

Figure S1: The age-standardized point prevalence of ischemic stroke in 2021 for the 21 Global Burden of Disease regions, by sex. (Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

SDI = Sociodemographic Index, GBD = Global Burden of Disease

Female Male

GBD Regions

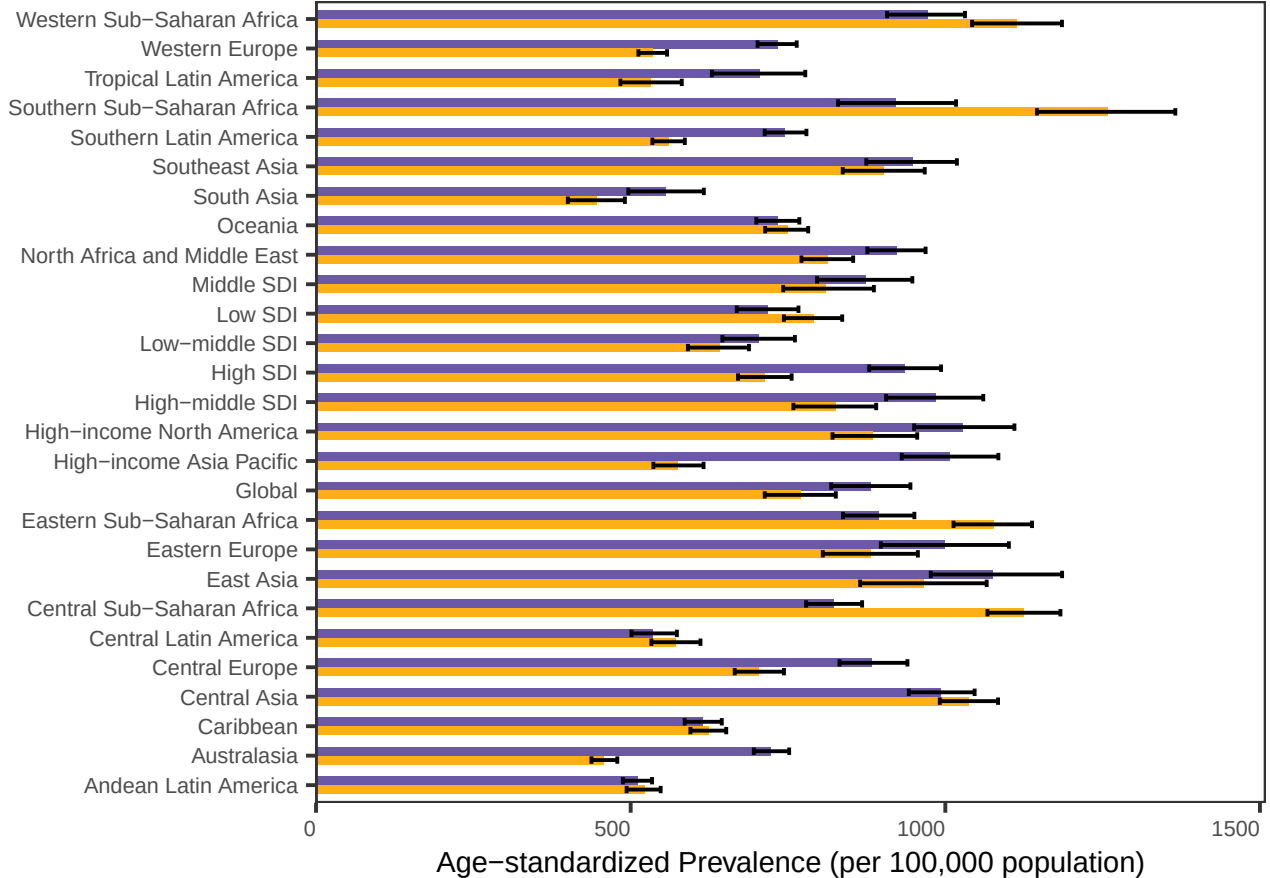


Figure S2: The age-standardized incidence of ischemic stroke in 2021 for the 21 Global Burden of Disease regions, by sex. (Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

SDI = Sociodemographic Index, GBD = Global Burden of Disease

Female Male

GBD Regions

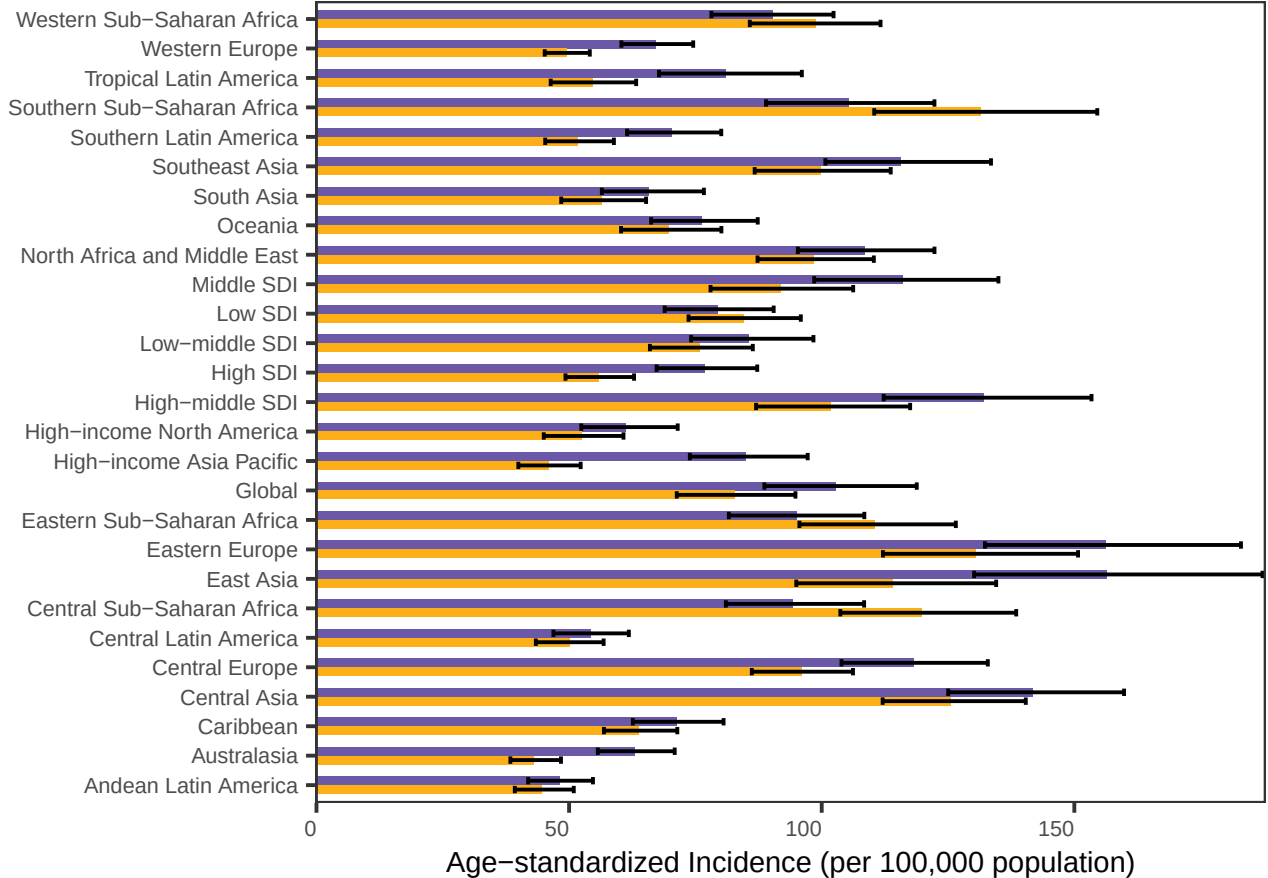


Figure S3: The age-standardized death rates of ischemic stroke in 2021 for the 21 Global Burden of Disease regions, by sex. (Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

SDI = Sociodemographic Index, GBD = Global Burden of Disease

Female Male

GBD Regions

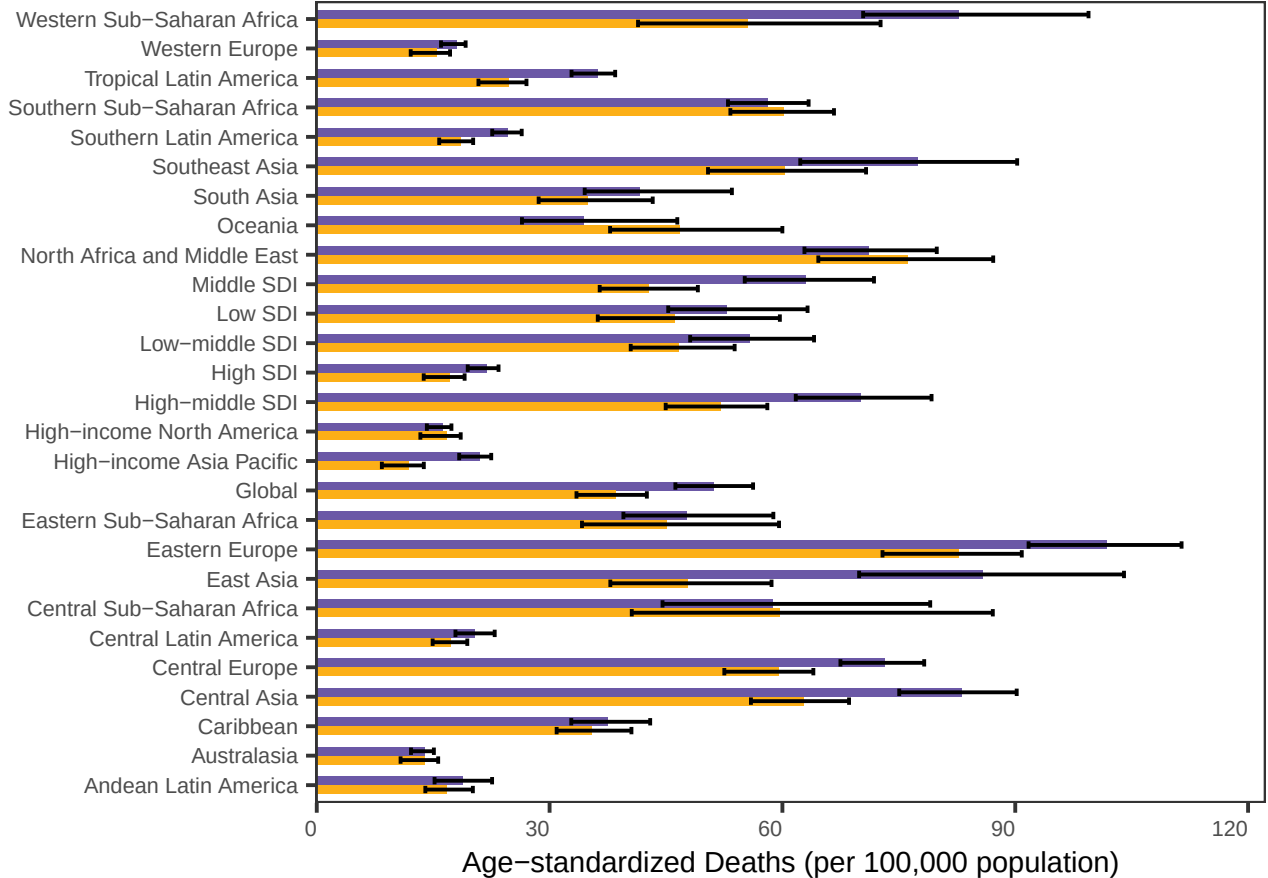


Figure S4: The age-standardized DALYs of ischemic stroke in 2021 for the 21 Global Burden of Disease regions, by sex.

DALYs = disability-adjusted life years (Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

SDI = Sociodemographic Index, GBD = Global Burden of Disease

Female Male

GBD Regions

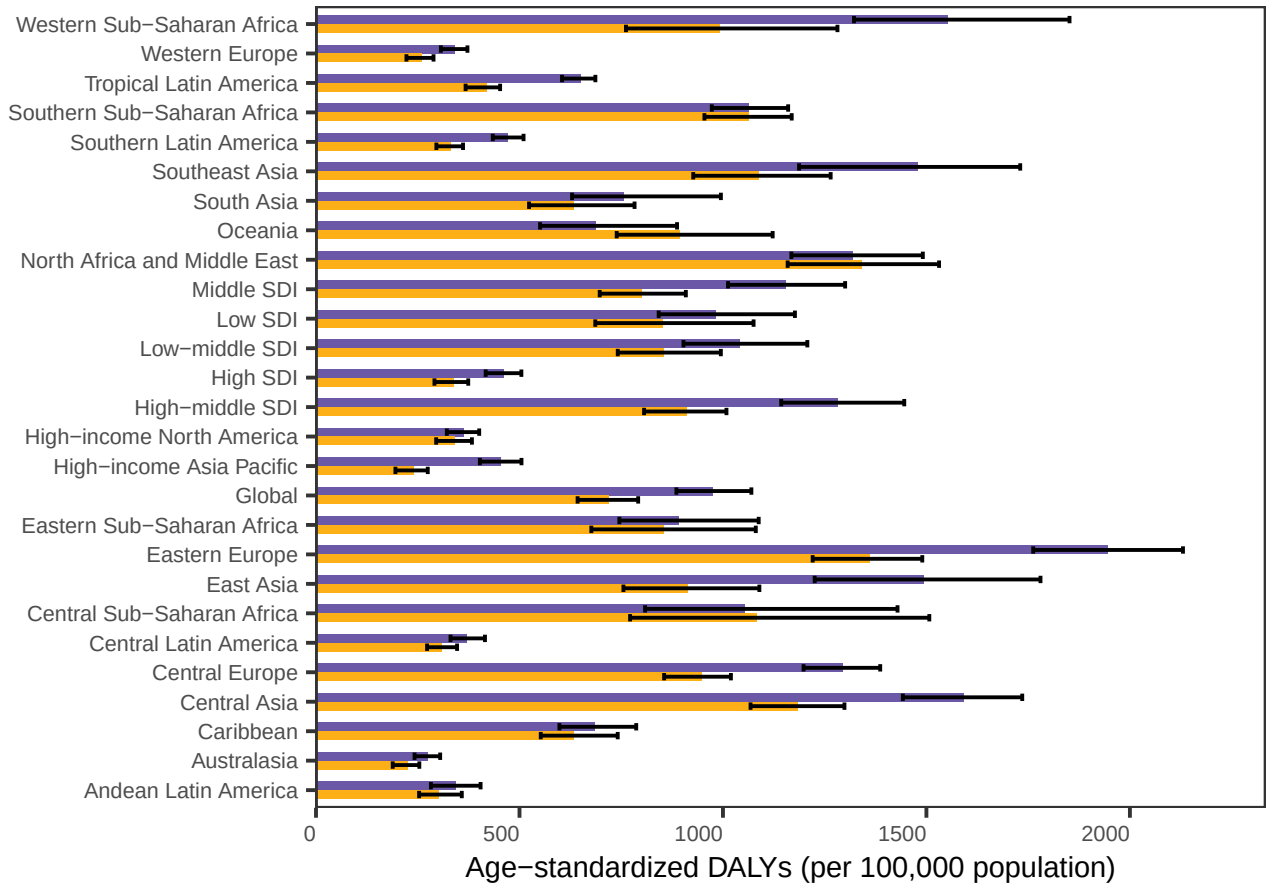


Figure S5: The percentage change in the age-standardized point prevalence of ischemic stroke from 1990 to 2021 for the 21 Global Burden of Disease regions, by sex.

(Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

SDI = Sociodemographic Index, GBD = Global Burden of Disease

Female Male

GBD Regions

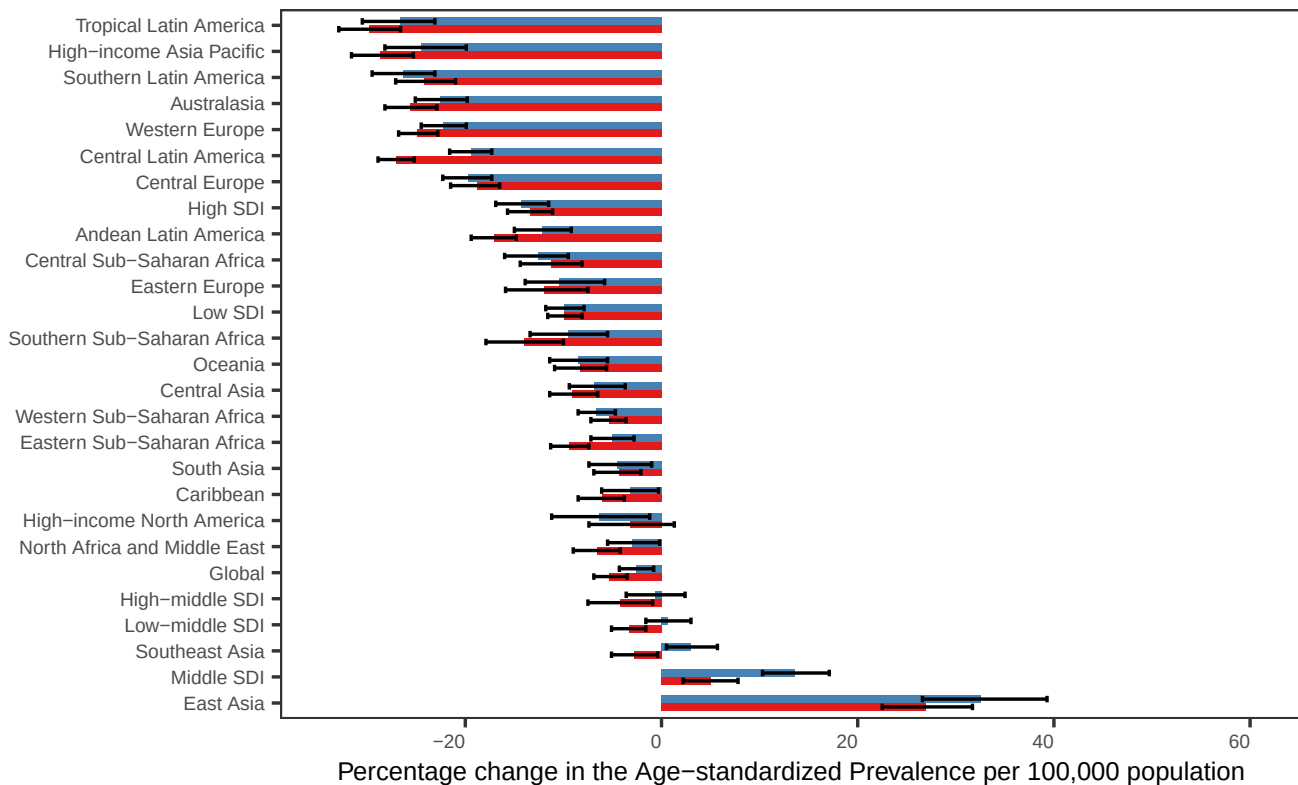


Figure S6: The percentage change in the age-standardized incidence of ischemic stroke from 1990 to 2021 for the 21 Global Burden of Disease regions, by sex. (Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

SDI = Sociodemographic Index, GBD = Global Burden of Disease

Female Male

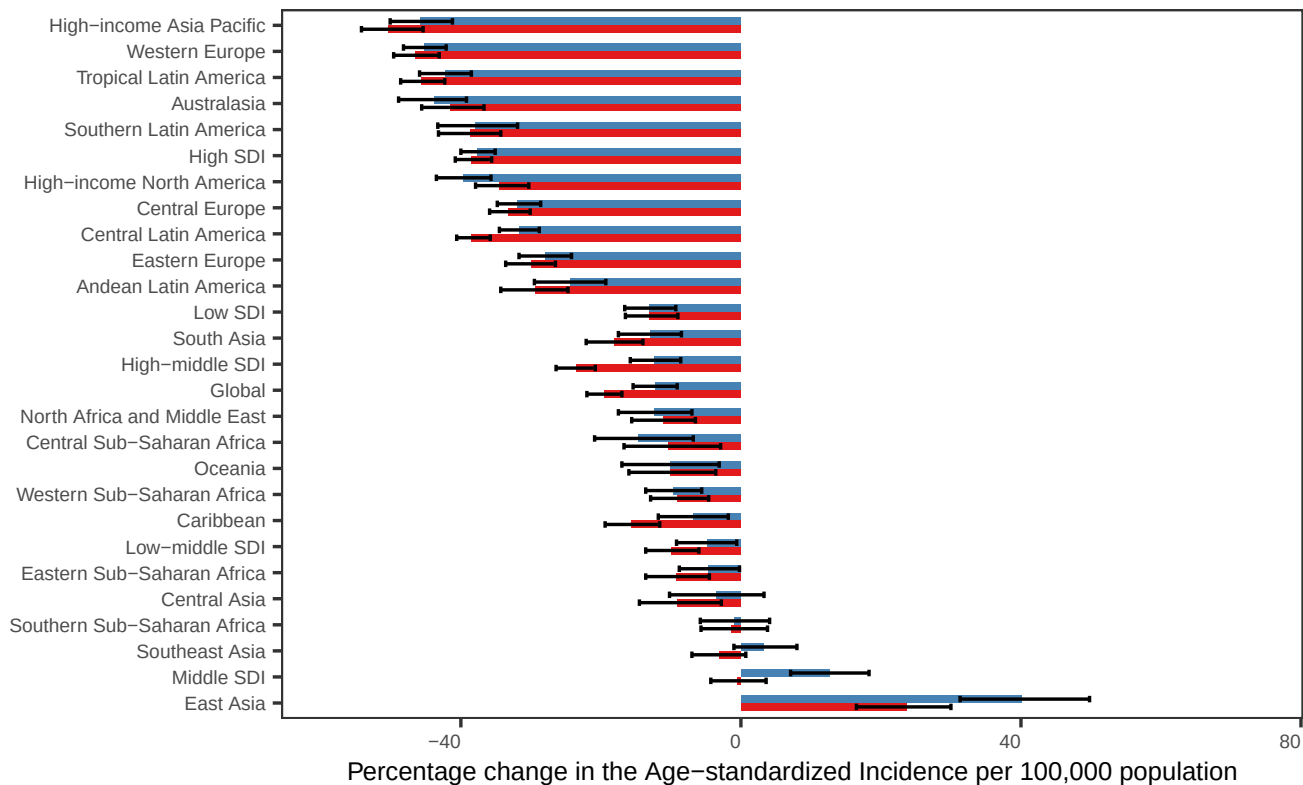


Figure S7: The percentage change in the age-standardized death rate of ischemic stroke from 1990 to 2021 for the 21 Global Burden of Disease regions, by sex. (Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

SDI = Sociodemographic Index, GBD = Global Burden of Disease

Female Male

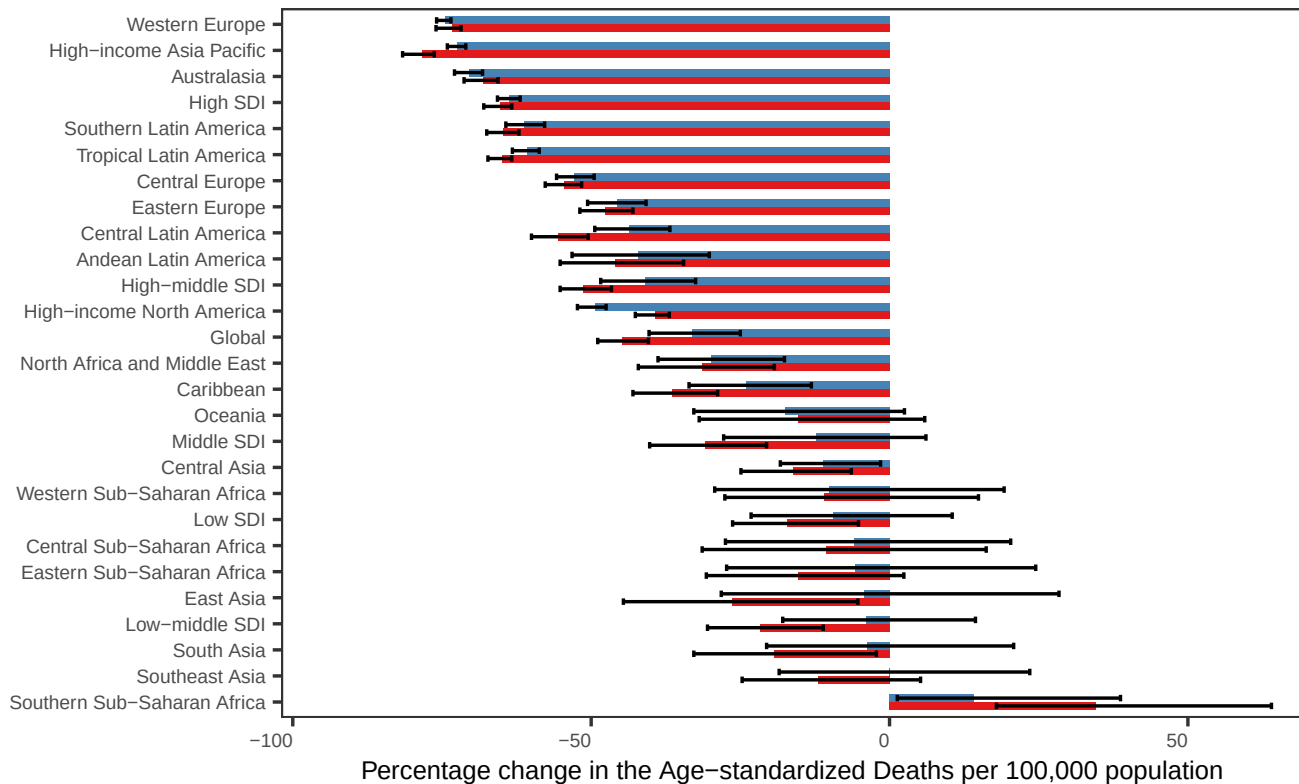


Figure S8: The percentage change in the age-standardized DALYs of ischemic stroke from 1990 to 2021 for the 21 Global Burden of Disease regions, by sex.

DALYs = disability-adjusted life years, SDI = Sociodemographic Index, GBD = Global Burden of Disease (Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

Female Male

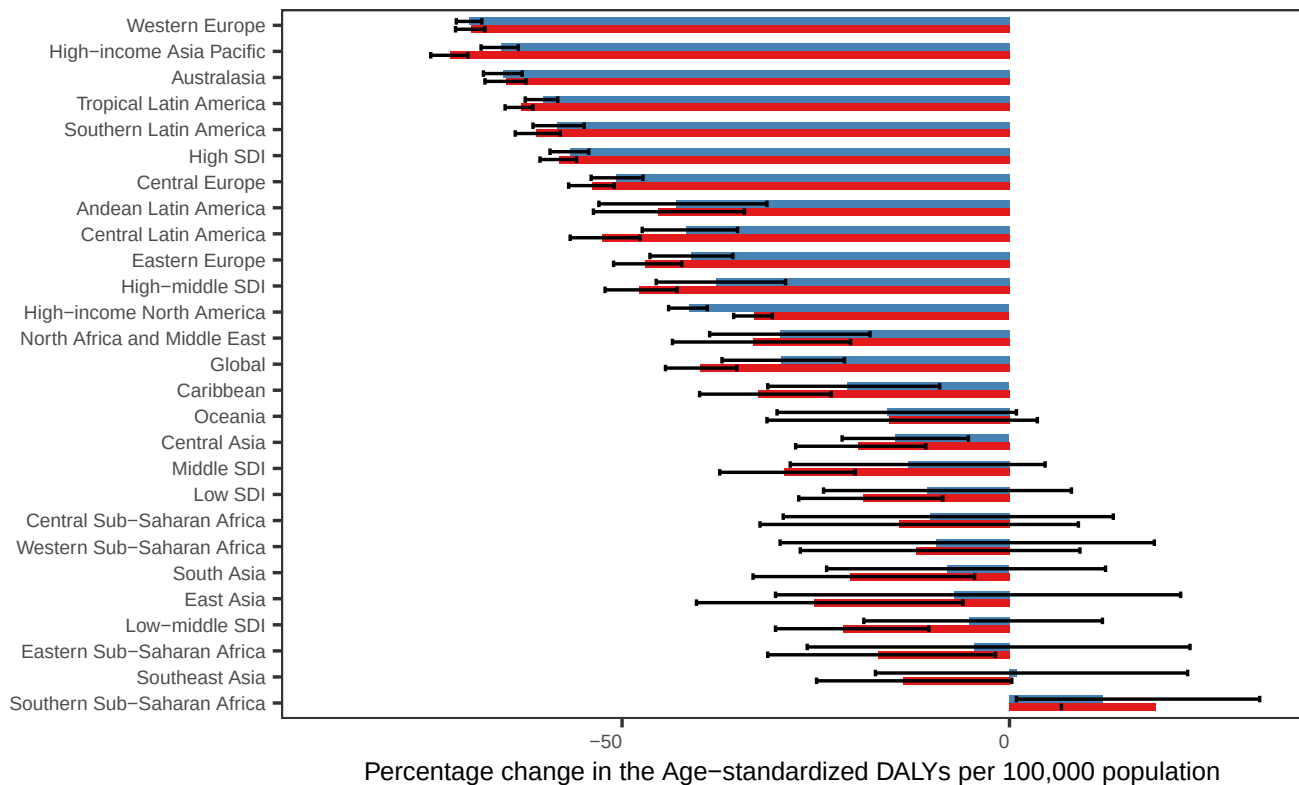
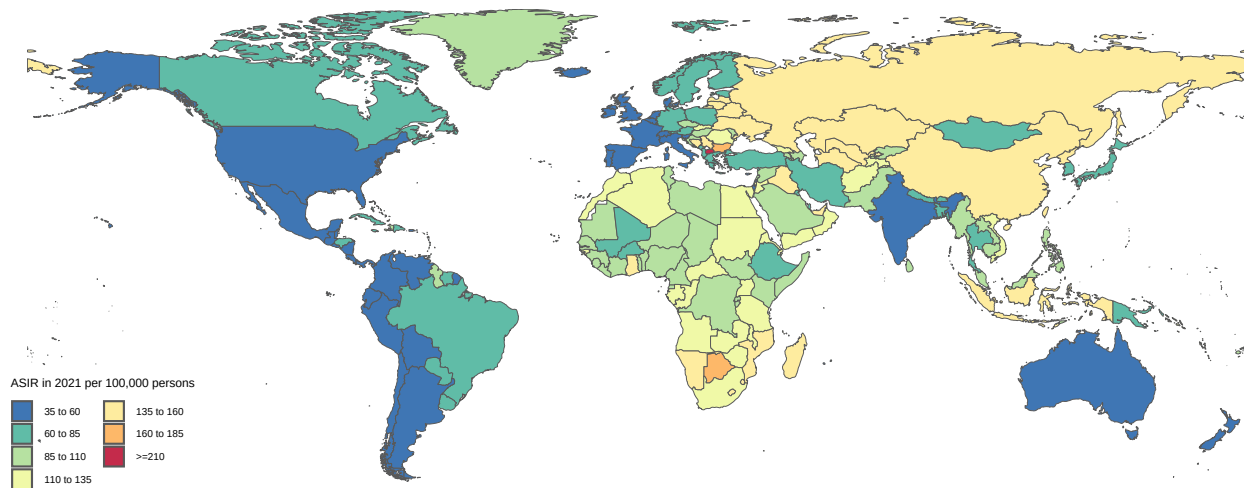


Figure S9: Age-standardized incidence of ischemic stroke per 100,000 population in 2021, by country. (Generated from data available at <https://ghdx.healthdata.org/gbd-results-tool>).

ASIR = Age-Standardized Incidence Rate



Caribbean and central America



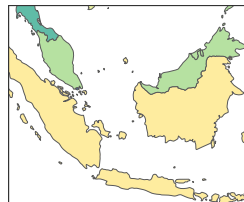
Persian Gulf



Balkan Peninsula



Southeast Asia



West Africa



Eastern Mediterranean

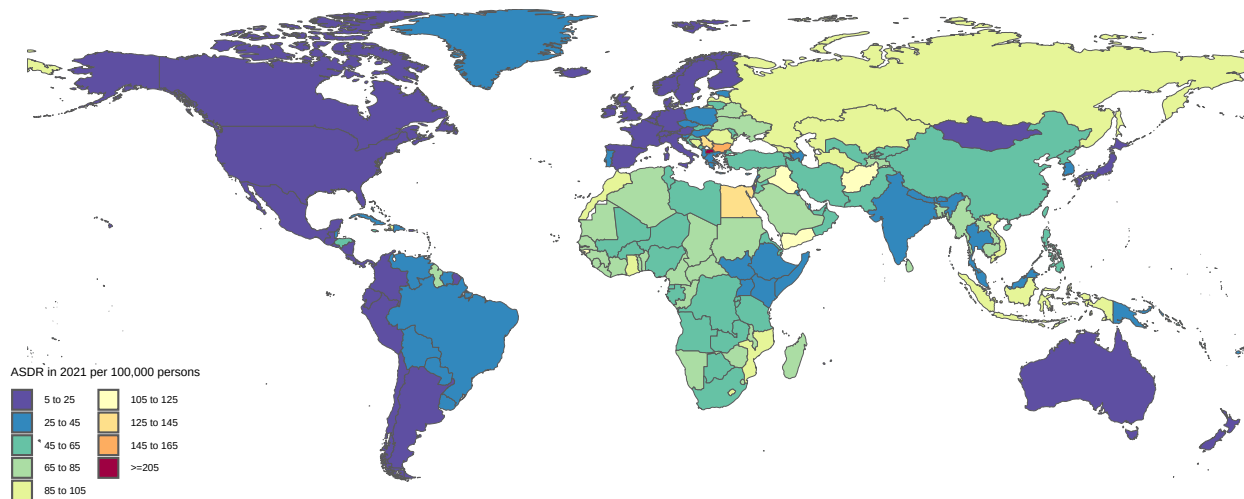


Northern Europe

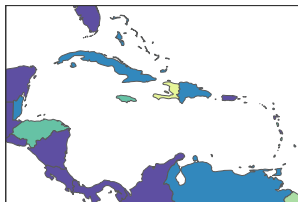


Figure S10: Age-standardized death rate of ischemic stroke per 100,000 population in 2021, by country. (Generated from data available at <https://ghdx.healthdata.org/gbd-results-tool>).

