

Tab	Summary measure name	Ordered / Non-ordered dimension	Absolute / Relative measure	Simple / Complex measure	Weighted / Unweighted measure	Unit of measure	Value of no inequality	Interpretation	Notation	Summary measure formula	Variance formula	Step-by-step summary measure calculation	Notes
p-1	Difference	Ordered	Absolute	Simple	Unweighted	Unit of indicator	Zero	For subgroups within an ordered dimension of inequality, D measures the difference between the most- and least-advantaged subgroup estimates. The larger the absolute value of D, the higher the level of inequality.	y_1 = estimate for subgroup 1 y_2 = estimate for subgroup 2 σ_1 = standard error for subgroup 1 σ_2 = standard error for subgroup 2	$D = y_1 - y_2$	$var(D) = \sigma_1^2 + \sigma_2^2$	1. Identify the two subgroups, y_1, and y_2 [see notes]. 2. Subtract the y_2 estimate from the y_1 estimate.	So that the result of D tends to be positive (making interpretation easier), the calculation depends on the indicator type: <i>Favourable indicator:</i> Most-advantaged subgroup – Least-advantaged subgroup <i>Adverse indicator:</i> Least-advantaged subgroup – Most-advantaged subgroup The variable <i>favourable_indicator</i> identifies whether the indicator is favourable (1) or adverse (0). See the glossary for the definition of favourable and adverse indicators. So that the result of R tends to be greater than 1 (making interpretation easier), the calculation depends on the indicator type: <i>Favourable indicator:</i> Most-advantaged subgroup / Least-advantaged subgroup <i>Adverse indicator:</i> Least-advantaged subgroup / Most-advantaged subgroup The variable <i>favourable_indicator</i> identifies whether the indicator is favourable (1) or adverse (0). See the glossary for the definition of favourable and adverse indicators.
r-1	Ratio	Ordered	Absolute	Simple	Unweighted	No unit	One	For subgroups within an ordered dimension of inequality, R measures the ratio between the most- and least-advantaged subgroup estimates. R assumes only positive values. The further the value of R from 1, the higher the level of inequality. Note that because ratio is a multiplicative measure, it should be reported on a logarithmic scale (for example, a ratio of 2 is equivalent to a ratio of 0.5).	y_1 = estimate for subgroup 1 y_2 = estimate for subgroup 2 σ_1 = standard error for subgroup 1 σ_2 = standard error for subgroup 2	$R = y_1 / y_2$	$var(R) = (1/y_2)^2 [\sigma_1^2 + (y_1/y_2)^2 \sigma_2^2]$	1. Identify the two subgroups, y_1, and y_2 [see notes]. 2. Divide the y_1 estimate by the y_2 estimate.	<i>Favourable indicator:</i> Most-advantaged subgroup / Least-advantaged subgroup <i>Adverse indicator:</i> Least-advantaged subgroup / Most-advantaged subgroup The variable <i>favourable_indicator</i> identifies whether the indicator is favourable (1) or adverse (0). See the glossary for the definition of favourable and adverse indicators.
ac1	Absolute concentration index	Ordered	Absolute	Complex	Weighted	Unit of indicator	Zero	The larger the absolute value of ACI, the higher the level of inequality. Positive values indicate a concentration of the indicator among advantaged subgroups, and negative values indicate a concentration of the indicator among disadvantaged subgroups.	p_j = population share of subgroup j out of the total population X_j = relative rank of subgroup j y_j = estimate for subgroup j σ_j = standard error for subgroup j	$ACI = \sum_j p_j (2X_j - 1) y_j$	$var(ACI) = \sum_j p_j^2 \sigma_j^2 (2X_j - 1)^2$	1. Sort the subgroups from the least-advantaged to the most-advantaged. 2. Calculate the population share of each subgroup out of the total population. 3. Calculate the cumulative total of the population share of each subgroup. 4. Calculate the relative rank of each subgroup by calculating the midpoint of each subgroup (0.5*popshare) and subtracting this from the cumulative population share. 5. Calculate the population share of each subgroup multiplied by two times the subgroup's relative rank minus one, multiplied by the subgroup estimate. 6. Calculate ACI as the sum of results from step 5.	This measure is only suitable when subgroups can be meaningfully ranked in order. The subgroups should be sorted from the least advantaged to the most-advantaged.
