of the University of Malaya's ethics committee, clearance from administrators of the Public Service Department, anonymity, informed consent, and withdrawal from the study, were all observed.

Translation

The translation process consists of two stages: (1) analyzing the original text and its meaning, and (2) re-expressing the meaning in the target language using the words or sentences received (Rietveld, 2019). Hence the researchers used two methods; the first word-for-word translation as the initial translation process to ensure the original meaning is preserved and researchers can understand the mechanism of the source language. Two professional translators had independently translated the 109 items of the original English item from the IPIP B5F, as stated in the early stages of translation using the word-for-word translation method. The next step is to adapt these items to achieve the original language's goals while also adhering to local cultural norms. Five expert judges reached a unanimous decision on the best translation.

Table 1
Summary Section of Narrow Personality Traits Items.

<i>Broad Facets</i>	<i>Narrow Trait</i>	<i>Description</i>	<i>No of Items</i>
Openness	Ideas	Narrow Traits	10
	Actions	Personality is described	8
Conscientiousness	Deliberation	as an individual's	6
	Competence	specific innate tendency	7
Extraversion	Assertiveness	to respond cognitively,	9
	Excitement	affectively, and	7
Agreeableness	Trust	behaviourally to various	6
	Compliance	situations and life	8
Neuroticism	Self-Consciousness	events.	8
	Vulnerability		10
	Impulsiveness		10

Source <http://ipip.ori.org/>. Internet Web Site Goldberg (2006)

Face and Content Validation

For this study, the researcher consulted ten experts, scholars, and lecturers in the disciplines of psychology (psychometrics and personality) to determine the instrument's validity. Content validity is the appropriateness of test items to the test content to be measured, where the content can evaluate mastery, domain skills, and the respondent's comprehension of the things to be tested. Face validity experts are selected based on two criteria: competence in developing psychological instruments, psychometric or psychological testing, and direct involvement in the practice of psychology for more than five years, as well as English and Malay proficiency. To enhance face validity, some modifications were made based on the recommendations of the consulted experts, including the use of relevant terminology from the outset. A panel of expert reviews of content validity reveals that the indicators of the various notions are truly distinct (Hair, Gabriel, & Patel, 2014). Therefore, it is possible to conclude that the study's face and content validity requirements were met.

Exploratory Factor Analysis

In this research, both the primary test and the pilot test were administered. The pilot study involved distributing a set of questions to 200 participants using Google Forms, and receiving 134 completed questionnaires in return. Due to the optimal quantity and selection, these participants represent a part of the population and exhibit comparable traits (Lis' Nacre, 2005; Mohd Majid, 2005; Chua, 2006). We administered the revised Malay version of IPIP B5F and 11-NPT to Malaysian civil officials, with the assistance of human resource departments from several Malaysian government ministries, to analyze narrow personality traits. We have substantially improved the instrument's reliability based on this analysis's findings.

JAMOV V2.0 software was used to analyse data from the pilot test. To investigate the factorial structure of the Big Five Factors, which consists of 11 facets. Each of the 134 items in the instrument underwent exploratory factor analysis with oblique rotation. According to Table 3 indicated value for Bartlett's test of sphericity, the correlation structure is suitable for factor analysis. Using maximum likelihood factor analysis with a cut-off point of 0.40 and the Kaiser's criteria of eigenvalues larger than 1. The correlation structure is deemed suitable for factor analysis, and the factors solution was determined to be the best match for the data.

The overall KMO value for each construct is 0.70, meeting the desired score as per Hoelzle & Meyer (2013) and Lloret et al. (2017). Values below 0.50 are unequivocally unacceptable, as stated by Child (2006), Hair et al. (2010), and Kaiser (1974), indicating that the correlation matrix is not factorable. Based on the data and analysis of the pilot study, the researcher determined that 25 items did not meet the study's measuring requirements and should be eliminated.

Quantitative content validity, on the other hand, maintains confidence in the selection of the most important and accurate information in an instrument, as measured by the content validity ratio (CVR). In this manner, experts are asked whether a specific item in a set of things is required to operate a construct. For that purpose, they are asked to rate each item from 1 to 4 on a four-point scale of "Not Relevant," "Somewhat Relevant," "Quite Relevant," and "Highly Relevant." The range for the content validity ratio is between 1 and -1. The higher the score, the more panellists concur on the significance of an instrument's item. The content validity ratio formula is $CVR = (N_e - N/2)/(N/2)$, where N_e is the number of panellists who indicated "Highly Relevant", and N is the total number of panellists. The Lawshe Table determines the content validity ratio's numerical value. In our study with ten panellists, for instance, if the CVR is greater than 0.6, the item with an acceptable significance level will be accepted. Following the theoretical definitions of the construct and its dimension, panel members are asked to rate instrument items on a 4-point ordinal scale regarding comprehensibility and relevance to the construct under study. It is the most frequently reported method for establishing content validity in instrument development reports.