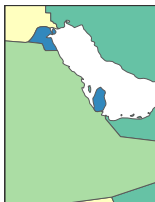
ASDR = Age-Standardized Death Rate



Caribbean and central America



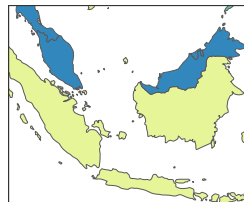
Persian Gulf



Balkan Peninsula



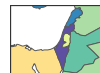
Southeast Asia



West Africa



Eastern Mediterranean

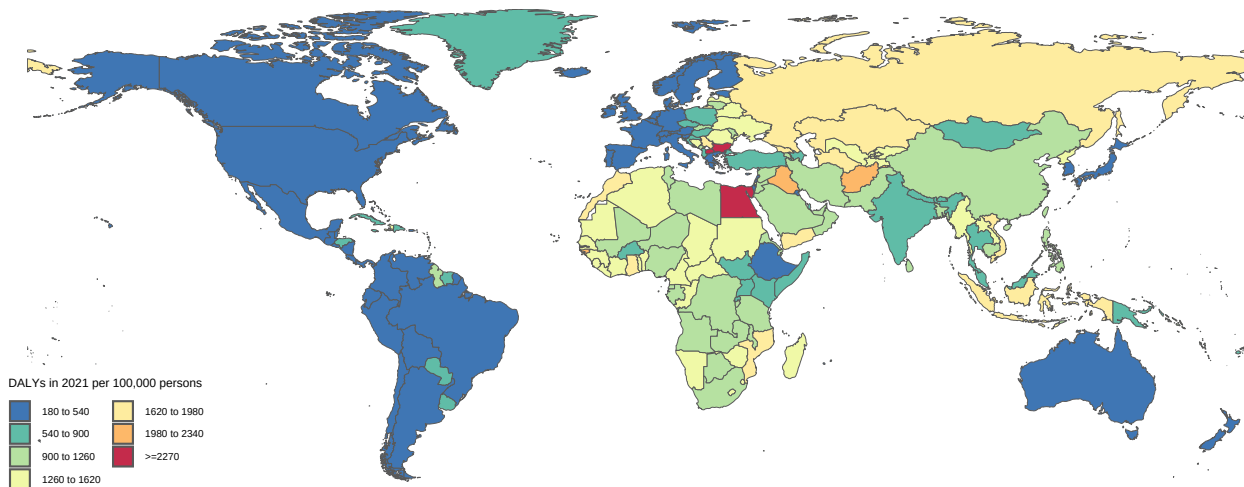


Northern Europe

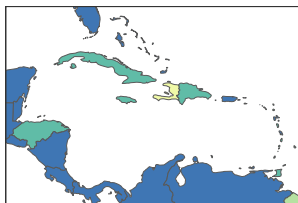


Figure S11: Age-standardized DALYs of ischemic stroke per 100,000 population in 2021, by country.

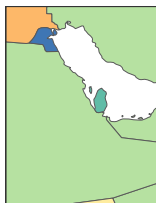
DALYs = disability-adjusted life years (Generated from data available from <https://ghdx.healthdata.org/gbd-results-tool>).



Caribbean and central America



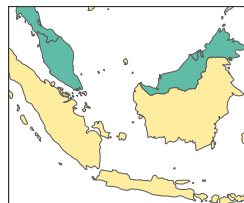
Persian Gulf



Balkan Peninsula



Southeast Asia



West Africa



Eastern Mediterranean



Northern Europe



Figure S12: Global prevalence rate of ischemic stroke, by age and sex, in 2021; Red and blue shadows indicate 95% upper and lower uncertainty intervals, respectively.
(Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

2021 Rate of Prevalence

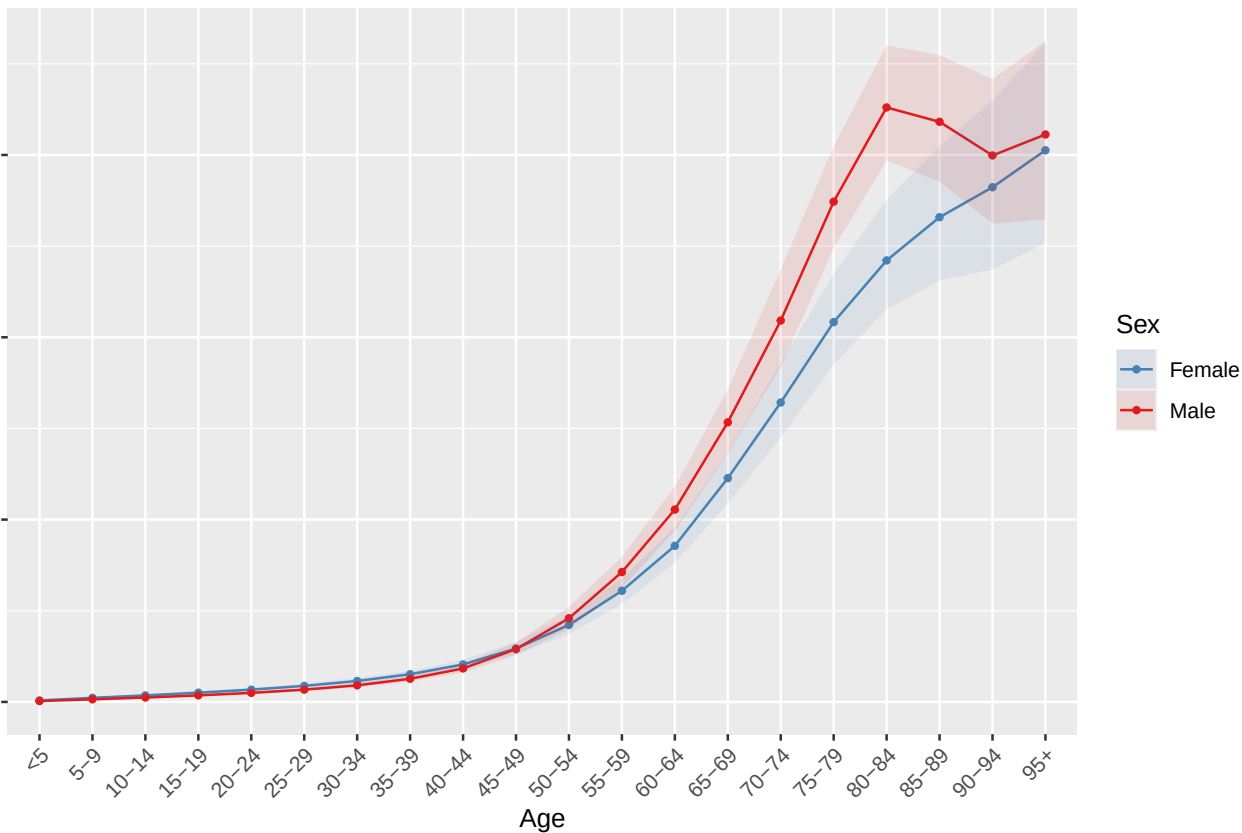


Figure S13: Global number of prevalence of ischemic stroke per 100,000 population, by age and sex, in 2021. (Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

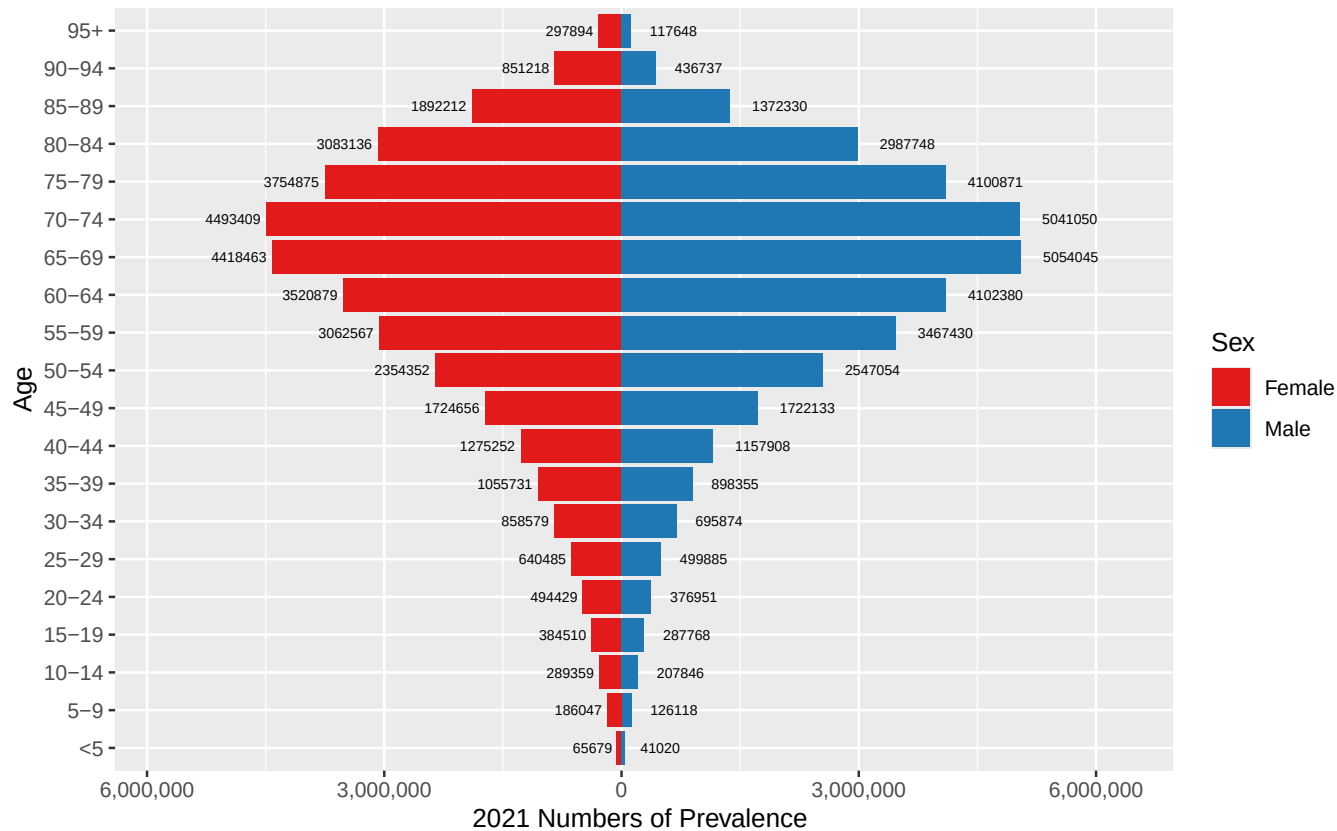


Figure S14: Number of incident cases globally and incidence of ischemic stroke, per 100,000 population by age and sex in 2021. Dotted and dashed lines indicate 95% upper and lower uncertainty intervals, respectively. (Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

Male Female Male Female

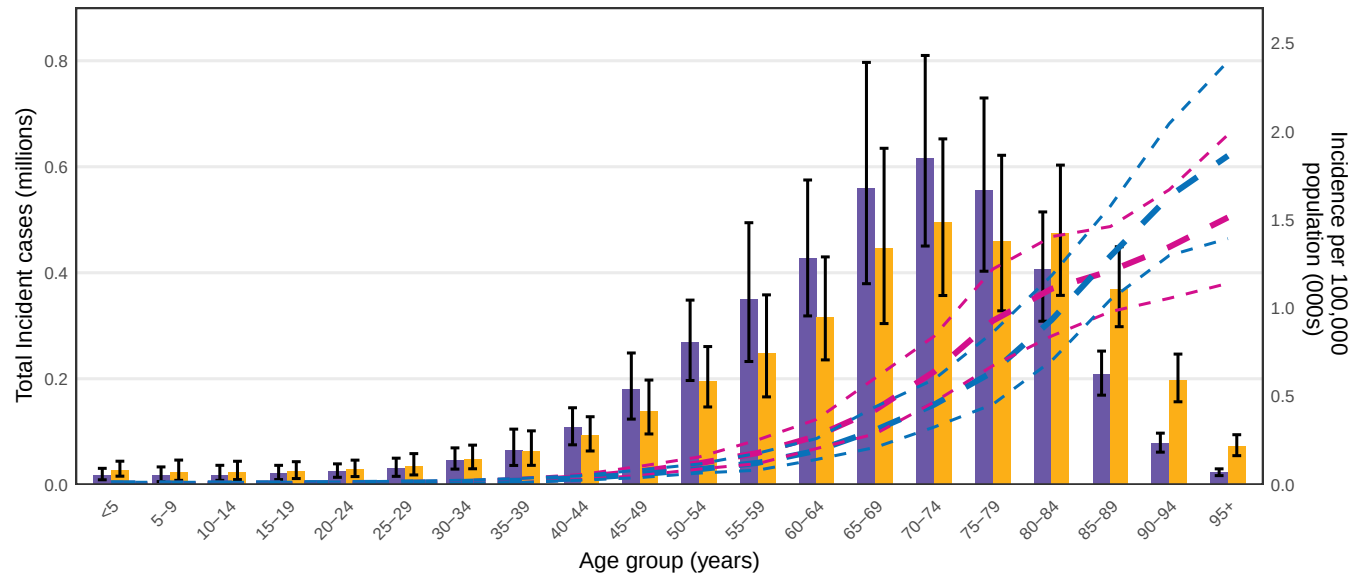
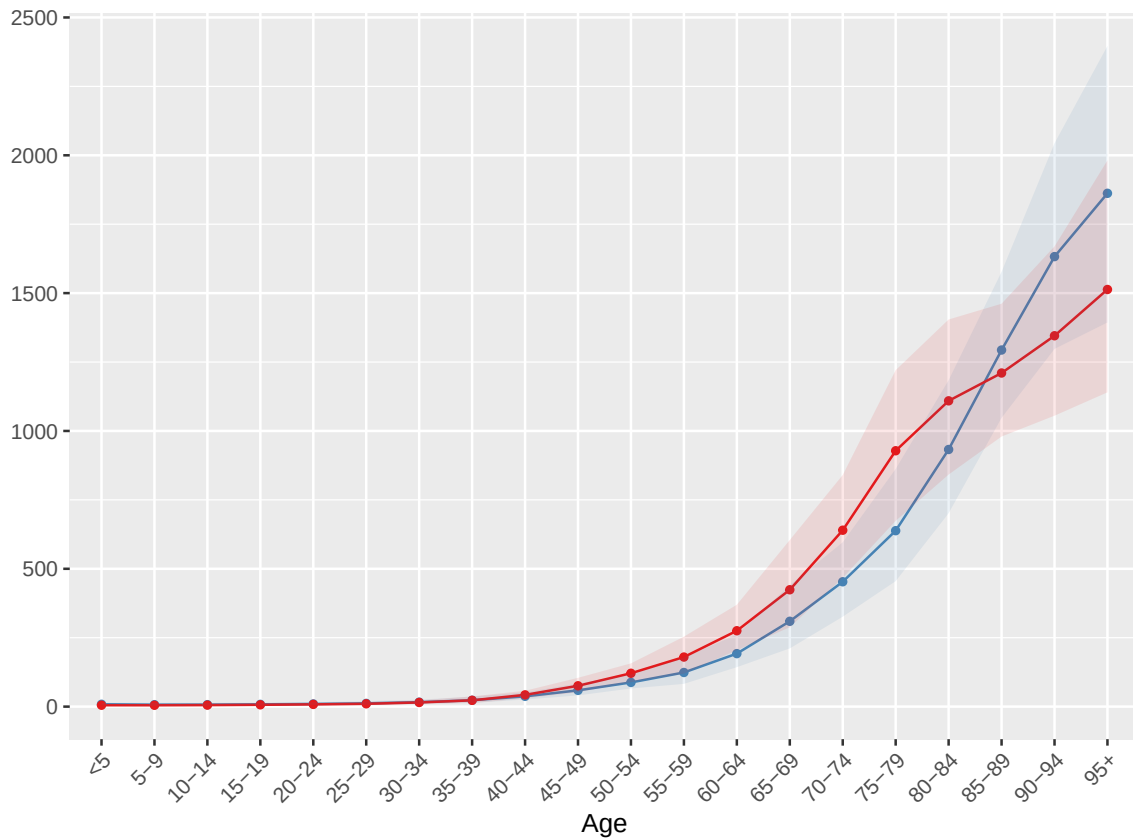


Figure S15: Global incidence rate of ischemic stroke, per 100,000 population, by age and sex, in 2021; Red and blue shadows indicate 95% upper and lower uncertainty intervals, respectively. (Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

2021 Rate of Incidence



Sex

- Female
- Male

Figure S16: Global number of incidence of ischemic stroke, per 100,000 population, by age and sex, in 2021. (Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