rc1	Relative concentration index	Ordered	Relative	Complex	Weighted	No unit	Zero	RCI is the relative version of ACI, dividing ACI by the setting average and multiplying the result by 100 to make it more easily interpretable. The larger the absolute value of RCI, the higher the level of inequality. Positive values indicate a concentration of the indicator among advantaged subgroups, and negative values indicate a concentration of the indicator among disadvantaged subgroups.	p_j = population share of subgroup j out of the total population X_j = relative rank of subgroup j y_j = estimate for subgroup j μ = weighted mean (setting average)	$RCI = 100 * (\sum_j p_j (2X_j - 1) y_j) / \mu$	$var(RCI) = (\sum_j p_j^2 \sigma_j^2 (2X_j - 1)^2) / \mu^2$	1. Sort the subgroups from the least-advantaged to the most-advantaged. 2. Calculate the population share of each subgroup out of the total population. 3. Calculate the cumulative total of the population share of each subgroup. 4. Calculate the relative rank of each subgroup by calculating the midpoint of each subgroup (0.5*popshare) and subtracting this from the cumulative population share. 5. Calculate the population share of each subgroup multiplied by two times the subgroup's relative rank minus one, multiplied by the subgroup estimate. 6. Calculate RCI as the sum of results from step 5 divided by the weighted mean of the subgroup estimates, times 100.	This measure is only suitable when subgroups can be meaningfully ranked in order. The subgroups should be sorted from the least advantaged to the most-advantaged.
si	Slope index of inequality	Ordered	Absolute	Complex	Weighted	Unit of indicator	Zero	SI measures the difference in predicted values of an indicator for the most-advantaged and most-disadvantaged subgroups, obtained by fitting a regression model. The larger the absolute value of SI, the higher the level of inequality. Positive values indicate that the level of the indicator is higher among advantaged subgroups, while negative values indicate that the level of the indicator is higher among disadvantaged subgroups.	y_i = predicted value of indicator estimate at rank 1 y_v = predicted value of indicator estimate at rank 0	$SI = y_v - y_i$	See notes	1. Sort the subgroups from the least-advantaged to the most-advantaged. 2. Calculate the population share of each subgroup out of the total population. 3. Calculate the cumulative total of the population share of each subgroup. 4. Calculate the relative rank of each subgroup (X) by calculating the midpoint of each subgroup (0.5*popshare) and subtracting this from the cumulative population share. 5. The health indicator estimates are converted to a 0-1 scale in preparation for the subsequent steps. The scaling factor to convert to a 0-1 scale is based on the maximum health indicator estimate. 6. The health indicator estimates are transformed to log odds (i.e., a logit transformation) [see notes]. 7. The square root of the weight variable (population) is calculated (W). 8. The weighted relative rank is calculated as X*W. 9. The weighted (logit transformed) health indicator estimate is calculated as Y*W. 10. The weighted health indicator (YW) is regressed against the weight (W) and the weighted relative rank (XW), without a constant term. The predicted log odds at rank 0 (the least-advantaged) is equal to the intercept (i.e. the second beta coefficient of the regression model). The predicted log odds at rank 1 (the most-advantaged) is equal to the sum of the intercept (i.e. the second beta coefficient) and the slope (i.e. the first beta coefficient). 11. Log odds are converted back to the indicator scale using the inverse logit formula, multiplied by the indicator scaling factor. 12. Calculate SI as the difference between the indicator estimates at rank 1 (y_i) and rank 0 (y_v).	This measure is only suitable when subgroups can be meaningfully ranked in order. The subgroups should be sorted from the least advantaged to the most-advantaged. Simple linear regression assumes a linear relationship between the health indicator and the population ranking (which is not always the case) and can result in estimated values outside of a 0-100% interval (which is inaccurate for indicators measured as percentages). Using logistic regression can sometimes solve these problems. In this workbook, a weighted least squares (WLS) regression method is used, with a logit transformation of the health indicator. With a logit transformation, the estimated values from the regression model will be bounded between 0 and 1 (or 0-100%), which is preferable for many indicators. The weighted (logit transformed) indicator estimates are regressed against the weighted relative rank of each subgroup and the values for the least-advantaged (rank 0) and most-advantaged (rank 1) groups are estimated from the regression model, taking all other subgroups into account. The predicted values are then converted back to the original indicator scale. Note that the SI results may differ from those produced using more sophisticated modelling techniques with statistical software. If 95% confidence intervals are required, upload the dataset to the Health Equity Assessment Toolkit Plus (HEAT Plus) software (available at www.who.int/data/inequality-monitor/assessment_toolkit_plus) or use statistical software package such as Stata or R (codes available at www.who.int/data/inequality-monitor/toolkit_plus/responses/statistical_codes).