Following a series of translations, adaptations, expert content validation, and quantification of the content validity ratio (CVR), the psychometric properties of the translated IPIP B5F and 11-NPT were measured using four different types of analysis: item analysis, domain and facet reliability, exploratory factor analysis (EFA) and measurement model analysis. On the data from the studies, item analyses, factor analyses, and analyses of facets reliabilities were performed.

Considering our validation criteria, we needed a sample that was:

1. Large enough to obtain a stable factor solution
2. Diverse in terms of demographic background
3. Large enough to compare factor solutions across age groups

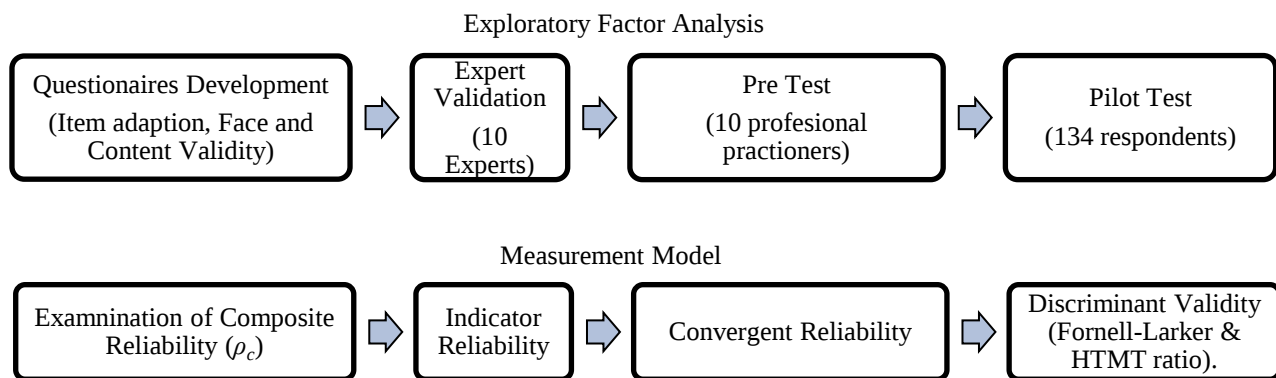
Internet studies are a quick and easy approach to collecting extensive samples, and they are at least as representative as the convenience samples commonly employed in psychology (Sosling, Vazire, Stivastava, & John, 2004).

Sample: Measurement Modelling

In 2021, during the widespread effects of the COVID-19 pandemic and the implementation of the Malaysian government's Movement Control Order (MCO), a survey was conducted using a Google Form questionnaire. Despite the challenges posed by the lockdown, 492 out of 513 respondents completed the survey through convenience sampling, which was the only viable method available at the time. A minimum and representative sample for evaluating the translated instrument was successfully gathered using the e-survey.

The measurement model (validity and reliability of measurements) was examined in this work using Anderson and Gerbing's (1988) multistage analytical approaches. (Putting theories to the test.) The researcher employed the bootstrapping approach to determine the importance of the route coefficients and loadings (5000 resamples). This study's and research methodology's variables are all multi-item constructs regarded as reflecting rather than formative. The reflecting construct seeks inter-correlated, unidimensional metrics with strong internal consistency. The outer model in PLS-SEM is a component of a route model that includes the indicators and their relationships with the constructs (Hair, Hult, et al., 2017).

Visual Process of Questionnaires Development and Validation



Measurement model tested on 492 respondent

Data Analysis Procedure

To ensure a robust measurement model, high internal consistency and intercorrelated measures within the reflecting construct are essential. The outer model in PLS-SEM connects indicators with constructs in the route model. Detailed steps for accessing the measurement model are provided, and measures illustrate the effects of an underlying concept. The relationship between a concept and its measures is termed as causality in PLS-SEM (Hair, Hult, et al., 2017).

A first assessment is carried out to ascertain reliability and internal consistency. The tests are conducted by the researcher using the Composite Reliability Index and Cronbach's Alpha. The composite dependability for all constructions should more than the minimal cut-off value of 0.7. This internal consistency reliability metric does not need equal indicator loadings, in contrast to Cronbach's alpha (Hair, Hult, et al., 2017). These results are important due to ensure that the measuring technique was sufficiently trustworthy.