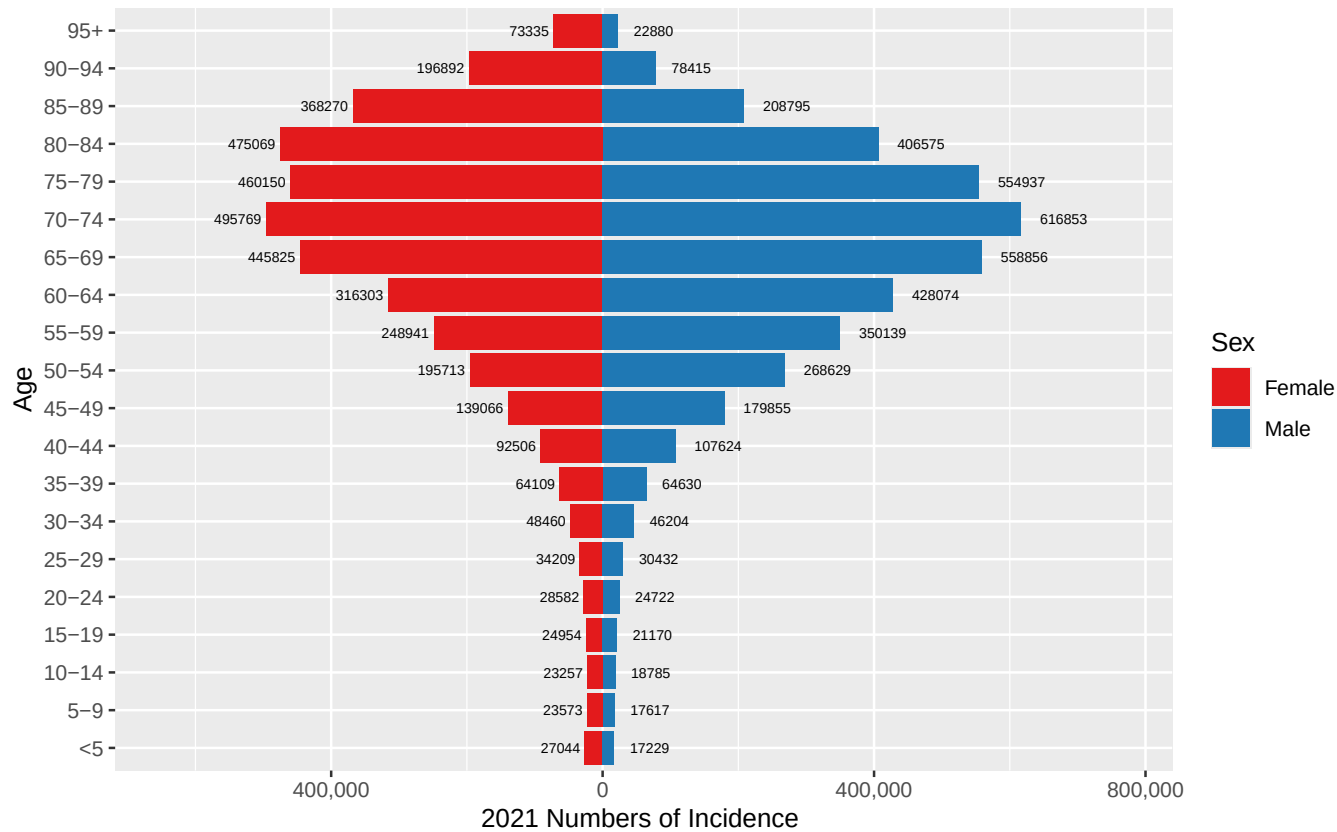


Figure S17: Number of death cases globally and death rate of ischemic stroke, per 100,000 population by age and sex in 2021. Dotted and dashed lines indicate 95% upper and lower uncertainty intervals, respectively. (Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

Male Female Male Female

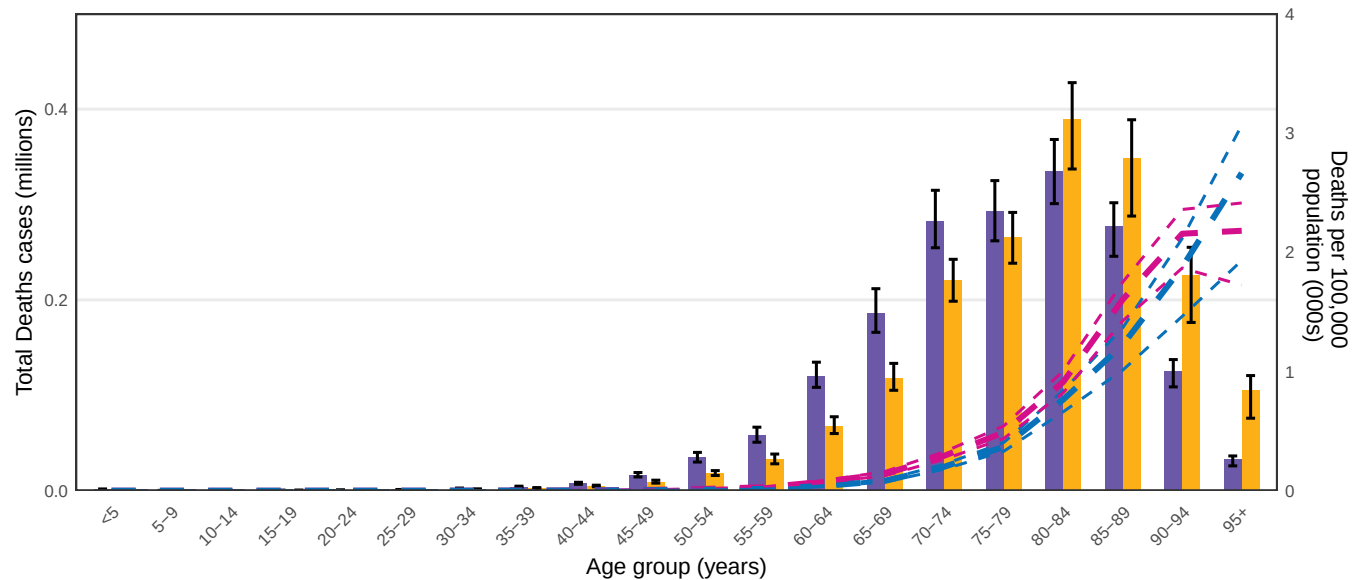


Figure S18: Global death rate of ischemic stroke, per 100,000 population, by age and sex, in 2021; Red and blue shadows indicate 95% upper and lower uncertainty intervals, respectively. (Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

2021 Rate of Deaths

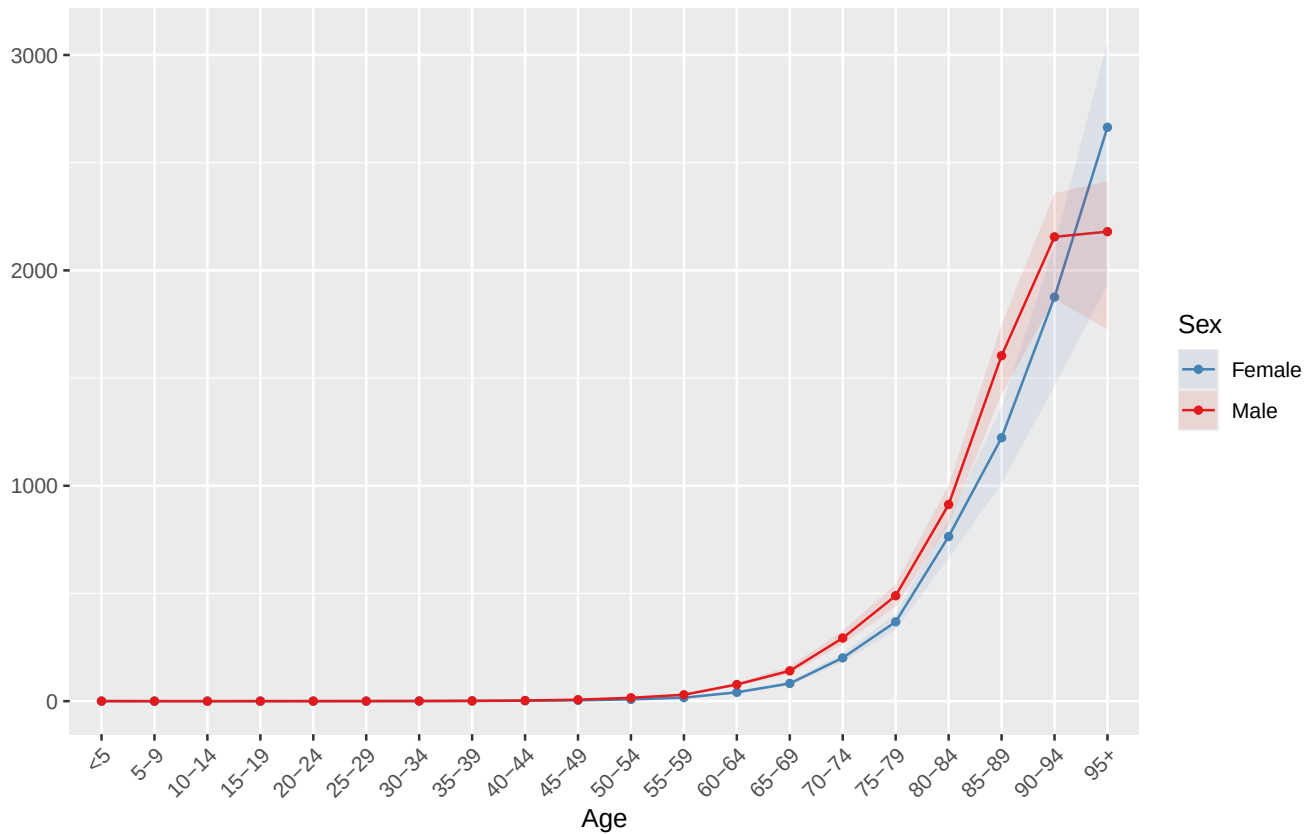


Figure S19: Global number of deaths of ischemic stroke, per 100,000 population, by age and sex, in 2021. (Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

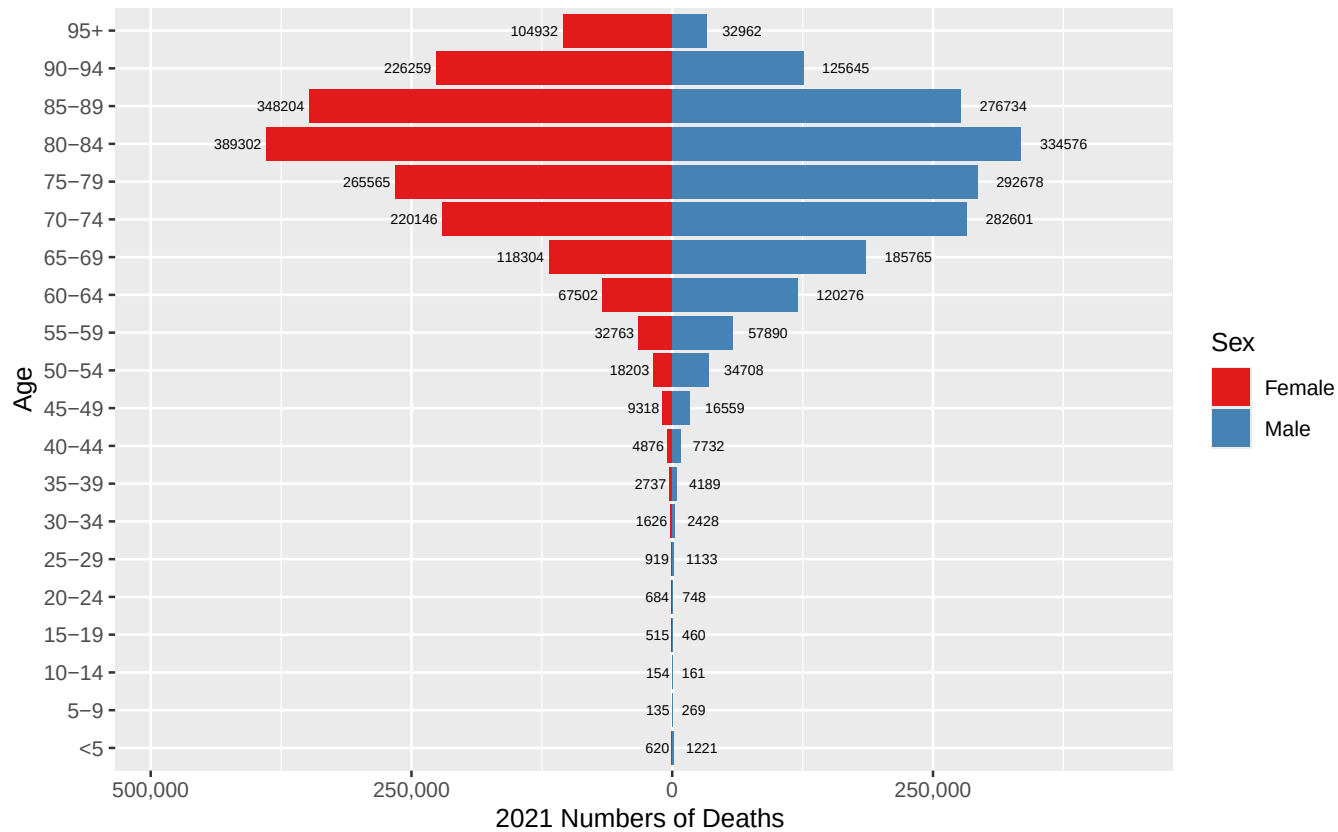


Figure S20: Number of DALYs globally and DALYs rate of ischemic stroke, per 100,000 population by age and sex in 2021. Dotted and dashed lines indicate 95% upper and lower uncertainty intervals, respectively.

DALYs = disability-adjusted life years (Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

Male Female Male Female

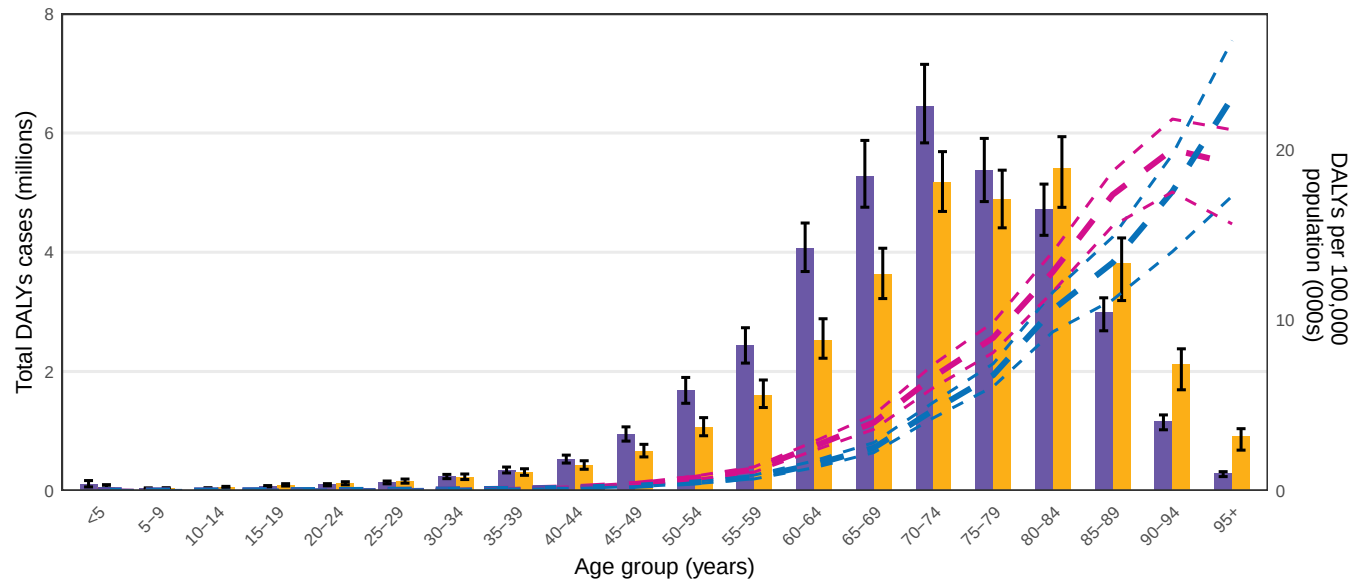
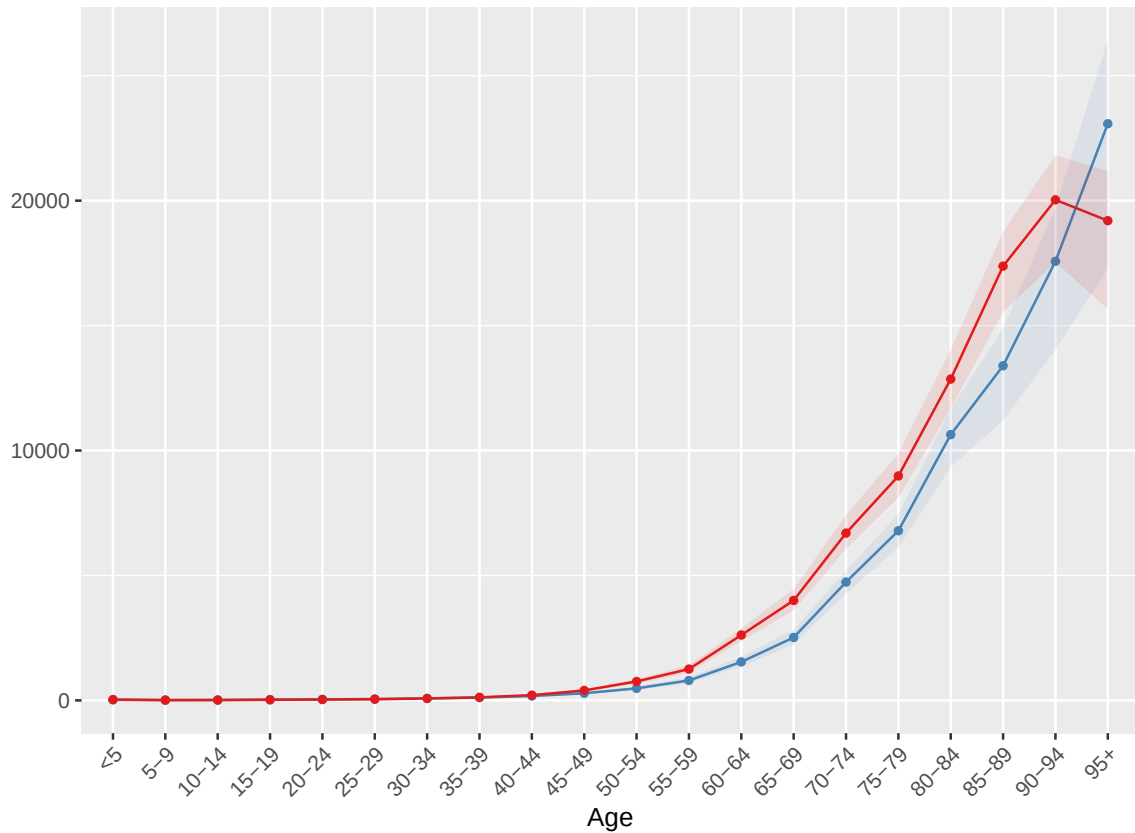


Figure S21: Global DALYs rate of ischemic stroke, per 100,000 population, by age and sex, in 2021; Red and blue shadows indicate 95% upper and lower uncertainty intervals, respectively.