ri	Relative index of inequality	Ordered	Relative	Complex	Weighted	No unit	One	RI measures the ratio between the predicted values of an indicator for the most-advantaged and most-disadvantaged subgroups, obtained by fitting a regression model. The further the value of RI from 1, the higher the level of inequality. Values larger than 1 indicate the level of the indicator is higher among advantaged subgroups, and values lower than 1 indicate the level of the indicator is higher among disadvantaged subgroups. Note that because RI is a multiplicative measure, it should be reported on a logarithmic scale (for example, a value of 2 is equivalent to a value of 0.5).	y_i = predicted value of indicator estimate at rank 1 y_v = predicted value of indicator estimate at rank 0	$RI = y_i / y_v$	See notes	1. Sort the subgroups from the least-advantaged to the most-advantaged. 2. Calculate the population share of each subgroup out of the total population. 3. Calculate the cumulative total of the population share of each subgroup. 4. Calculate the relative rank of each subgroup (X) by calculating the midpoint of each subgroup (0.5*popshare) and subtracting this from the cumulative population share. 5. The health indicator estimates are converted to a 0-1 scale in preparation for the subsequent steps. The scaling factor to convert to a 0-1 scale is based on the maximum health indicator estimate. 6. The health indicator estimates are transformed to log odds (i.e., a logit transformation) [see notes]. 7. The square root of the weight variable (population) is calculated (W). 8. The weighted relative rank is calculated as X*W. 9. The weighted (logit transformed) health indicator estimate is calculated as Y*W. 10. The weighted health indicator (YW) is regressed against the weight (W) and the weighted relative rank (XW), without a constant term. The predicted log odds at rank 0 (the least-advantaged) is equal to the intercept (i.e. the second beta coefficient of the regression model). The predicted log odds at rank 1 (the most-advantaged) is equal to the sum of the intercept (i.e. the second beta coefficient) and the slope (i.e. the first beta coefficient). 11. Log odds are converted back to the indicator scale using the inverse logit formula, multiplied by the indicator scaling factor. 12. Calculate RI as the ratio between the indicator estimates at rank 1 (y_i) and rank 0 (y_v).	This measure is only suitable when subgroups can be meaningfully ranked in order. The subgroups should be sorted from the least advantaged to the most-advantaged. Simple linear regression assumes a linear relationship between the health indicator and the population ranking (which is not always the case) and can result in estimated values outside of a 0-100% interval (which is inaccurate for indicators measured as percentages). Using logistic regression can sometimes solve these problems. In this workbook, a weighted least squares (WLS) regression method is used, with a logit transformation of the health indicator. With a logit transformation, the estimated values from the regression model will be bounded between 0 and 1 (or 0-100%), which is preferable for many indicators. The weighted (logit transformed) indicator estimates are regressed against the weighted relative rank of each subgroup and the values for the least-advantaged (rank 0) and most-advantaged (rank 1) groups are estimated from the regression model, taking all other subgroups into account. The predicted values are then converted back to the original indicator scale. Note that the RI results will differ from those produced using more sophisticated modelling techniques with statistical software. If 95% confidence intervals are required, upload the dataset to the Health Equity Assessment Toolkit Plus (HEAT Plus) software (available at www.who.int/data/inequality-monitor/assessment_toolkit_plus) or use statistical software package such as Stata or R (codes available at www.who.int/data/inequality-monitor/toolkit_plus/responses/statistical_codes).
par-1	Population attributable risk	Ordered	Absolute	Complex	Weighted	Unit of indicator	Zero	For ordered dimensions of inequality, PAR shows the potential improvement in the average of an indicator, in absolute terms, that could be achieved if all population subgroups had the same level of the indicator as the most-advantaged subgroup. PAR assumes only positive values for favourable indicators and only negative values for adverse indicators. The larger the absolute value, the higher the level of inequality.	y_{ref} = estimate for a reference subgroup μ = weighted mean (setting average) PAF = population attributable fraction PAF_{ref} = standard error of population attributable fraction	$PAR = y_{ref} - \mu$	$var(PAR) = ([1/\mu] [PAF + 1.96 PAF_{ref}] - (PAF - 1.96 PAF_{ref}) / ((2*1.96)^2)$	1. Identify the reference subgroup, y_{ref} [see notes]. 2. Calculate the population share of each subgroup out of the total population. 3. Calculate the weighted mean of the subgroup estimates. 4. Calculate PAF as the difference between the reference subgroup estimate and the weighted mean.	For ordered dimensions, the reference subgroup (y_{ref}) is the most-advantaged subgroup (identified as having the largest subgroup_order value). If the indicator is favourable (favourable_indicator=1) and PAR has a negative value, PAR is equal to zero. If the indicator is adverse (favourable_indicator=0) and PAR has a positive value, PAR is equal to zero. See the glossary for the definition of favourable and adverse indicators. The variable <i>indicator_scale</i> is required for the calculation of 95% confidence intervals and identifies the scale of the indicator. See the glossary for the definition of indicator scale.