This stage then moves on to convergent validity. According to Urbach & Ahlemann (2010), convergent validity is the degree to which certain indicators accurately reflect the constructs in comparison to indicators measuring other constructs. It is demonstrated when there is a strong correlation between the scores obtained from two different instruments measuring the same concept (Sekaran & Bougie, 2016). Convergent validity is determined by the use of the Average Variance Extracted (AVE). The AVE measures how well a latent construct matches the variance of its indicators (Hair, Hult, et al., 2017). $AVE \geq 0.5$ (Fornell & Larcker, 1981; Gefen, Straub, & Boudreau, 2000; Hair, Hult, Ringle, & Sarstedt, 2014)

Next, the determination of discriminant validity involves three (3) types of tests. The first is a comparison of cross-loadings (table 6), which examines the correlations between indicators (items) and other model components. To show discriminant validity, the indicator's (item's) outer loading (table 7) on the associated construct must be bigger than any of its cross-loadings on other constructs (Hair et al., 2014; Hair, Hult, et al., 2017). It indicates that all of the indicator's outer loading has met the loading threshold score. Secondly, Fornell and Larcker's criteria (table 8) is a discriminant validity measure that compares the square root of each construct's AVE to its correlations with the other constructs in the model. Thus, the Fornell-Larcker results on discriminant validity in this study fit the requirement since discriminant validity exists when the square root of a construct's AVE is greater than its correlation with other constructs in the same model.

And then, the Heterotrait-Monotrait (HTMT) Ratio is the final test, it estimates the underlying correlation between two constructs assuming they were fully assessed (totally reliable) (Hair, Hult, et al., 2017). HTMT is the average of all indicators across constructs measuring different constructs compared to the average correlations of indicators measuring the same construct (Henseler, Ringle, & Sarstedt, 2015). The HTMT approach introduced by Henseler et al. (2015) is used to examine discriminant validity. The "ratio of between-trait correlations to within-trait correlations" is known as HTMT (Hair et al., 2017). This study employed two ways to determine discriminant validity. If the HTMT value exceeds 0.85 (Kline, 2011), discriminant validity issues exist.

Results

Participants reported their age, gender, marital status, and educational level after completing the translated IPIP B5F and 11-NPT. In the first-level exam sample, there were 134 Malaysian civil servants, 57.5 per cent of whom were men and 45.2 per cent women. Nearly seventy-six per cent of responders are between the ages of 31 and 40, while only 14 per cent are between 18 and 30. Seventy-seven per cent of those polled were married, while the remaining 22.4% were classified as either categorically single or still in high school. A total of 492 Malaysian public servants took the second-level test. The demographic details show that 165 were male (33.5%) and 327 were female (66.5%). In terms of age, the majority were between 31 and 40 years (44.1%), followed by 41 to 50 years (32.7%), 51 years and over (11.8%), and 18 to 30 years (11.4%). Of the 492 respondents, 295 had bachelor's degrees or postgraduate studies (60%), 95 had a high school certificate (SPM) or vocational certifications (19.3%), and 102 had diplomas (20.76%). Eighty-two percent of the respondents were married, and 85.2% had served for more than five years. Furthermore, 265 respondents (53.9%) and 227 respondents (46.1%) belonged to the supporting staff group and management and professional group, (table 2).

Table 2
Participants' Demographics

	Variables	Frequency (f)	Percentage (%)
Gender	Male	165	33.5
	Female	327	66.5
Age	18 – 30	56	11.4
	31 – 40	217	44.1
	41 – 50	161	32.7
	51 and above	58	11.8
Length of Service	Below 5 years	73	14.8
	6 years above	419	85.2
Academic Level	Postgrad/ Bachelor's Degree	295	60
	Diploma/ STPM/STAM	102	20.7
	SPM/Vocational/Skills	95	19.3
	Certificate		
Marital Status	Single	75	15.2
	Married	400	81.3
	Single Mother/Father/Divorcee	17	3.4
Service Group	Management & Professional	227	46.1
	Supporting Staff	265	53.9
Total		492	100.0

In pilot stage, items were subjected to reliability analysis. The reliability of a measuring instrument in measuring a concept in a study measures its accuracy and stability (Creswell, 2012). According to Mohd Faizal and Leow (2017), reliability is essential to determine whether a questionnaire item should be retained or removed. Din et al. (2009) stated that the higher the reliability value of the questionnaire, the more precise and reliable the data obtained. Based on the pilot study's data and analysis, the researcher determined that 16 items did not meet the study's measurement standards and should be removed. A summary of factor analyses on the narrow traits is shown in Table 3 below.