DALYs = disability-adjusted life years (Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

2021 Rate of DALYs



Sex

Female

Male

Figure S22: Global number of DALYs of ischemic stroke, per 100,000 population, by age and sex, in 2021.

DALYs = disability-adjusted life years (Generated from data available from <http://ghdx.healthdata.org/gbd-results-tool>).

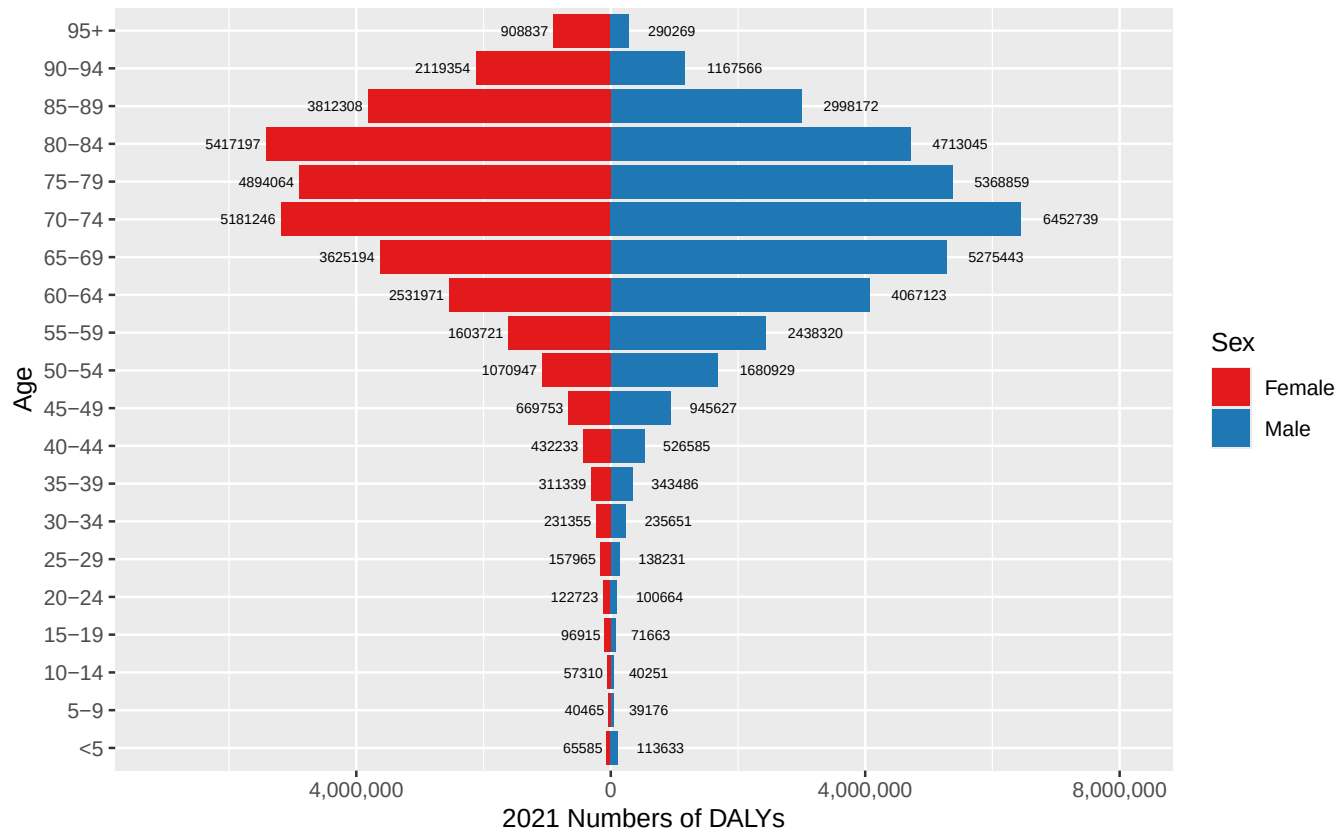


Figure S23: Disability-adjusted life years (DALYs) rates of ischemic stroke for the 22 Global Burden of Disease regions by sociodemographic index, 1990–2021. Thirty-two points are plotted for each region and show the observed age-standardized DALY rates from 1990 to 2021 for that region.

DALYs = disability-adjusted life years. Expected values, based on sociodemographic index and disease rates in all locations, are shown as a solid line. Regions above the solid line represent a higher-than-expected burden and regions below the line show a lower-than-expected burden (Generated from data available at <https://ghdx.healthdata.org/gbd-results-tool>).

- Andean Latin America
- ▽ Central Latin America
- ⌘ Global
- South Asia
- Western Europe
- △ Australasia
- ⊠ Central Sub-Saharan Africa
- ⊞ High-income Asia Pacific
- ▲ Southeast Asia
- Western Sub-Saharan Africa
- + Caribbean
- * East Asia
- ⊠ High-income North America
- ◆ Southern Latin America
- × Central Asia
- ◇ Eastern Europe
- ⊠ North Africa and Middle East
- Southern Sub-Saharan Africa
- ◇ Central Europe
- ⊕ Eastern Sub-Saharan Africa
- Oceania
- Tropical Latin America

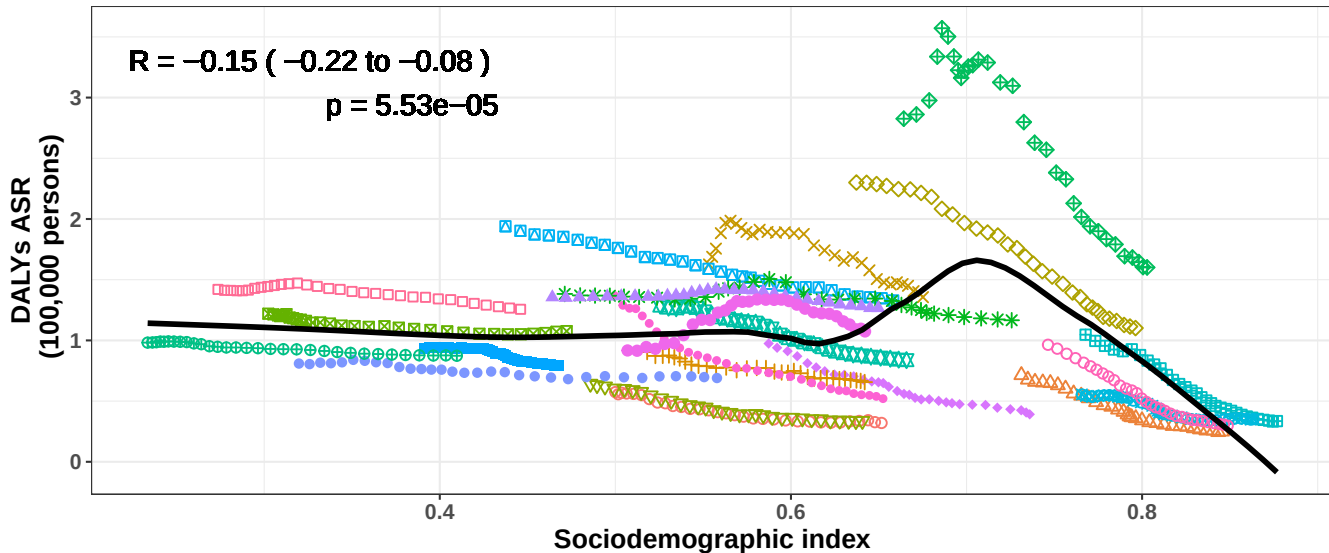


Figure S24: Percentage of disability-adjusted life years (DALYs) due to ischemic stroke attributable to each risk factor for the 21 Global Burden of Disease regions in 2021. (Top 1–6)

DALYs = disability-adjusted life years. (Generated from data available at <https://ghdx.healthdata.org/gbd-results-tool>).

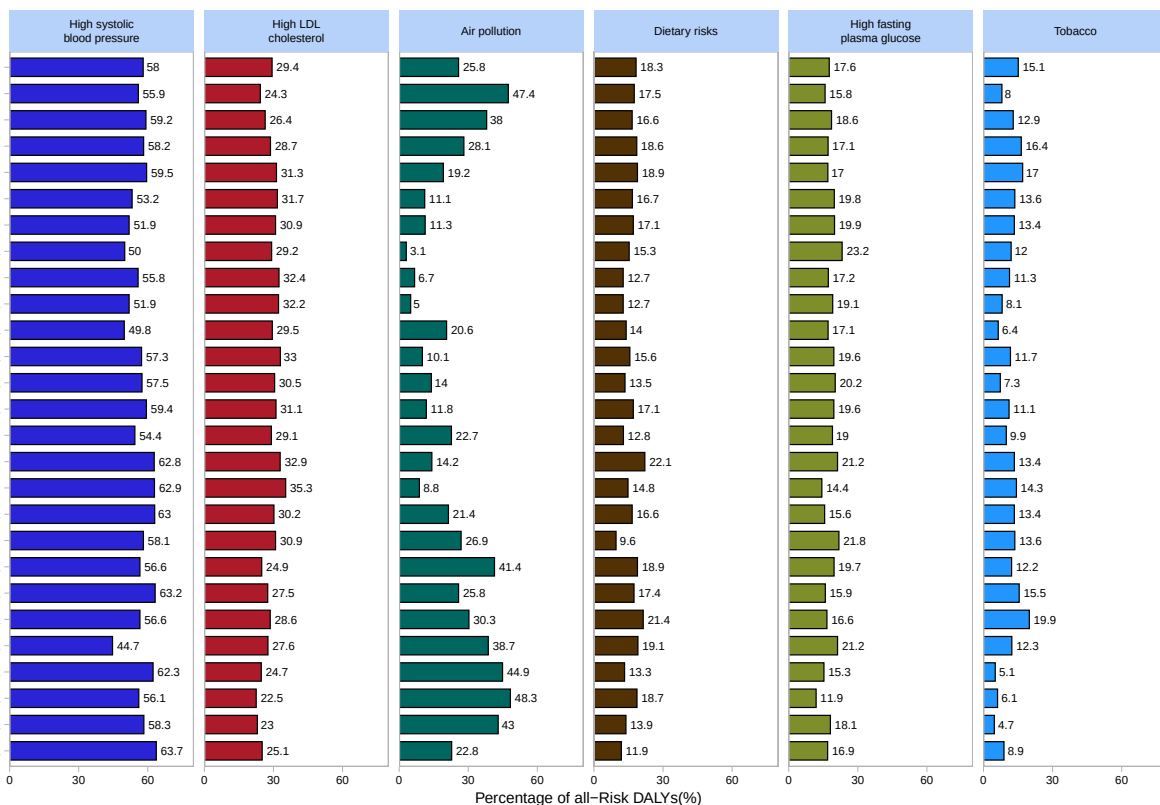
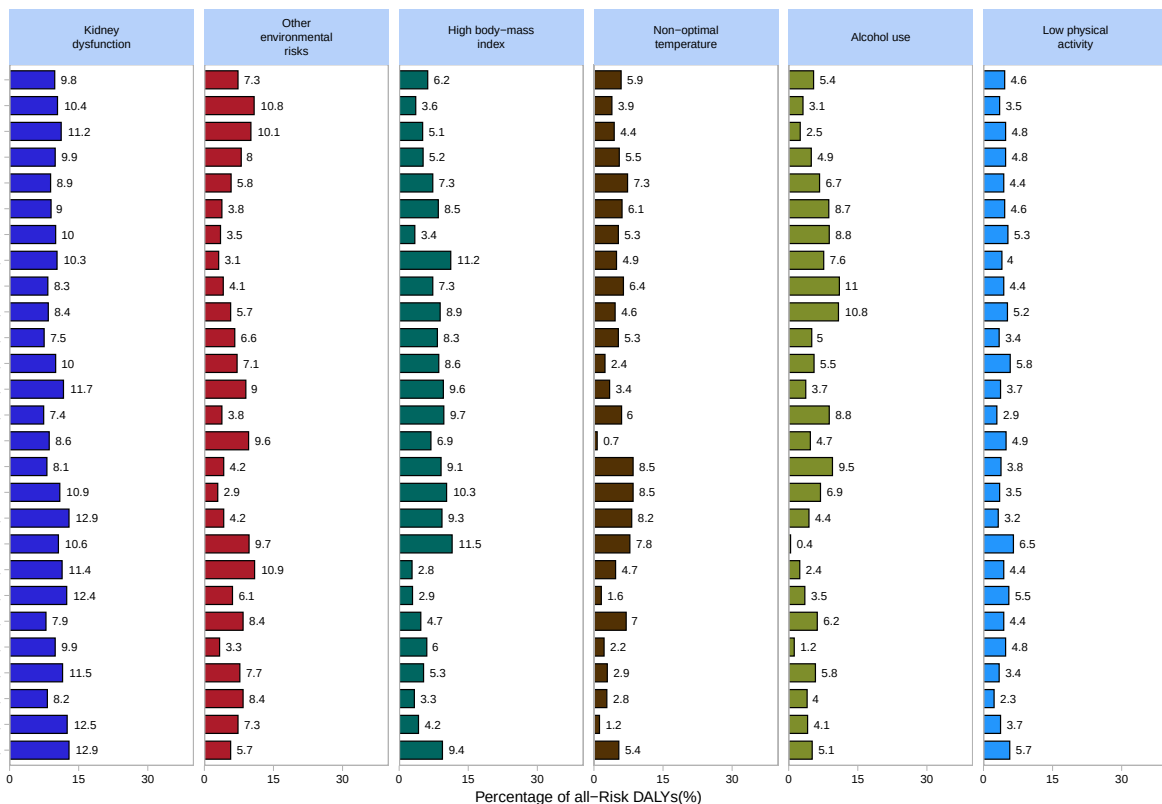


Figure S25: Percentage of disability-adjusted life years (DALYs) due to ischemic stroke attributable to each risk factor for the 21 Global Burden of Disease regions in 2021. (Top 7–12)

DALYs = disability-adjusted life years, GBD = Global Burden of Disease. (Generated from data available at <https://ghdx.healthdata.org/gbd-results-tool>).



Percentage of all-Risk DALYs(%)

Figure S26: Percentage of disability-adjusted life years (DALYs) due to ischemic stroke attributable to each risk factor by age groups in 2021. (Top 1–6)

DALYs = disability-adjusted life years, LDL = Low-Density Lipoprotein. (Generated from data available at <https://ghdx.healthdata.org/gbd-results-tool>).

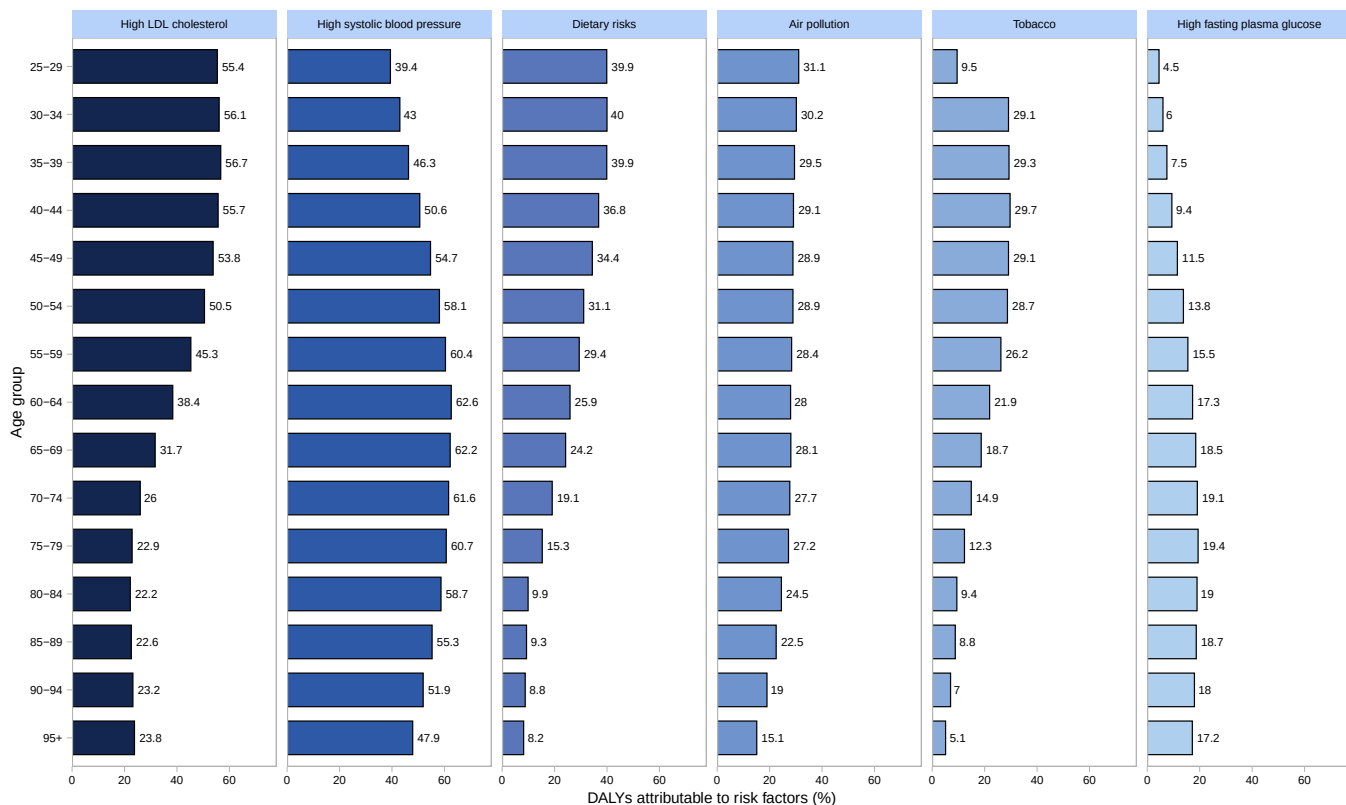


Figure S27: Percentage of disability-adjusted life years (DALYs) due to ischemic stroke attributable to each risk factor by age groups in 2021. (Top 7–12)

DALYs = disability-adjusted life years. (Generated from data available at <https://ghdx.healthdata.org/gbd-results-tool>).