faf-1	Population attributable fraction	Ordered	Relative	Complex	Weighted	No unit	Zero	For ordered dimensions of inequality, PAF shows the potential improvement in the average of an indicator, in relative terms, that could be achieved if all population subgroups had the same level of the indicator as the most-advantaged subgroup. PAF assumes only positive values for favourable indicators and only negative values for adverse indicators. The larger the absolute value of PAF, the larger the level of inequality.	y_{ref} = estimate for a reference subgroup μ = weighted mean (setting average) a = number of people not in the reference subgroup for whom the indicator was achieved b = number of people not in the reference subgroup for whom the indicator was not achieved c = number of people in the reference subgroup for whom the indicator was achieved d = number of people in the reference subgroup for whom the indicator was not achieved	$PAF = 100 * (y_{ref} - \mu) / \mu$	$var(PAF) = 100^2 [(a/(bN - c) + c^2) / (a + c)^2 + d^2]$	1. Identify the reference subgroup, y_{ref} [see notes]. 2. Calculate the population share of each subgroup out of the total population. 3. Calculate the weighted mean of the subgroup estimates. 4. Calculate PAF as the difference between the reference subgroup estimate and the weighted mean, divided by the weighted mean, times 100.	For ordered dimensions, the reference subgroup (y_{ref}) is the most-advantaged subgroup (identified as having the largest subgroup_order value). If the indicator is favourable (favourable_indicator=1) and PAF has a negative value, PAF is equal to zero. If the indicator is adverse (favourable_indicator=0) and PAF has a positive value, PAF is equal to zero. See the glossary for the definition of favourable and adverse indicators. The variable <i>indicator_scale</i> identifies the scale of the indicator. See the glossary for the definition of indicator scale.

Tab	Summary measure name	Ordered / Non-ordered dimension	Absolute / Relative measure	Simple / Complex measure	Weighted / Unweighted measure	Unit of measure	Value of no inequality	Interpretation	Notation	Summary measure formula	Variance formula	Step-by-step summary measure calculation	Notes
d-2	Difference	Non-ordered	Absolute	Simple	Unweighted	Unit of indicator	Zero	For subgroups within an non-ordered dimension of inequality, D measures the difference between the two selected subgroup estimates. The larger the absolute value of D, the higher the level of inequality.	y_1 = estimate for subgroup 1 y_2 = estimate for subgroup 2 σ_1 = standard error for subgroup 1 σ_2 = standard error for subgroup 2	$D = y_1 - y_2$	$var(D) = \sigma_1^2 + \sigma_2^2$	- Identify the two subgroups, y_1 and y_2 (see notes). - Subtract the y_2 estimate from the y_1 estimate.	The selection of the two subgroups depends on the characteristics of the inequality dimension, the purpose of the analysis and the indicator type: <i>No reference subgroup / favourable indicator</i>: Highest estimate – lowest estimate <i>No reference subgroup / adverse indicator</i>: Lowest estimate – highest estimate <i>Reference subgroup / favourable indicator</i>: Reference group estimate – lowest estimate <i>Reference subgroup / adverse indicator</i>: Highest estimate – Reference group estimate The variable <i>favourable_indicator</i> identifies whether the indicator is favourable (1) or adverse (0). See the glossary for the definition of favourable and adverse indicators. The variable <i>reference_subgroup</i> can optionally be used to identify a reference subgroup with the value 1.