Table 3
Total Items Retained and Dropped

Construct [No. of Item]	Bartlett's Test of Sphericity	KMO Measure of Sampling Adequacy	Parallel Analysis	Eigenvalue	Item Deleted	Cronbach's α
Ideas [10]	$\chi^2_{360} < .001$ df (45) p	0.752	2 Factors	1 Factor	Nul	0.783
Actions [10]	$\chi^2_{289} < .001$ df (28) p	0.693	2 Factors	1 Factors	No.5,8	0.738
Deliberation [8]	$\chi^2_{488} < .001$ df (15) p	0.833	1 Factor	1 Factor	No.1,2	0.864
Competence [8]	$\chi^2_{362} < .001$ df (21) p	0.773	3 Factors	1 Factor	No.8	0.801

Assertiveness [10]	χ^2_{362} <.001	df (36)	p	0.729	3 Factors	2 Factors	No.9	0.766
Excitement [10]	χ^2_{249} <.001	df (21)	p	0.736	3 Factors	1 Factor	No.5,9,10	0.761
Trust [9]	χ^2_{209} <.001	df (15)	p	0.690	3 Factors	1 Factor	No.1,2,3	0.729
Compliance [10]	χ^2_{810} <.001	df (28)	p	0.887	1 Factor	1 Factor	No.1,2	0.913
Self-Consciousness [10]	χ^2_{297} <.001	df (28)	p	0.751	3 Factors	1 Factor	No.3,9	0.774
Vulnerability [10]	χ^2_{686} <.001	df (45)	p	0.810	3 Factors	2 Factors	Nul	0.857
Impulsiveness [10]	χ^2_{576} <.001	df (45)	p	0.709	3 Factors	2 Factors	Nul	0.796

As mentioned, Cronbach's alpha and Exploratory Factor Analysis were used to evaluate the new translated IPIP B5F and 11-NPT for reliability. The value obtained from each construct were all >0.7, which indicated high reliability.

Then Smart PLS 3.0 Measurement Model analysis was used to assess the second phase for this item validation. Ringle, Wende, and Becker (2015). Hair et al. (2019) and Purwanto. The initial examination is carried out to determine internal consistency and reliability. The researcher again used Cronbach's Alpha and the Composite Reliability Index to execute the experiments. Table 4 shows that the Cronbach alpha values in this study range from 0.600 to 0.940. The composite reliability should be more than 0.7 to demonstrate appropriate internal consistency. Hair, Ringle, and Sarstedt (2011). The results also show that for each batch of data, the composite reliability for all constructions surpassed the minimal cut-off value of 0.7, ranging from 0.700 to 0.950. Unlike Cronbach's alpha, this measure of internal consistency reliability does not necessitate equal indicator loadings (Hair, Hult, et al., 2017). These data project that the measuring approach was sufficiently reliable.

Table 4
Construct Reliability and Validity

	<i>Cronbach's Alpha</i>	<i>rho_A</i>	<i>ρ_c</i>	<i>Average Variance Extracted (AVE)</i>
ACT	0.794	0.822	0.865	0.616
AST	0.733	0.810	0.816	0.530
COM	0.881	1.007	0.907	0.661
CPL	0.944	0.966	0.954	0.721
DEL	0.931	0.955	0.947	0.783
EXC	0.607	0.626	0.834	0.716
IDE	0.753	0.818	0.829	0.551
IMP	0.817	0.833	0.870	0.575
SCO	0.850	0.861	0.892	0.623
TRU	0.649	0.710	0.846	0.735
VUL	0.814	0.832	0.877	0.643

The indicator reliability is assessed after establishing the reliability of each internal consistency reliability. As seen in table 6 (Cross/Outer Loading), all items have satisfactory indicator reliability (ranging from 0.560 to 0.980), with all AVE ratings more than 0.5 exceeding the Byrne (2016) threshold value. The variance derived from the item reveals how much of an item's variation is explained by the concept (Hair, Hult, et al., 2017). However, as shown in table 5, several items fell short of the satisfactory AVE score; the result resulted in the deletion of 14 entries, and the following items have been removed.

Table 5
Deleted Item

No	Construct	Item Deleted
1	Ideas	IDE03
2	Action	ACT04
3	Excitement	EXC03, EXC04, EXC05
4	Assertiveness	AST05
5	Trust	TRU01, TRU02, TRU05
6	Deliberation	DEL0
7	Self-Conscientiousness	SCO05, SCO07
8	Vulnerability	VUL05
9	Ideas	IDE03

According to the result, each batch of data's composite reliability—which ranged from 0.700 to 0.950—exceeded the minimal cut-off value of 0.7. This internal consistency reliability metric does not need equal indicator loadings, in contrast to Cronbach's alpha (Hair, Hult, et al., 2017).