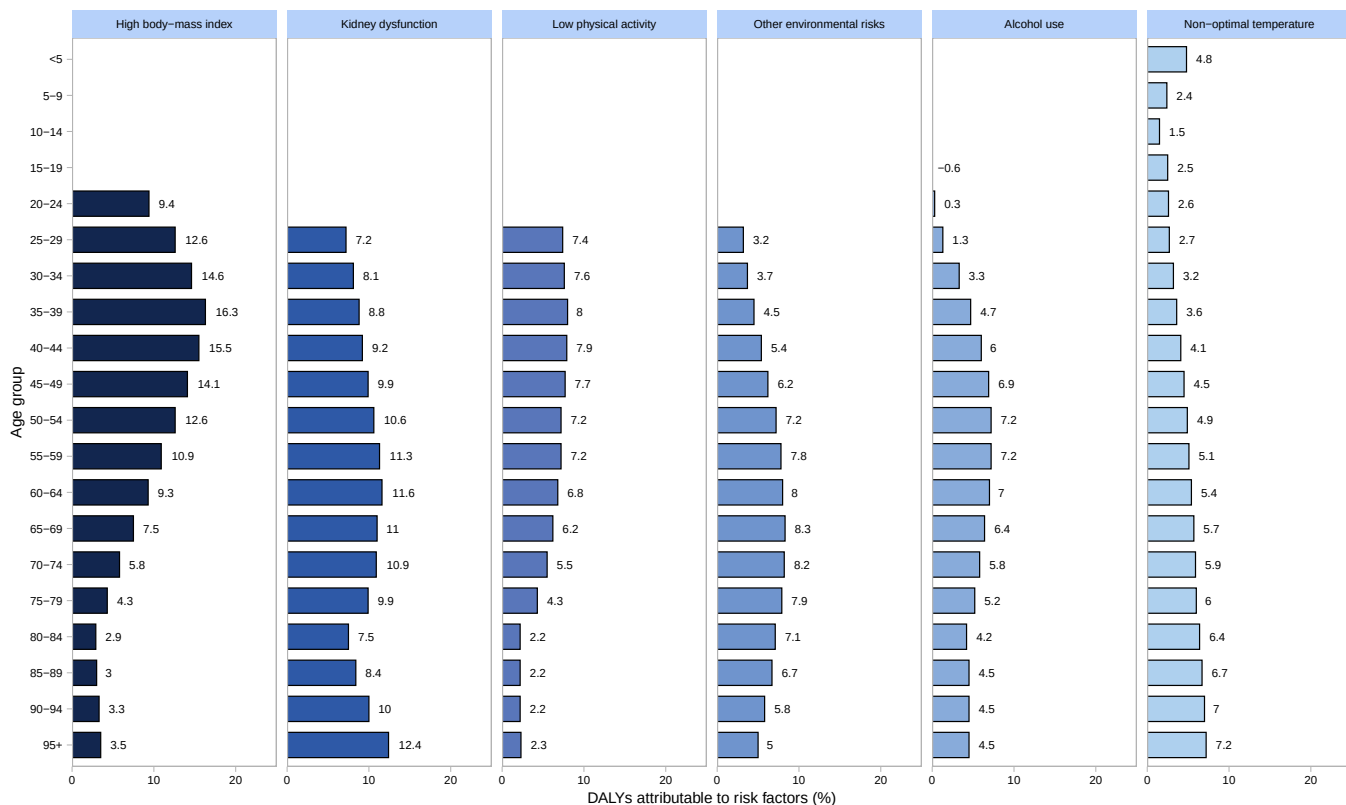


Figure S28: A frontier analysis based on SDI and ischemic stroke DALYs rate from 1990 to 2021.

DALYs = disability-adjusted life years. SDI = sociodemographic index.

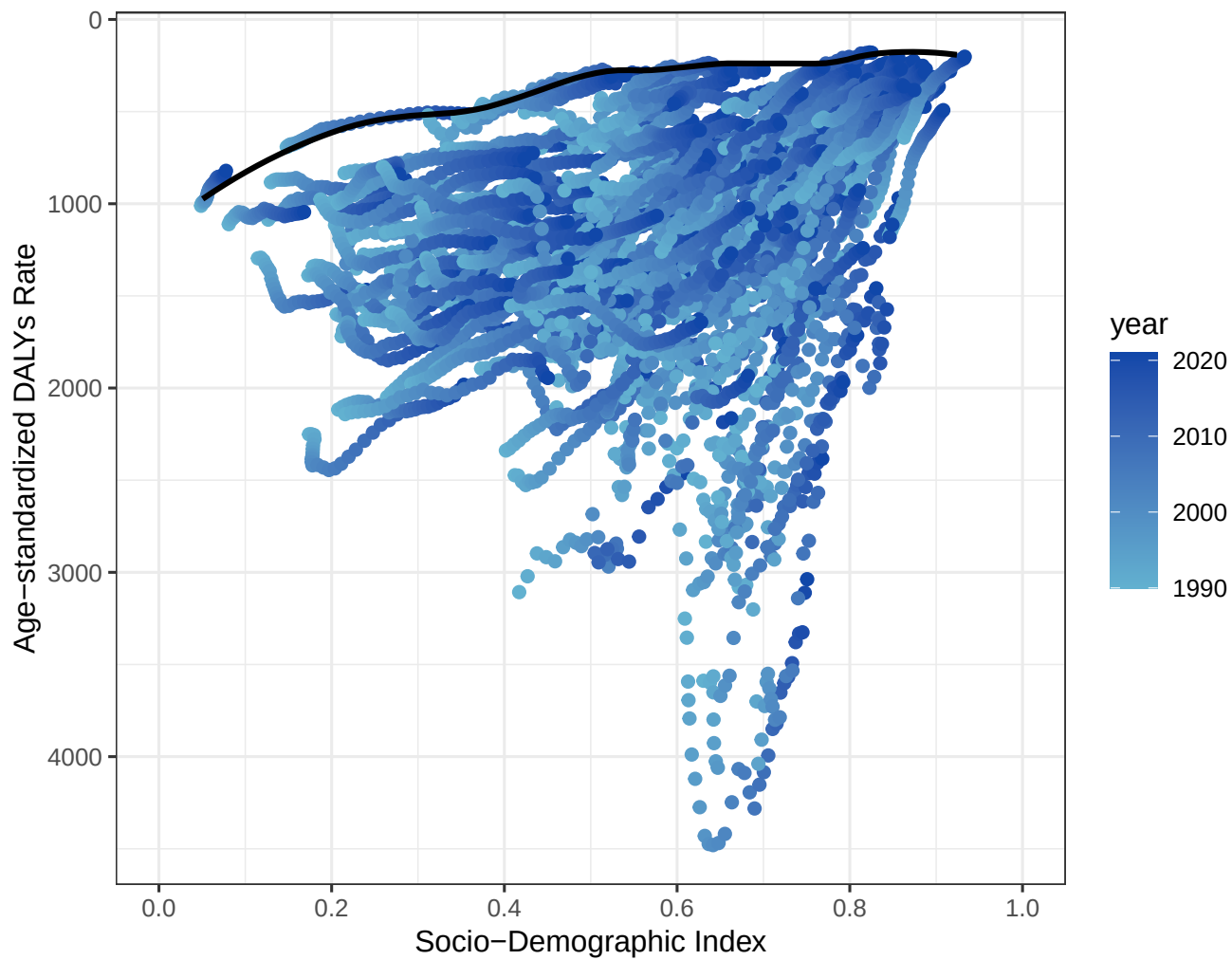


Figure S29: Projected number of ASPR cases for ischemic stroke in males in 2036 according to the SDI.

SDI = sociodemographic index. ASPR = Age-Standardized Prevalence Rate.

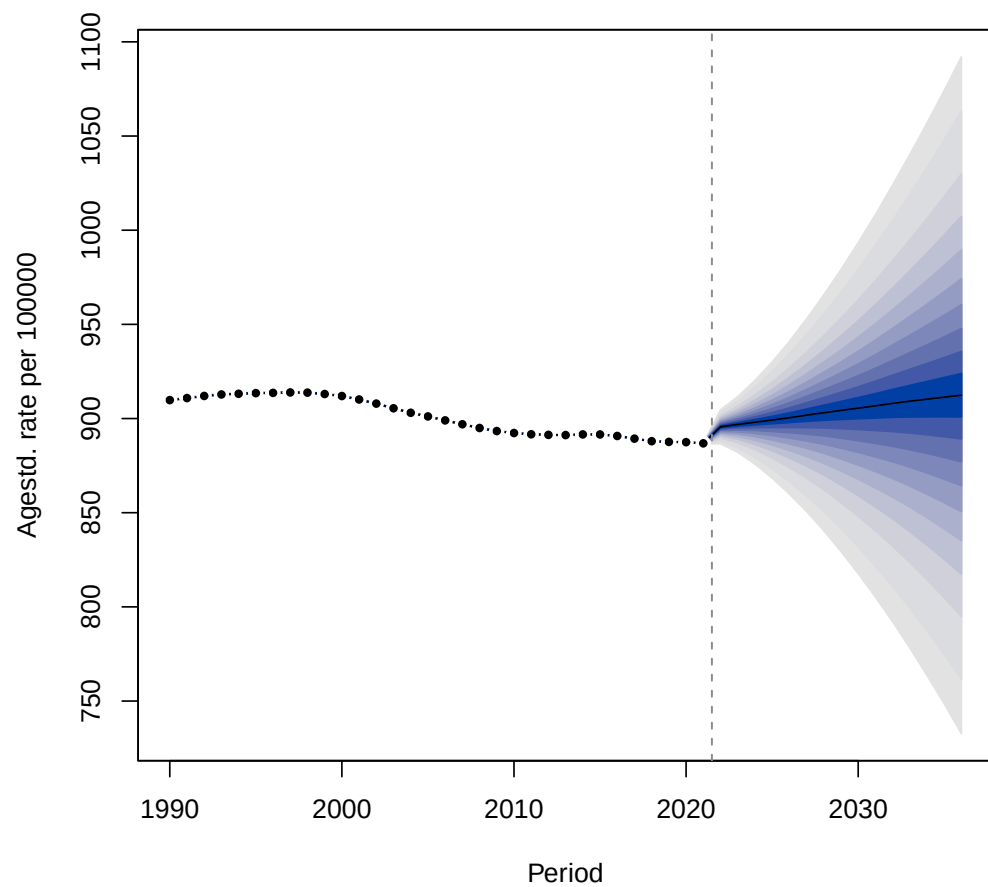


Figure S30: Projected number of ASIR cases for ischemic stroke in males in 2036 according to the SDI.

SDI = sociodemographic index. ASIR = Age-Standardized Incidence Rate.

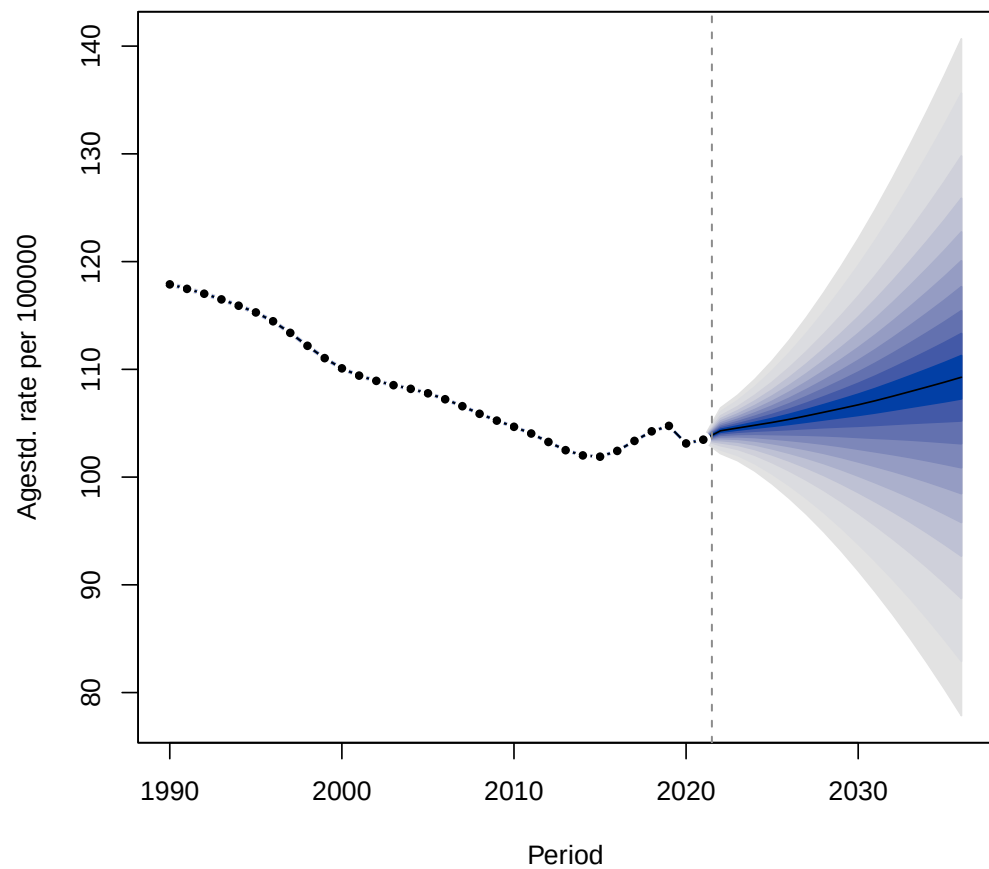


Figure S31: Projected number of ASDR cases for ischemic stroke in males in 2036 according to the SDI.

SDI = sociodemographic index. ASDR = Age-Standardized Death Rate.

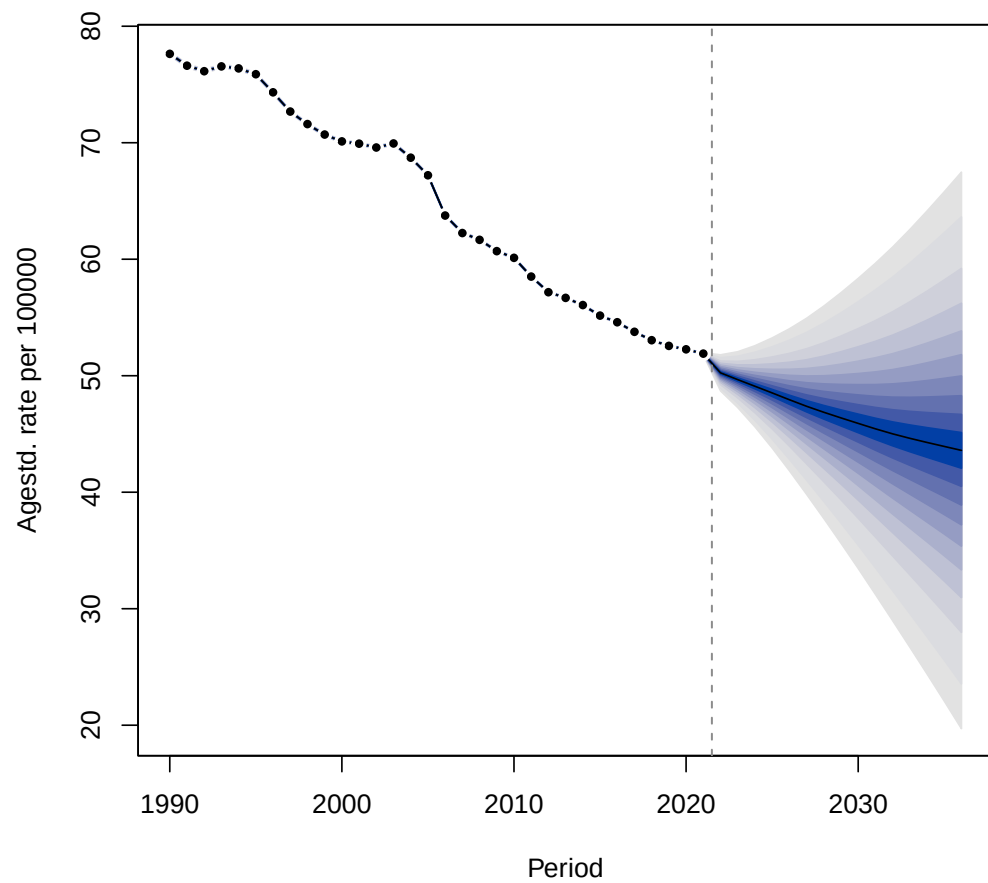


Figure S32: Projected number of new age-standardized DALY cases for ischemic stroke in males in 2036 according to the SDI.