r-2	Ratio	Non-ordered	Absolute	Simple	Unweighted	No unit	One	For subgroups within an ordered dimension of inequality, R measures the ratio between two selected subgroup estimates. R assumes only positive values. The further the value of R from 1, the higher the level of inequality. Note that because ratio is a multiplicative measure, it should be reported on a logarithmic scale (for example, a ratio of 2 is equivalent to a ratio of 0.5).	y_1 = estimate for subgroup 1 y_2 = estimate for subgroup 2 σ_1 = standard error for subgroup 1 σ_2 = standard error for subgroup 2	$R = y_1 / y_2$	$var(R) = (1/y_2^2)\sigma_1^2 + (y_1/y_2^3)\sigma_2^2$	- Identify the two subgroups, y_1 and y_2 (see notes). - Divide the y_1 estimate by the y_2 estimate.	The selection of the two subgroups depends on the characteristics of the inequality dimension, the purpose of the analysis and the indicator type: <i>No reference subgroup</i>: Highest estimate / Lowest estimate <i>Reference subgroup / favourable indicator</i>: Reference group estimate / Lowest estimate <i>Reference subgroup / adverse indicator</i>: Highest estimate / Reference group estimate The variable <i>favourable_indicator</i> identifies whether the indicator is favourable (1) or adverse (0). See the glossary for the definition of favourable and adverse indicators. The variable <i>reference_subgroup</i> can optionally be used to identify a reference subgroup with the value 1.
mdm	Mean difference from mean	Non-ordered	Absolute	Complex	Weighted / unweighted	Unit of indicator	Zero	MDM shows the mean difference between each subgroup and the setting average. It has only positive values, with larger values indicating higher levels of inequality.	n = number of subgroups y_i = estimate for subgroup i μ = weighted mean (setting average) p_i = population share of subgroup i out of the total population	$MDMU = 1/n * \sum_i y_i - \mu $ $MDMW = \sum_i p_i y_i - \mu $	See notes	- Calculate the population share of each subgroup out of the total population. - Calculate the weighted mean of the subgroup estimates. - Calculate MDMU as the absolute difference between each estimate and the weighted mean, and then sum the results and divide by the total number of subgroups. - Calculate MDMW as the absolute difference between each estimate and the weighted mean, multiplied by the population share, and then sum the results.	95% confidence intervals are not calculated in this workbook, as they require a simulation methodology. If 95% confidence intervals are required, upload the dataset to the Health Equity Assessment Toolkit Plus (HEAT Plus) software (available at www.who.int/data/inequality-monitor/assessment_toolkitplus) or use statistical software package such as Stata or R (codes available at www.who.int/data/inequality-monitor/toolkitplus/resources/statistical_codes).
mdb	Mean difference from best-performing subgroup	Non-ordered	Absolute	Complex	Weighted / unweighted	Unit of indicator	Zero	MDB shows the mean difference between each subgroup and the best-performing subgroup. It has only positive values, with larger values indicating higher levels of inequality.	n = number of subgroups y_i = estimate for subgroup i y_{best} = highest subgroup estimate for favourable indicator or lowest subgroup estimate for adverse indicators p_i = population share of subgroup i out of the total population	$MDBU = 1/n * \sum_i y_i - y_{best} $ $MDBW = \sum_i p_i y_i - y_{best} $	See notes	- Identify the best-performing subgroup estimate, y_{best}. This is the highest subgroup estimate for favourable indicators or lowest subgroup estimate for adverse indicators. Use the variable <i>favourable_indicator</i> to identify whether an indicator is favourable (1) or adverse (0). - Calculate the population share of each subgroup out of the total population. - Calculate MDBU as the absolute difference between each estimate and the best-performing subgroup estimate, and then sum the results and divide by the total number of subgroups. - Calculate MDBW as the absolute difference between each estimate and the best-performing subgroup estimate, multiplied by the population share, and then sum the results.	MDM can be calculated as an unweighted or weighted measure. For the unweighted version (MDMU) all subgroups are weighted equally. For the weighted version (MDMW) subgroups are weighted according to their population share. MDR can be calculated as an unweighted or weighted measure. For the unweighted version (MDBU) all subgroups are weighted equally. For the weighted version (MDBW) subgroups are weighted according to their population share. The best-performing subgroup is the subgroup with the highest estimate for favourable indicators, or the subgroup with the lowest estimate for adverse indicators. The variable <i>favourable_indicator</i> identifies whether the indicator is favourable (1) or adverse (0). See the glossary for the definition of favourable and adverse indicators. 95% confidence intervals are not calculated in this workbook, as they require a simulation methodology. If 95% confidence intervals are required, upload the dataset to the Health Equity Assessment Toolkit Plus (HEAT Plus) software (available at www.who.int/data/inequality-monitor/assessment_toolkitplus) or use statistical software package such as Stata or R (codes available at www.who.int/data/inequality-monitor/toolkitplus/resources/statistical_codes).