Likewise, the purpose of studying indicator reliability is to discover the extent to which an indicator or set of indications is consistent with what it is supposed to signify (Urbach & Ahlemann, 2010). The indicator reliability denotes the fraction of indicator variance explained by the hidden variable. Loading values equal to or more than 0.5 are acceptable for indicator reliability if the aggregate of loadings results in high loading scores, contributing to an AVE score greater than 0.5. (Byrne 2016). SmartPLS 3.3.3 calculates the AVE value using the PLS Algorithm, and table 6 provides the AVE values for all constructs and items. All constructs had AVE values greater than 0.5 for each set of data. The lowest AVE value given is for Assertiveness (0.530) followed by Ideas (0.551) and impulsiveness (0.575). Excitement (0.530). Meanwhile, Action (0.616) Self-consciousness (0.623), Vulnerability (0.643) and Competence (0.661). Following that, Excitement (0.716), Compliance (0.721), Trust (0.735) and Deliberation score at (0.783). If the value of AVE is more than 0.5 and explains at least 50% of the variation in the supplied indicators (Chin, 2010a; Hair, Hult, Ringle, & Sarstedt, 2017), the measurement model has adequate convergent validity.

Meanwhile the criteria proposed by Fornell and Larcker (1981) (Table 8) relates the model correlations of each construct to its square root of AVE. Since each concept has discriminant validity when the square root of its AVE is bigger than its correlation with other constructs in

the same model, the results of Fornell and Larcker's (1981) discriminant validity research match the necessary criteria.

When using the PLS Algorithm, none of the associated constructs, as indicated in Table 9, contradict the HTMT value, suggesting that construct validity is attained in the measurement.

Table 6
Cross/Outer Loading

Column	ACT	ASR	COM	CPL	DEL	EXC	IDE	IMP	SCS	TRU	VUL
ACT01	0.700										
ACT02	0.580										
ACT03	0.760										
ACT05	0.720										
ASR01		0.620									
ASR02		0.650									
ASR03		0.600									
ASR04		0.640									
COM01			0.740								
COM02			0.810								
COM03			0.840								
COM04			0.830								
COM05			0.800								
CPL01				0.770							
CPL02				0.890							
CPL03				0.700							
CPL04				0.840							
CPL05				0.810							
CPL06				0.880							
CPL07				0.920							
CPL08				0.900							
DEL01					0.880						
DEL02					0.920						
DEL03					0.890						
DEL04					0.590						
DEL05					0.800						
DEL06					0.870						
EXC01						0.820					
EXC02						0.560					
IDE06							0.620				
IDE07							0.600				
IDE08							0.600				
IDE09							0.730				
IDE10							0.700				

Column	ACT	ASR	COM	CPL	DEL	EXC	IDE	IMP	SCS	TRU	VUL
IMP01								0.650			
IMP02								0.780			
IMP03								0.790			
IMP04								0.780			
IMP05								0.730			
SCS01									0.740		
SCS02									0.730		
SCS03									0.750		
SCS04									0.790		
SCS05									0.650		
SCS06									0.770		
TRU03										0.980	
TRU04										0.730	
VUL01											0.660
VUL02											0.720
VUL03											0.740
VUL04											0.720

Table 7
Results Summary for Reflective Measurement Models

Constructs	Items	Indicator Reliability	Convergent Validity	Internal Consistency Reliability	
		Outer Loadings >0.60	AVE >0.50	ρ_c >0.7	α >0.7
ACT	ACT01	0.733	0.616	0.865	0.794
	ACT02	0.843			
	ACT03	0.822			
	ACT05	0.737			
ASR	ASR01	0.629	0.530	0.816	0.733
	ASR02	0.659			
	ASR03	0.757			
	ASR04	0.847			
COM	COM01	0.750	0.661	0.907	0.881
	COM02	0.813			
	COM03	0.850			
	COM04	0.840			
	COM05	0.809			