DALYs = disability-adjusted life years. SDI = sociodemographic index.

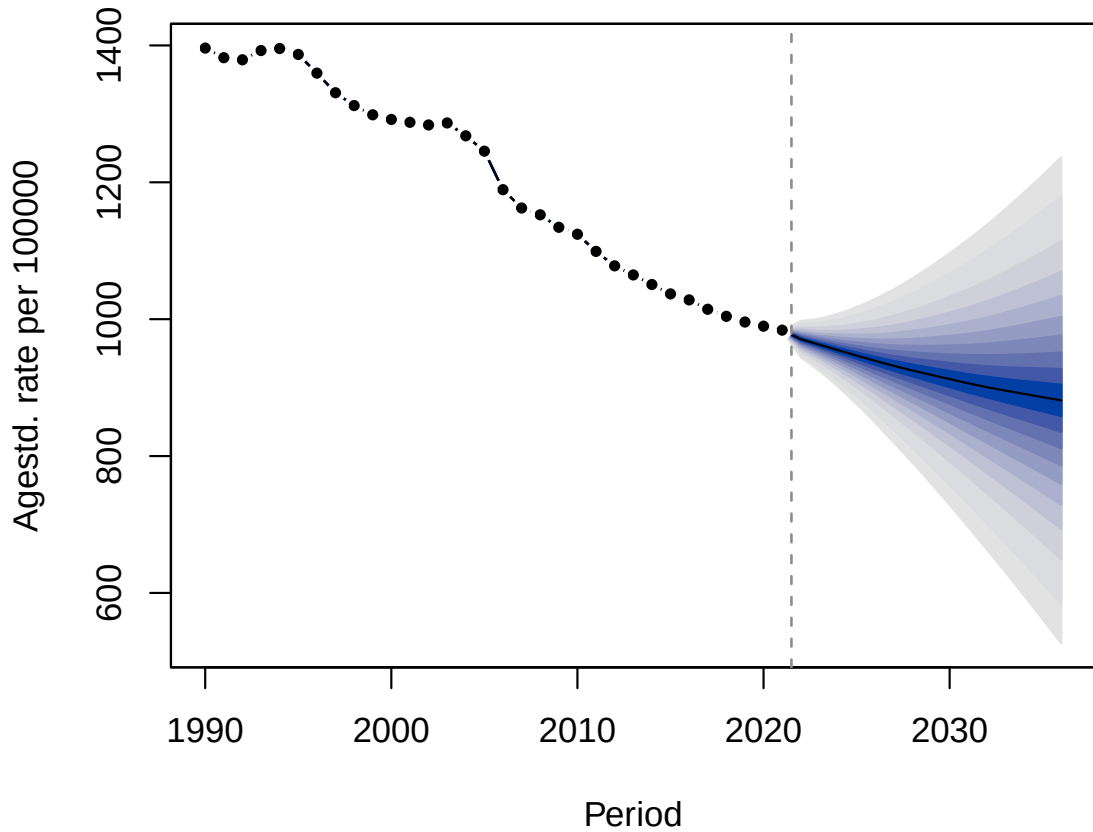


Figure S33: Projected number of ASPR cases for ischemic stroke in females in 2036 according to the SDI.

SDI = sociodemographic index. ASPR = Age-Standardized Prevalence Rate.

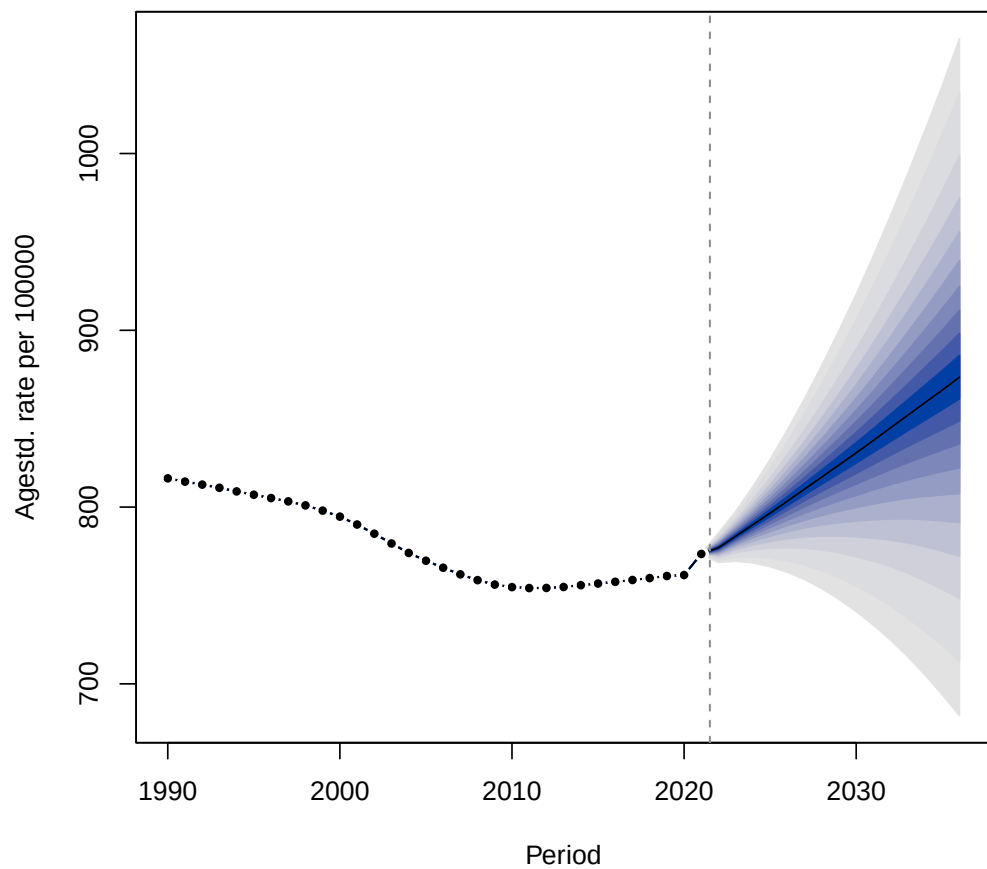


Figure S34: Projected number of ASIR cases for ischemic stroke in females in 2036 according to the SDI.

SDI = sociodemographic index. ASIR = Age-Standardized Incidence Rate.

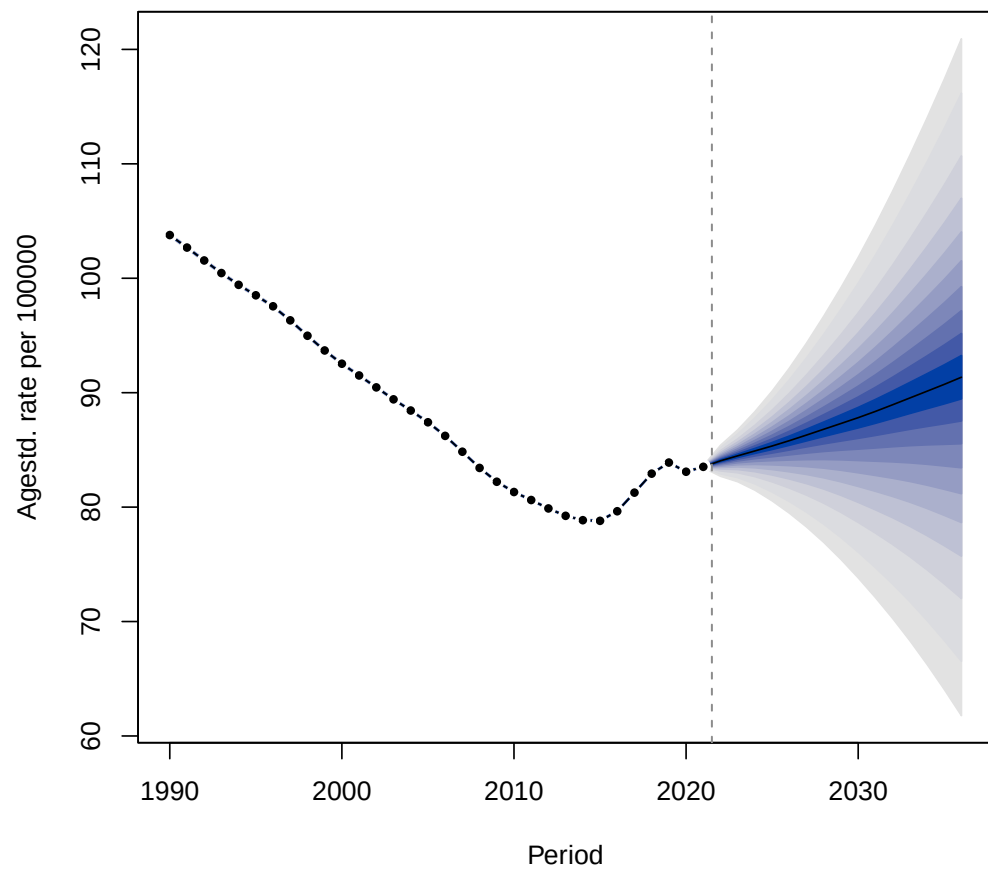


Figure S35: Projected number of ASDR cases for ischemic stroke in females in 2036 according to the SDI.

SDI = sociodemographic index. ASDR = Age-Standardized Death Rate.

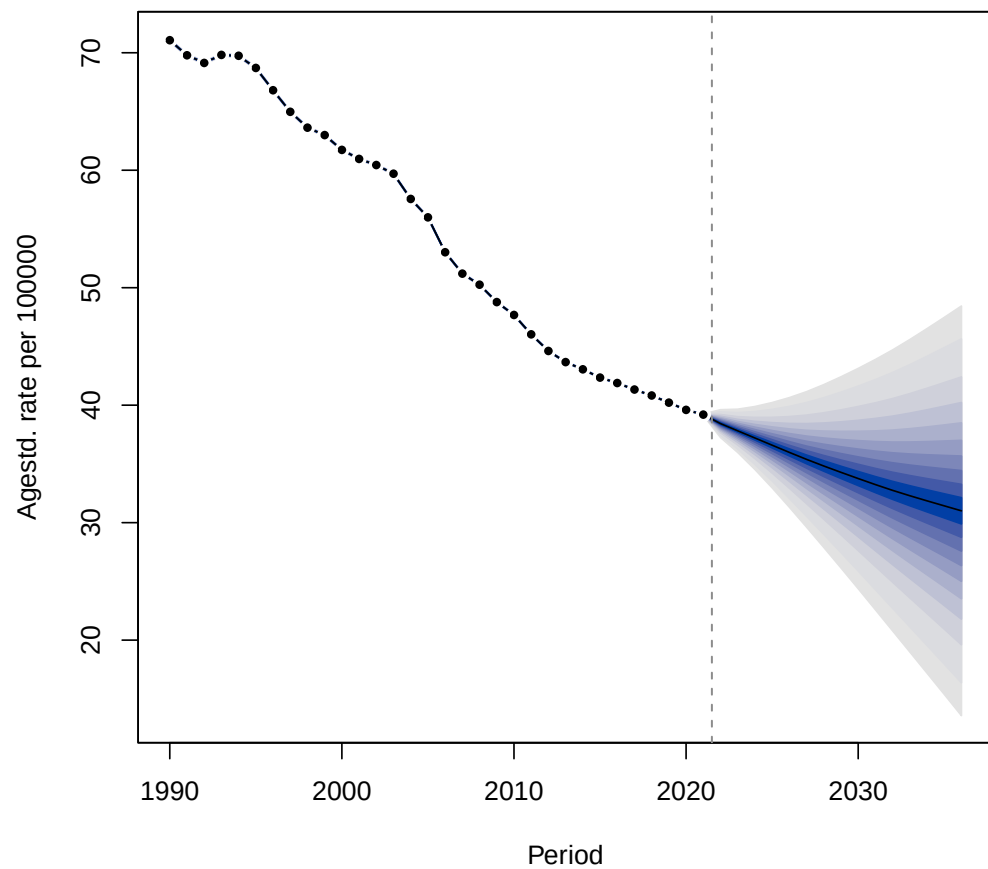
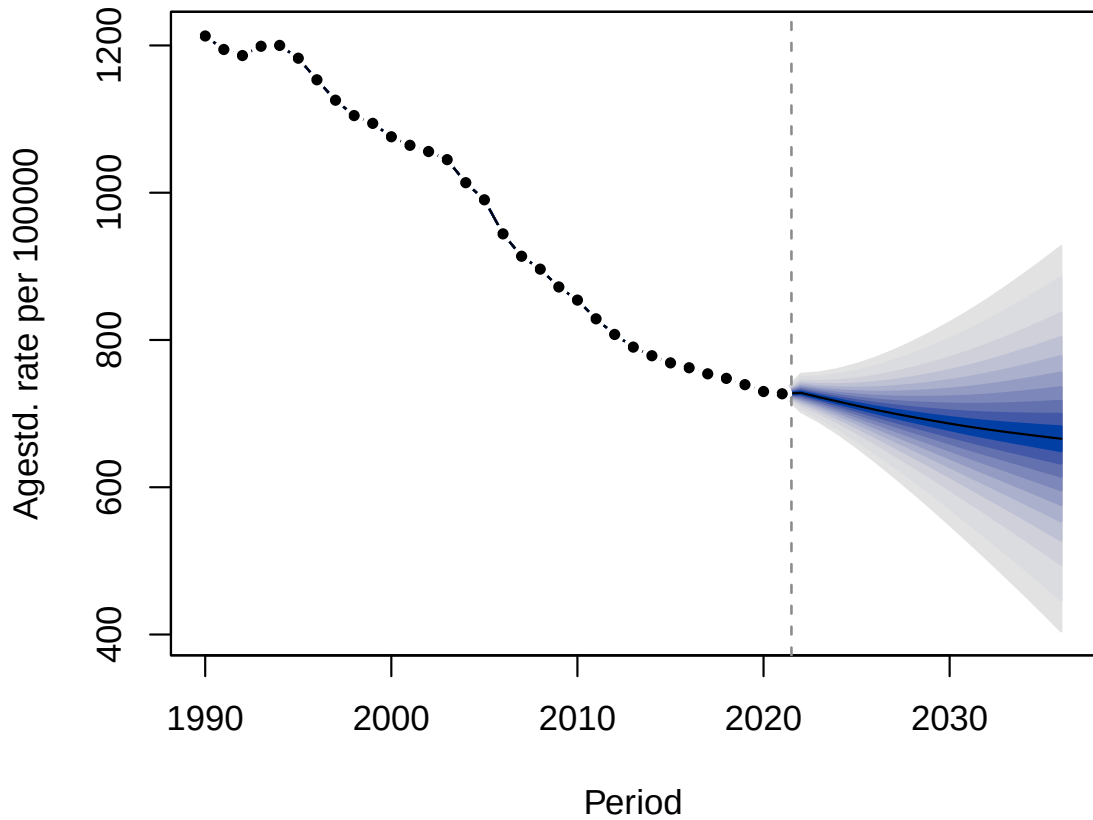


Figure S36: Projected number of new age-standardized DALYs cases for ischemic stroke in females in 2036 according to the SDI.

DALYs = disability-adjusted life years. SDI = sociodemographic index.



Supplementary Methods

1.Query data

Data indicators: GBD database contained all GBD diseases, risk, etiology, injury, etc. Indicators to measure the global burden of disease include death, years of life lost, disability-adjusted life years, prevalence, incidence, and life expectancy. The units of data indicators extracted included number, ratio, percentage, year, probability of death, etc. Annual results for all measures and all GBD age groups from 1990 to 2017 were available for extractable data years. The gender can be selected as male, female, or the combination of the two. The study areas were divided into GBD super regions, regions, countries, selected subnational units, and custom regions (WHO regions, World Bank income levels, etc.).

Search criteria: Click on the "Data tools and practices" drop-down arrow at the top of the GBD database homepage and then select "Data sources". In the new page, click on "VISIT OUR DATA CATALOG". In this interface, Select the data of "GBD 2021 data" or "All IHME data" according to the requirements, click "GBD 2021 data", click "GBD 2019 Data Resources" after seeing the interface, and click "query tool" on the new page. New users can register first and log in directly if they are already registered. The new interface can be divided into two parts, the left side is the search bar, the right side is the result bar, the left side can through GBD Estimate, Measure, Metric, Risk, Cause, Location, Age, Sex, Year of different options to limit the search conditions to obtain their desired data.

2.Data download

Download operation: After obtaining the required data through search, click "Charts" to obtain the trend chart, select "Download" under the left search bar, select "Submit" after the pop-up interface, you can directly select "here" in the new interface, follow the prompts to obtain the desired data, and you can also view your email. Click the link to jump to the download screen.