mdr	Mean difference from a reference subgroup	Non-ordered	Absolute	Complex	Weighted / unweighted	Unit of indicator	Zero	MDR shows the mean difference between each subgroup and a selected reference subgroup. It has only positive values, with larger values indicating higher levels of inequality.	n = number of subgroups y_i = estimate for subgroup i y_{ref} = reference subgroup estimate p_i = population share of subgroup i out of the total population	$MDRU = 1/n * \sum_i y_i - y_{ref} $ $MDRW = \sum_i p_i y_i - y_{ref} $	See notes	- Identify the best-performing subgroup estimate, y_{ref}. This is the highest subgroup estimate for favourable indicators or lowest subgroup estimate for adverse indicators. Use the variable <i>reference_subgroup</i> to identify a reference subgroup with the value 1. - Calculate the population share of each subgroup out of the total population. - Calculate MDRU as the absolute difference between each estimate and the reference subgroup estimate, and then sum the results and divide by the total number of subgroups. - Calculate MDWR as the absolute difference between each estimate and the reference subgroup estimate, multiplied by the population share, and then sum the results.	MDR can be calculated as an unweighted or weighted measure. For the unweighted version (MDRU) all subgroups are weighted equally. For the weighted version (MDRW) subgroups are weighted according to their population share. The variable <i>reference_subgroup</i> is used to identify a reference subgroup with the value 1. 95% confidence intervals are not calculated in this workbook, as they require a simulation methodology. If 95% confidence intervals are required, upload the dataset to the Health Equity Assessment Toolkit Plus (HEAT Plus) software (available at www.who.int/data/inequality-monitor/assessment_toolkitplus) or use statistical software package such as Stata or R (codes available at www.who.int/data/inequality-monitor/toolkitplus/resources/statistical_codes).
ids	Index of disparity	Non-ordered	Relative	Complex	Weighted / unweighted	No unit	Zero	IDIS is the relative version of MDM, dividing MDM by the setting average and multiplying the result by 100 to make it more easily interpretable. It has only positive values, with larger values indicating higher levels of inequality.	n = number of subgroups y_i = estimate for subgroup i μ = weighted mean (setting average) p_i = population share of subgroup i out of the total population	$IDISU = 100 * (1/n * \sum_i y_i - \mu) / \mu$ $IDISW = 100 * (\sum_i p_i y_i - \mu) / \mu$	See notes	- Calculate the population share of each subgroup out of the total population. - Calculate the weighted mean of the subgroup estimates. - Calculate IDISU as the absolute difference between each estimate and the weighted mean, and then sum the results and divide by the total number of subgroups. Multiply the result by 100. - Calculate IDISW as the absolute difference between each estimate and the weighted mean, multiplied by the population share, and then sum the results. Multiply the result by 100.	IDIS can be calculated as an unweighted or weighted measure. For the unweighted version (IDISU) all subgroups are weighted equally. For the weighted version (IDISW) subgroups are weighted according to their population share. 95% confidence intervals are not calculated in this workbook, as they require a simulation methodology. If 95% confidence intervals are required, upload the dataset to the Health Equity Assessment Toolkit Plus (HEAT Plus) software (available at www.who.int/data/inequality-monitor/assessment_toolkitplus) or use statistical software package such as Stata or R (codes available at www.who.int/data/inequality-monitor/toolkitplus/resources/statistical_codes).
bgv	Between-group variance	Non-ordered	Absolute	Complex	Weighted	Squared unit of indicator	Zero	BGV has only positive values, with larger values indicating higher levels of inequality.	p_i = population share of subgroup i out of the total population y_i = estimate for subgroup i μ = weighted mean (setting average) σ^2 = standard error for subgroup i	$BGV = \sum_i p_i (y_i - \mu)^2$	$var(BGV) = \sum_i p_i^2 \sigma_i^2 (y_i - \mu)^2 + 2(\sum_i p_i \sigma_i^2 (y_i - \mu) - (\sum_i p_i \sigma_i^2)^2) / \mu^2$	- Calculate the population share of each subgroup out of the total population. - Calculate the weighted mean of the subgroup estimates. - Subtract the weighted mean from each subgroup estimate and square these differences. - Calculate BGV as the sum of the results from step #3.	BGV is more sensitive to outlier estimates as it gives more weight to the estimates that are further from the setting average.