Constructs	Items	Indicator Reliability	Convergent Validity	Internal Consistency Reliability	
		Outer Loadings >0.60	AVE >0.50	ρ_c >0.7	α >0.7
CPL	CPL01	0.781	0.721	0.954	0.944
	CPL02	0.899			
	CPL03	0.713			
	CPL04	0.852			
	CPL05	0.803			
	CPL06	0.883			
	CPL07	0.926			
	CPL08	0.912			
DEL	DEL01	0.882	0.783	0.947	0.931
	DEL02	0.930			
	DEL03	0.901			
	DEL05	0.813			
	DEL06	0.892			
EXC	EXC01	0.882	0.716	0.834	0.607
	EXC02	0.809			
IDE	IDE01	0.702	0.551	0.829	0.753
	IDE02	0.618			
	IDE04	0.855			
IMP	IDE05	0.772	0.575	0.870	0.817
	IMP01	0.654			
	IMP02	0.791			
	IMP03	0.800			
	IMP04	0.788			
SCS	IMP05	0.747	0.623	0.892	0.850
	SCO01	0.809			
	SCO02	0.809			
	SCO03	0.787			
	SCO04	0.797			
TRU	SCO06	0.742	0.735	0.846	0.649
	TRU03	0.800			
	TRU04	0.911			
VUL	VUL01	0.686	0.643	0.877	0.814
	VUL02	0.883			
	VUL03	0.864			
	VUL04	0.758			

Table 8
Fornell & Larcker (1981) Criterion

		1	2	3	4	5	6	7	8	9	10	11
1	ACT	0.785										
2	ASR	0.098	0.728									
3	COM	0.440	0.222	0.813								
4	CPL	0.167	0.213	0.181	0.849							
5	DEL	0.146	0.358	0.286	0.426	0.885						
6	EXC	0.509	0.167	0.447	0.136	0.107	0.846					
7	IDE	0.142	0.268	0.214	0.098	0.200	0.182	0.742				
8	IMP	-0.031	-0.391	-0.164	-0.347	-0.521	-0.036	-0.102	0.758			
9	SCS	0.374	0.133	0.422	0.285	0.347	0.372	0.209	-0.325	0.789		
10	TRU	-0.056	-0.045	-0.128	-0.157	-0.118	-0.079	-0.026	0.068	-0.193	0.857	
11	VUL	-0.288	-0.255	-0.545	-0.266	-0.416	-0.341	-0.231	0.457	-0.534	0.286	0.802

A bootstrapping test was also utilised to assess whether the HTMT value is significantly different from 1.00 (Henseler, Ringle, & Sarstedt, 2015), as suggested by Hair, Risher, Sarstedt, & Ringle (2019). If the confidence interval contains the value, discriminant validity is lacking (Henseler, Ringle, & Sarstedt, 2015). More specifically, none of the upper bounds of the 95 per cent confidence interval of HTMT is more than 0.85 or 0.9, as indicated in the table. The results are less than the needed threshold values of HTMT.85 (Kline, 2011) and HTMT.90 (Gold et al., 2001), demonstrating that discriminant validity for the constructs in this investigation has been established. As indicated in table 9, when applying the PLS Algorithm, none of the associated constructs contradicts HTMT, showing that construct validity is achieved in the measurement model.

Table9
HTMT Results

		1	2	3	4	5	6	7	8	9	10	11
1	ACT											
2	ASR	0.135										
3	COM	0.553	0.265									
4	CPL	0.178	0.234	0.178								
5	DEL	0.168	0.427	0.295	0.442							
6	EXC	0.739	0.224	0.618	0.172	0.144						
7	IDE	0.185	0.357	0.243	0.123	0.240	0.257					
8	IMP	0.097	0.490	0.175	0.384	0.585	0.085	0.138				
9	SCS	0.459	0.140	0.463	0.310	0.382	0.513	0.250	0.362			
10	TRU	0.100	0.089	0.160	0.174	0.155	0.151	0.093	0.126	0.238		
11	VUL	0.357	0.304	0.629	0.288	0.476	0.470	0.296	0.535	0.623	0.366	

Discussion

After going through different stages of studies, the objective of producing a Bahasa Melayu version of the instrument was achieved as evident by the findings from the expert review and online surveys. The translated items show good psychometric properties as per the PLS SEM analysis. It was found that the psychometric evidence to support the use of the instrument among Malaysian civil servants is sufficient.

There are limitations for discussing the findings from the present study due to the lack of comparable studies using narrow personality traits. Some comparisons can be made to the other versions of IPIP instrument tested by researchers in Malaysia. For example, in the present study, satisfactory factor structure was achieved without the need to item-parcelling as done upon the Mini IPIP among individuals with drug abuse problem in Malaysia (Leong et al., 2019). Combining several items measuring broad personality traits and turning them into indicators may help to improve model fit, however it would introduce more complexities into the interpretation of the scores. This is especially pertinent for measures of narrow personality traits as the items should have less divergence into overlapping constructs.

Compared to the Big Five IPIP (50 items) which was tested with a smaller ($n=112$) student sample (Heng et al., 2017), the present study had a better psychometric performance. In their study, Heng et al. had to delete many items and the entire Extraversion items were removed due to low factor loading. This excessive removal of items might be due to the inappropriateness of the items for university students in Malaysia or due to the relatively small sample size. The problems of low factor loading were not observed to the same extent in the present study.

Further test of the translated version would be desirable considering the lack of studies that tested narrow personality traits. The psychometric performance of the NPT should be tested with different sub-populations to examine measurement invariance. Furthermore, the functionality of the rating scale used should be examined. Abd Hamid (2004) found that there are national differences in the meaning of rating scale categories used by Malaysians, Indonesians and Singaporean for measuring ethical values. To what extent does the difference exist within Malaysia sub-population? This detailed analysis using Rasch Rating Scale Model would further add to our understanding of the cultural and sub-cultural influences in the measurement of personality.

In conclusion, the present study has added to the depth and richness of the IPIP items by producing a Bahasa Melayu version of items capable of measuring narrow personality traits. Overall, this study demonstrates a rigorous adaptation and validation procedure, ensuring that the instrument is not only linguistically appropriate but also psychometrically sound for usage among Malaysian civil servants. This supports the growth of research and practice in the area by giving a helpful tool for examining numerous facets within this population, which then can be extended to other populations to allow us to investigate many diverse viewpoints.

Acknowledgment

We would like to offer our deepest gratitude to the participants and those who took the time and effort to complete the questionnaire for this research. Not to mention the Public Service Department, which helped and funded us to complete this study. Thank you very much.

References

- Abdullah, M. F. N. L., & Wei, L. T. (2017). Kesahan dan kebolehppercayaan instrumen penilaian sendiri pembelajaran geometri tingkatan satu. *Malaysian Journal of Learning and Instruction*, 14(1), 211
- Abd Hamid, H. S. (2024). Functionality of Rating Categories for a Measure of Ethical Values. In H. S. Abd Hamid. (Ed). Application of Rasch Rating Scale Model (pp 97-102). Mind Metric Solution.
- Aithal, A., & Aithal, P. S. (2020). Development and validation of survey questionnaire & experimental data—a systematical review-based statistical approach. *International Journal of Management, Technology, and Social Sciences (IJMTS)*, 5(2), 233-251.
- Almanasreh, E., Moles, R., & Chen, T. F. (2019). Evaluation of methods used for estimating content validity. *Research in social and administrative pharmacy*, 15(2), 214-221.
- American Psychological Association (Ed.). Washington, DC: American Psychology.
- Ben-Porath, Y. S., & Waller, N. G. (1992). Five big issues in clinical personality assessment: A rejoinder to Costa and McCrae.
- Bleidorn, W., Schwaba, T., Zheng, A., Hopwood, C. J., Sosa, S. S., Roberts, B. W., & Briley, D. A. (2022). Personality stability and change: A meta-analysis of longitudinal studies. *Psychological bulletin*, 148(7-8), 588.
- Child, D. (2006). *The essentials of factor analysis*. A&C Black.
- Chin, W. W. (1998). Commentary: Issues and opinion on structural equation modeling. *MIS quarterly*, vii-xvi.
- Elleman, L. G., Condon, D. M., Holtzman, N. S., Allen, V. R., & Revelle, W. (2020). Smaller is better: Associations between personality and demographics are improved by examining narrower traits and regions. *Collabra: Psychology*, 6(1), 17210.
- Field, A. P., & Miles, J. (2009). Discovering statistics using SPSS:(and sex and drugs and rock'n'roll).
- Fornell, C., & Larcker, D. F. (1981). Structural equation models with unobservable variables and measurement error: Algebra and statistics.
- Gefen, D., Straub, D., & Boudreau, M. C. (2000). Structural equation modeling and regression: Guidelines for research practice. *Communications of the association for information systems*, 4(1), 7.
- Goldberg, L. R. (1999). A broad-bandwidth, public domain, personality inventory measuring the lower- level facets of several five-factor models. *Personality psychology in Europe*, 7(1), 7-28.
- Goldberg, L. R., Johnson, J. A., Eber, H. W., Hogan, R., Ashton, M. C., Cloninger, C. R., & Gough, H. G. (2006). The international personality item pool and the future of public-domain personality measures. *Journal of Research in personality*, 40(1), 84-96.
- Gosling, S. D., Vazire, S., Srivastava, S., & John, O. P. (2004). Should we trust web-based studies? A comparative analysis of six preconceptions about internet questionnaires. *American psychologist*, 59(2), 93.
- Groth-Marnat, G. (2009). *Handbook of psychological assessment*. John Wiley & Sons.
- Habidin, N. F., Omar, C. M. Z. C., Kamis, H., Latip, N. A. M., & Ibrahim, N. (2012). Confirmatory factor analysis for lean healthcare practices in Malaysian healthcare industry. *Journal of Contemporary Issues and Thought*, 2, 17-29.
- Hair, J. F., Gabriel, M., & Patel, V. (2014). AMOS covariance-based structural equation modelling (CB-SEM): Guidelines on its application as a marketing research tool. *Brazilian Journal of Marketing*, 13(2).