bgstd	Between-group standard deviation	Non-ordered	Absolute	Complex	Weighted	Unit of indicator	Zero	BGSD is calculated as the square root of between-group variance (BGV). It has only positive values, with larger values indicating higher levels of inequality.	p_i = population share of subgroup i out of the total population y_i = estimate for subgroup i μ = weighted mean (setting average) n = number of subgroups	$BGSD = \sqrt{\sum_i p_i (y_i - \mu)^2}$	See notes	- Calculate the population share of each subgroup out of the total population. - Calculate the weighted mean of the subgroup estimates. - Subtract the weighted mean from each subgroup estimate and square these differences. - Calculate BGSD as the square root of the sum of the results from step #3.	Since BGSD is the squared root of BGV, it is reported in the unit of the indicator, which may be more easily interpretable than BGV and will not be as sensitive to outlier estimates that are further from the setting average. 95% confidence intervals are not calculated in this workbook, as they require a simulation methodology. If 95% confidence intervals are required, upload the dataset to the Health Equity Assessment Toolkit Plus (HEAT Plus) software (available at www.who.int/data/inequality-monitor/assessment_toolkitplus) or use statistical software package such as Stata or R (codes available at www.who.int/data/inequality-monitor/toolkitplus/resources/statistical_codes).
cov	Coefficient of variation	Non-ordered	Relative	Complex	Weighted	No unit	Zero	COV is the relative version of BGSD, dividing BGSD by the setting average and multiplying the result by 100 to make it more easily interpretable. It only has positive values, with larger values indicating higher levels of inequality.	p_i = population share of subgroup i out of the total population y_i = estimate for subgroup i μ = weighted mean (setting average) n = number of subgroups	$COV = 100 * (\sum_i p_i y_i - \mu) / \mu$	See notes	- Calculate the population share of each subgroup out of the total population. - Calculate the weighted mean of the subgroup estimates. - Subtract the weighted mean from each subgroup estimate and square these differences. - Calculate COV as the square root of the sum of the results from step #3, divided by the weighted mean, times 100.	COV calculates BGSD in relative terms, by dividing BGSD by the setting average. 95% confidence intervals are not calculated in this workbook, as they require a simulation methodology. If 95% confidence intervals are required, upload the dataset to the Health Equity Assessment Toolkit Plus (HEAT Plus) software (available at www.who.int/data/inequality-monitor/assessment_toolkitplus) or use statistical software package such as Stata or R (codes available at www.who.int/data/inequality-monitor/toolkitplus/resources/statistical_codes).
ti	Theil Index	Non-ordered	Relative	Complex	Weighted	No unit	Zero	TI expresses inequality as a function of how the subgroup's share of the health indicator compares to their share of the population. The larger the value of TI, the greater the level of inequality.	p_i = population share of subgroup i out of the total population l_i = share of the indicator in each subgroup, equal to the estimate for subgroup i divided by μ σ_i = standard error for subgroup i μ = weighted mean (setting average)	$TI = 1000 * \sum_i p_i (l_i \ln(l_i))$	$var(TI) = (\sum_i p_i \sigma_i^2 (l_i + \ln(l_i)) - (\sum_i p_i \sigma_i^2 (l_i + \ln(l_i)))^2) / \mu^2$	- Replace any estimates equal to zero with a very small number (0.000001) to avoid issues in calculating the log of zero. - Calculate the population share of each subgroup out of the total population. - Calculate the weighted mean of the subgroup estimates. - Calculate the share of the indicator in each subgroup by dividing the estimate of each subgroup by the weighted mean. - Multiply the share of the indicator in each subgroup by the natural logarithm (log) of the share of the indicator in each subgroup and the population share. - Calculate TI as the sum of the results from step #5, multiplied by 1000 for easier interpretation.	TI is more sensitive to differences further from the setting average (by the use of the logarithm).
mid	Mean log deviation	Non-ordered	Relative	Complex	Weighted	No unit	Zero	The larger the absolute value of MID, the higher the level of inequality.	p_i = population share of subgroup i out of the total population l_i = share of the indicator in each subgroup, equal to the estimate for subgroup i divided by μ σ_i = standard error for subgroup i μ = weighted mean (setting average)	$MID = 1000 * \sum_i p_i \ln(l_i)$	$var(MID) = \sum_i p_i \sigma_i^2 (\ln(l_i) + 1 - 1/\ln(l_i))^2$	- Replace any estimates equal to zero with a very small number (0.000001) to avoid issues in calculating the log of zero. - Calculate the population share of each subgroup out of the total population. - Calculate the weighted mean of the subgroup estimates. - Calculate the share of the indicator in each subgroup by dividing the estimate of each subgroup by the weighted mean. - Multiply the negative natural logarithm (log) of the share of the indicator in each subgroup by the population share. - Calculate MID as the sum of the results from step #5, multiplied by 1000 for easier interpretation.	MID is more sensitive to differences further from the setting average (by the use of the logarithm).

Tab	Summary measure name	Ordered / Non-ordered dimension	Absolute / Relative measure	Simple / Complex measures	Weighted / Unweighted measures	Unit of measure	Value of no inequality	Interpretation	Notation	Summary measure formula	Variance formula	Step-by-step summary measure calculation	Notes
par-2	Population attributable risk	Non-ordered	Absolute	Complex	Weighted	Unit of indicator	Zero	For non-ordered dimensions of inequality, PAR shows the potential improvement in the average of an indicator, in absolute terms, that could be achieved if all population subgroups had the same level of the indicator as the subgroup with the best estimate. PAR assumes only positive values for favourable indicators and only negative values for adverse indicators. The larger the absolute value, the higher the level of inequality.	y_{ref} = estimate for a reference subgroup μ = weighted mean (setting average) PAF = population attributable fraction PAF_{se} = standard error of population attributable fraction	$PAR = y_{ref} - \mu$	$var(PAR) = ((\mu(PAF + 1.96 PAF_{se}) - (PAF - 1.96 PAF_{se})) / 2 * 1.96)^2$	1. Identify the reference subgroup, y_{ref} (see notes). 2. Calculate the population share of each subgroup out of the total population. 3. Calculate the weighted mean of the subgroup estimates. 3. Calculate PAR as the difference between the reference subgroup estimate and the weighted mean.	For ordered dimensions, the reference subgroup (y_{ref}) is the subgroup with the highest estimate for favourable indicators, or the subgroup with the lowest estimate for adverse indicators. The variable <i>favourable_indicator</i> identifies whether the indicator is favourable (1) or adverse (0). See the glossary for the definition of favourable and adverse indicators. If the indicator is favourable (<i>favourable_indicator</i> =1) and PAR has a negative value, PAR is equal to zero. If the indicator is adverse (<i>favourable_indicator</i> =0) and PAR has a positive value, PAR is equal to zero. See the glossary for the definition of favourable and adverse indicators. The variable <i>indicator_scale</i> identifies the scale of the indicator. See the glossary for the definition of indicator scale.
pafr-2	Population attributable fraction	Non-ordered	Relative	Complex	Weighted	No unit	Zero	For non-ordered dimensions of inequality, PAF shows the potential improvement in the average of an indicator, in relative terms, that could be achieved if all population subgroups had the same level of the indicator as the subgroup with the best estimate. PAF assumes only positive values for favourable indicators and only negative values for adverse indicators. The larger the absolute value of PAF, the larger the level of inequality.	y_{ref} = estimate for a reference subgroup μ = weighted mean (setting average) a = number of people not in the reference subgroup for whom the indicator was achieved b = number of people not in the reference subgroup for whom the indicator was not achieved c = number of people in the reference subgroup for whom the indicator was achieved d = number of people in the reference subgroup for whom the indicator was not achieved N = total number of people ($a + b + c + d$)	$PAF = 100 * (y_{ref} - \mu) / \mu$	$var(PAF) = (N \mu^2 (N - c) + bc^2) / (a + c)^2 (c + d)^2$	1. Identify the reference subgroup, y_{ref} (see notes). 2. Calculate the population share of each subgroup out of the total population. 3. Calculate the weighted mean of the subgroup estimates. 3. Calculate PAF as the difference between the reference subgroup estimate and the weighted mean, divided by the weighted mean, times 100.	For ordered dimensions, the reference subgroup (y_{ref}) is the subgroup with the highest estimate for favourable indicators, or the subgroup with the lowest estimate for adverse indicators. The variable <i>favourable_indicator</i> identifies whether the indicator is favourable (1) or adverse (0). See the glossary for the definition of favourable and adverse indicators. If the indicator is favourable (<i>favourable_indicator</i> =1) and PAF has a negative value, PAF is equal to zero. If the indicator is adverse (<i>favourable_indicator</i> =0) and PAF has a positive value, PAF is equal to zero. See the glossary for the definition of favourable and adverse indicators. The variable <i>indicator_scale</i> identifies the scale of the indicator. See the glossary for the definition of indicator scale.