- Hair, J. F., Hult, G. T. M., Ringle, C. M., Sarstedt, M., & Thiele, K. O. (2017). Mirror, mirror on the wall: a comparative evaluation of composite-based structural equation modeling methods. *Journal of the academy of marketing science*, 45(5), 616-632.
- Hair, J. F., Risher, J. J., Sarstedt, M., & Ringle, C. M. (2019). When to use and how to report the results of PLS-SEM. *European business review*, 31(1), 2-24.
- Henseler, J., Ringle, C. M., & Sarstedt, M. (2015). A new criterion for assessing discriminant validity in variance-based structural equation modeling. *Journal of the academy of marketing science*, 43(1), 115-135.
- Heng, G. W. Y., Hui, C. Y., Ting, T. P., Yin, C. T., & Wider, W. (2017). The Relationship Between Birth Order and The Big-Five Personality Dimensions Among Psychology Students in Universiti Malaysia Sabah. *Jurnal Psikologi dan Kesihatan Sosial (JPsiKS)*, 1, 71-75.
- Hoelzle, J. B., & Meyer, G. J. (2013). Exploratory factor analysis: Basics and beyond.
- John, O. P., & Srivastava, S. (1999). The Big Five trait taxonomy: History, measurement, and theoretical perspectives. *Handbook of personality: Theory and research*, 2(1999), 102-138.
- Kline, R. B. (2011). Convergence of structural equation modeling and multilevel modeling.
- Leong, F. W., Mohd Yasin, M. A., Muhd Ramli, E. R., Fadzil, N. A., & Kueh, Y. C. (2019). Validation of the Malay version of mini-IPIP among substance use disorder patients attending methadone clinics in Malaysia. *International journal of environmental research and public health*, 16(22), 4434.
- Matthews, G., Deary, I. J., & Whiteman, M. C. (2003). *Personality traits*. Cambridge University Press.
- McCrae, R. R., & Costa Jr, P. T. (2004). A contemplated revision of the NEO Five-Factor Inventory. *Personality and individual differences*, 36(3), 587-5
- Ogunsemi, J. O., Akinlawo, E. O., Akinbobola, O. I., Ariyo, J. O., Babatunde, S. I., & Akpunne, B. C. (2022). Psychometric properties and validation of Mini-International Personality Item Pool (mini- IPIP) among Nigerian population. *Advances in Research*, 23(4), 49-57.
- Paunonen, S. V., & Ashton, M. C. (2001). Big Five factors and facets and the prediction of behavior. *Journal of Personality and Social Psychology*, 81(3), 524-539.
- Purwanto, A. (2021). Partial least squares structural equation modeling (PLS-SEM) analysis for social and management research: a literature review. *Journal of Industrial Engineering & Management Research*.
- Ringle, C. M., Wende, S., & Becker, J. M. (2015). SmartPLS 3. SmartPLS GmbH, Boenningstedt. *Journal of Service Science and Management*, 10(3), 32-49.
- Rietveld, L., & Van Harmelen, F. (2019). Use of vocabulary translation strategies: A semantic translation analysis. *Applied Translation*, 13(2), 1-7.
- Roberts, B. W., & Yoon, H. J. (2022). Personality psychology. *Annual review of psychology*, 73, 489-516.
- Santacreu, J., Romero, M., Casadevante, C., & Hernández, J. M. (2024). How to Design an ObjectiveTest to Assess Personality: Step by Step. *Authorea Preprints*.
- Sekaran, U., & Bougie, R. (2016). *Research methods for business: A skill building approach*. John Wiley & Sons.
- Thabane, L., Ma, J., Chu, R., Cheng, J., Ismaila, A., Rios, L. P., ... & Goldsmith, C. H. (2010). A tutorial on pilot studies: the what, why and how. *BMC medical research methodology*, 10(1), 1-10.

Urbach, N., & Ahlemann, F. (2010). Structural equation modeling in information systems research using partial least squares. *Journal of Information Technology Theory and Application (JITTA)*, 11(2), 2.