3.CODEm model calculation method and assumption basis

1. Model calculation method

CODEm (Cause of Death Ensemble model) is an ensemble model used to estimate and predict cause-specific mortality. The calculation method mainly includes the following steps:

Data collection and preprocessing:

Mortality data were collected from different data sources, including national Dynamic registries and oral autopsy reports.

Data were cleaned and preprocessed to remove ambiguous data and sites with alignment gaps.

Covariate selection:

Covariate selection algorithms were used to screen out combinations of covariates associated with a particular cause of death from a large number of possible models.

For example, when studying maternal mortality, 1984 possible models were considered.

Model construction and evaluation:

Multiple models, including mixed-effects linear models and spatio-temporal models, were constructed to account for temporal, geographic, and age correlations.

Each model was evaluated using out-of-sample predictive validity to select the best performing model.

Integrated model construction:

The best performing individual models were combined into an ensemble model (CODEm) to improve the accuracy and stability of prediction.

The performance of the integrated model was evaluated by the root mean square error, the frequency of correctly predicting the time trend, and the 95% coverage of the prediction interval.

2. Model assumption basis

The assumption basis of CODEm model mainly includes the following aspects:

Data distribution assumptions:

Assuming correlations in the data across time, geography, and age groups, spatial-temporal models were used to capture these correlations.

Associations of covariates with cause of death:

Assuming a linear or nonlinear relationship between covariates and cause-specific mortality, a covariate selection algorithm was used to identify the most relevant covariates.

Independence of the model:

It is assumed that the different individual models are independent of each other, and the accuracy of the prediction can be improved by combining these models.

Prediction interval assumptions:

The prediction interval of the model was assumed to cover 95% of the true value, which was used to evaluate the uncertainty of the model.

With these computational methods and assumptions, CODEm models can effectively estimate and predict cause-specific mortality and provide scientific basis for public health decision making.

4. DisMod-MR 2.1 Assumption basis and calculation method

1. Assumptions

DisMod-MR 2.1 is a Bayesian meta-regression tool for estimating epidemiological data on nonfatal diseases. Its assumption basis mainly includes the following aspects:

Data distribution assumptions:

It is assumed that there is an intrinsic link between indicators such as disease incidence, prevalence, remission rate, and case fatality rate, which are not independent.

Bayesian framework:

Bayesian methods were used to deal with the uncertainty and heterogeneity of the data, and the model parameters were estimated through the prior and posterior distributions.

Spatial and temporal correlation:

Assuming correlations between disease data across time, geographic regions, and age groups, spatial-temporal Gaussian process regression models were used to capture these correlations.

Data adjustment:

The meta-regression procedure of Bayesian regularized pruning (MR-BRT) was used to adjust the data bias under different case definitions and research methods, assuming bias between different data sources.

Hierarchy of the model:

It is assumed that the estimates of the model can be progressively refined from the

global level to the national and subnational levels, using a hierarchical model structure to progressively adjust the estimates.

2. Calculation method

The calculation method of DisMod-MR 2.1 mainly includes the following steps:

Data collection and preprocessing:

Epidemiological data were collected from different data sources, including vital registration, oral autopsy reports, population representative surveys and surveillance data.

Data were cleaned and preprocessed to remove ambiguous data and sites with alignment gaps.

Model construction:

The Bayesian meta-regression method was used to construct the model, considering the incidence, prevalence, remission rate and mortality of the disease.

Polynomial network meta-regression, logistic regression, and spatio-temporal Gaussian process regression were used to fit quantities and ratios.

Model fitting and adjustment:

The stochastic Markov chain Monte Carlo (MCMC) method was used to fit the model and estimate the model parameters.

Bayesian regularized pruning (MR-BRT) was used to adjust the bias between different data sources.

Hierarchical estimation:

Starting at the global level and progressively refining to the super-regional, regional, national, and subnational levels, a hierarchical model structure was used to progressively adjust the estimates.

Assessment of uncertainty:

The 95% uncertainty interval (95%UI) of each indicator was estimated to evaluate the uncertainty of the model.

Summary of results:

All subnational level estimates were pooled to the country level, then to the district and super-district levels, and finally to the global level.

With these assumptions and computational methods, DisMod-MR 2.1 can effectively estimate and predict the epidemiological data of non-fatal diseases, and provide scientific basis for public health decision-making.

5. ST-GPR

1. ST-GPR definition

Spatiotemporal Gaussian Process Regression (ST-GPR) is a spatiotemporal Gaussian process regression model that is mainly used to estimate and predict risk factor exposure and mortality in the Global Burden of Disease (GBD) project. It estimates a single indicator by exploiting the correlation between time and space, borrowing the strength of the data between different geographical locations and time points.

2. ST-GPR calculation method

The ST-GPR calculation method is a three-stage process:

Linear mixed-effects model:

A linear mixed-effects model was first run using a set of predictive covariates (e.g., GDP per capita, women's education, and energy consumption). These predicted values provide a general trend of the data.

Spatiotemporal pattern estimates:

In the second step, the spatio-temporal patterns are estimated by applying a series of spatio-temporal weights to average the residuals from the linear model of the first step. These spatio-temporal patterns are added to the linear predictions to generate spatio-temporal predictions.

Gaussian process regression:

Finally, the spatio-temporal prediction is regressed on the data across time as a function of the mean of the Gaussian process regression. The estimated results of the Gaussian process regression serve as the final ST-GPR forecast and generate the full time series estimates for 195 countries from 1990 to 2017.

3. ST-GPR assumption basis

The assumptions underlying ST-GPR include:

Spatial-temporal correlations: Disease data were assumed to be correlated across time, geographic regions, and age groups, and spatial-temporal Gaussian process regression models were used to capture these correlations.

Data adjustment: Assuming bias between different data sources, a meta-regression procedure using Bayesian regularized pruning (MR-BRT) was used to adjust data bias under different case definitions and study methods.

Hierarchical estimation: Assuming that the estimates of the model can be progressively refined from the global level to the national and subnational levels, a hierarchical model structure is used to progressively adjust the estimates.

With these calculation methods and assumptions, ST-GPR can effectively estimate and predict various indicators in the global Burden of Disease project, and provide scientific basis for public health decision-making.

4.BAPC

1. The assumption basis of BAPC

The Bayesian Age-Period-Cohort (BAPC) model is a Bayesian statistical model used to analyze and predict disease burden data, taking into account the effects of Age, Period, and Cohort. Its assumption basis mainly includes the following aspects:

Multiplicative effects of age, period, and cohort:

Suppose that the burden of disease (e.g., incidence, mortality, etc.) can be expressed as a multiplicative effect of age, period, and cohort, namely:

$$\ln(\text{ASR}) = \alpha + \mu_i + \beta_j + \gamma_k$$

Bayesian framework:

Bayesian method was used to estimate the model parameters, and the posterior distribution was obtained by integrating the prior information of the unknown parameters with the sample information.

Random walk model:

A second-order random walk (RW2) model was used to adjust for the effects of age, period, and cohort, assuming that the expected effects of being adjacent in time would be similar.

Integrated nested Laplace approximation (INLA) :

The use of the INLA method to directly approximate the posterior marginal distribution avoids mixing and convergence problems and shows a relatively low error rate.

Prior distribution:

The prior distributions of age, period, and cohort effects were assumed to be inverse gamma distributions, and a second-order random walk model was applied to adjust for overdispersion phenomena.

2. Calculation method of BAPC

The BAPC calculation method is based on the Bayesian statistical framework, and the specific steps are as follows:

Data preparation:

Data on disease burden were collected, including year, age group, incidence, and mortality.

Model construction:

The BAPC package and INLA package in R language were used to build the model.

The model can be expressed as follows.

$$\ln(\text{ASR}) = \alpha + \mu_i + \beta_j + \gamma_k$$

Model fitting:

The INLA method was used to fit the model and estimate the effects of age, period, and cohort.

Forecasting:

Use the well-fitted model to make predictions and generate forecasted values of disease burden for future time periods.

Model Evaluation:

The prediction performance of the model is evaluated using indicators such as Mean Squared Error (MSE), Mean Absolute Error (MAE), and Mean Absolute Percentage Error (MAPE).

Through these assumptions and calculation methods, the BAPC model can effectively analyze and predict disease burden data, providing scientific basis for epidemiological research and public health decision-making.

5.EAPC

EAPC (Estimated Annual Percentage Change) is used to measure the trend of a disease index (such as incidence, mortality, disability-adjusted life years, etc.) in a

specific time period. Its assumption basis mainly includes the following aspects:

Linearity assumption:

Assuming a linear relationship between the natural log value of the disease indicator and time, namely:

$$\ln(\text{ASR}) = \alpha + \beta \times \text{year} + \sigma$$

Where α is the intercept, β is the coefficient of the time variable, and σ is the error term.

Time continuity:

It is assumed that the annual change continues to be stable over the entire period of observation, that is, the trend of change is continuous.

Logarithmic transformation:

Assuming a linear relationship in the change of the disease indicator after natural log transformation helps to stabilize the variance and simplify the model.

2. Calculation method of EAPC

EAPC is calculated based on a linear regression model and the steps are as follows:

Data preparation:

Time series data of disease indicators (such as incidence, mortality, etc.) are collected, usually including year and corresponding index values.

Logarithmic transformation:

The natural logarithmic transformation of disease indicators was calculated as $\ln(\text{ASR})$, where ASR is the age-standardized rate.

Linear regression:

Linear regression models were used to fit $\ln(\text{ASR})$ as a function of year, namely:

$$\ln(\text{ASR}) = \alpha + \beta \times \text{year}$$

Here, β is the coefficient of the time variable.

Calculate EAPC:

244 EAPC was calculated from the regression coefficient β using the formula:

$$\text{EAPC} = (e^{\beta} - 1) \times 100$$

245

246 Where e is the base of the natural logarithm (approximately 2.718).

247 Confidence intervals were calculated as follows:

248 The 95% confidence interval of the EAPC was calculated, usually from the
249 confidence interval of the regression model:

$$\text{CI}_{\text{lower}} = (e^{\text{CI}_{\text{lower}}} - 1) \times 100$$

$$\text{CI}_{\text{upper}} = (e^{\text{CI}_{\text{upper}}} - 1) \times 100$$

250

251 Hypothesis testing:

252 Hypothesis testing on EAPC is usually translated into a t-test for the regression
253 coefficient β , and the test statistic is:

$$t = \frac{\beta}{\text{SE}(\beta)}$$

254

255 Where $\text{SE}(\beta)$ is the standard error of β .

256 We add the ARIMA model for prediction and compare it with the BAPC model, and
257 report the assumption basis and limitations of the two prediction models.

258 The ARIMA model was used to predict the same outcomes (incidence, prevalence,
259 mortality, and DALYs) from 2021 to 2036, and the results were compared to the
260 projections from the BAPC model. Both models produced consistent trends,
261 suggesting robustness in the long-term predictions.

262 The Bayesian age-period-cohort (BAPC) model relies on several key assumptions
263 that are critical to the validity and accuracy of the model:

264 Separability of age, period, and cohort effects

265 Basic Assumptions: The BAPC model assumes that time-related changes in health
266 outcomes can be separated into age, period, and cohort effects. The age effect refers
267 to the changes in an individual's life course as he or she ages. Period effects reflect
268 the impact of events occurring within a given year on health outcomes; Cohort effect
269 is the variation due to the different characteristics of the population in different birth

270 years.

271 The prior distribution of the parameters

272 Use of prior information: In a Bayesian framework, the BAPC model allows prior
273 information to be incorporated into the parameter estimation process. This means
274 that for an unknown parameter in the model, an appropriate prior distribution needs
275 to be specified for it. For example, a highly concentrated Laplace distribution can be
276 used as a prior for parameters with linear effects, whereas a diffusive normal
277 distribution can be used for slopes of age and period.

278 Similarity of neighboring effects

279 Adjacent effect Assumption: The BAPC model assumes that the adjacent effects of
280 age, period, and cohort are similar in time. This means that the model considers the
281 effect changes to be smooth and continuous between adjacent time points or age
282 groups.

283 Poisson distribution assumption of the data

284 Poisson distribution: In some cases, the BAPC model assumes that the data follow a
285 Poisson distribution. This assumption applies to count data, such as disease
286 incidence or mortality, etc., and it allows the model to use Poisson regression models
287 when dealing with such data.

288 These assumptions provide the theoretical basis and computational framework for
289 the BAPC model, but also limit the scope of model application and the interpretation
290 of results. In practice, careful consideration needs to be given to whether these
291 assumptions are applicable to specific research contexts and data characteristics.

292 The assumptions of the BAPC model have significant impacts on the analysis results.
293 Here is a detailed explanation:

294 Separability of Age, Period, and Cohort Effects

295 Identification of Influencing Factors: The assumption that age, period, and cohort
296 effects can be separated allows the model to identify and quantify the roles of these
297 three effects in health outcomes. For example, if the incidence of a certain disease
298 increases over time, the model can isolate the period effect to determine whether this
299 is due to increased diagnoses from medical advancements or an actual rise in
300 disease occurrence due to environmental changes.

301 Result Interpretation: This assumption enables clearer interpretation of the results by
302 distinctly attributing changes in health outcomes to different factors. For instance, it

can clearly indicate that a high disease incidence in a specific age group is due to physiological changes unique to that age (age effect) or due to particular environmental conditions when that age group was born (cohort effect).

Prior Distributions of Parameters

Robustness of Estimation Results: The choice of prior distributions affects the estimation of parameters. If the prior information is accurate and consistent with the actual data, it can enhance the robustness and accuracy of the estimates. For example, using highly concentrated prior distributions can reduce the uncertainty in parameter estimation, bringing the results closer to the true values.

Subjectivity of Results: The setting of prior distributions may introduce subjectivity, as different priors can lead to different analysis results. This requires researchers to be cautious when setting priors and to provide sufficient justification for the choice of prior distributions.

Similarity of Adjacent Effects

Smoothness Constraint: Assuming similarity between adjacent effects essentially imposes a smoothness constraint on the changes in effects. This causes the model to tend to produce smooth curves in estimating effects, rather than abrupt fluctuations. For example, when estimating age effects, the model assumes that changes in incidence between adjacent age groups are gradual, not sudden jumps.

Reasonableness of Results: In many cases, the smoothness constraint is reasonable because changes in health outcomes are often continuous. However, in some special situations, if the actual effects do indeed change dramatically, this assumption may prevent the model from accurately capturing the true pattern of effects.

Poisson Distribution Assumption for Data

Applicability Limitations: The Poisson distribution assumption is suitable for count data, such as disease incidence or mortality rates. If the data do not conform to the characteristics of the Poisson distribution (such as overdispersion or zero inflation), this assumption can lead to poor model fit, thereby affecting the accuracy of the analysis results.

Bias in Results: When data do not meet the Poisson distribution assumption, it can cause bias in the estimation results. For example, if data exhibit overdispersion, using Poisson regression models may underestimate the variance, thus affecting the significance tests of parameters and the reliability of prediction results.

In summary, the assumptions of the BAPC model significantly impact the analysis

results, helping to identify and interpret the roles of different factors while also potentially limiting the model's applicability and accuracy in certain situations. In practical applications, it is necessary to test and adjust these assumptions based on specific data and research contexts to ensure the reliability and validity of the analysis results.

There are some limitations in the application of Bayesian age-period-cohort (BAPC) prediction model, which are analyzed as follows:

Difficulty in parameter estimation

Identification problems caused by linear relationships: There is a linear relationship between the three factors of age, period and cohort, which makes parameter estimation more difficult. The trend of change estimated using the APC model may even be reversed without certain conditional constraints.

Data dependency

Impact of data quality: The predictive accuracy of the BAPC model depends to some extent on the quality of the input data. If the data is missing, wrong or incomplete, it will affect the fitting effect of the model and the prediction result.

Limitations of data sources: In some studies, data may come from GBD and other databases, but these data themselves are estimated based on system dynamics model and statistical model, which has certain limitations.

Limitations of model assumptions

No support for inclusion of covariates: BAPC model does not support inclusion of covariates in the analysis process, which means that other influencing factors (such as socioeconomic factors and environmental factors) other than age, period and cohort cannot be considered, which limits the ability of the model to explain the complex causes of changes in disease burden.

Computational complexity

High computational resource demand: BAPC model usually adopts Bayesian Markov chain Monte Carlo (MCMC) algorithm for parameter estimation, which has high computational complexity and requires a large amount of computing resources and time.

Limitations of the prediction range

Uncertainty in long-term projections: Although the BAPC model can be used to predict future trends in disease burden, the prediction results of the model may be

370 more uncertain due to the greater uncertainty of future influencing factors when
371 making long-term projections.

372 These limitations need to be considered and overcome when applying the BAPC
373 model to improve the accuracy and reliability of the model.

374 To overcome the problem of data dependence of BAPC models, several strategies
375 can be adopted:

376 Data quality improvement

377 Data cleaning and preprocessing: The original data were thoroughly cleaned to
378 remove outliers, duplicate records, and inconsistent data to ensure the accuracy and
379 integrity of the data.

380 Data calibration: Calibration of data to reduce systematic and random errors by
381 comparison with other data sources or known facts.

382 Diversity of data sources

383 Integrating multiple data sources: In addition to traditional medical records and death
384 registration data, multiple data sources such as censuses, socioeconomic surveys,
385 and environmental monitoring can also be integrated to improve data
386 comprehensivity and representativity.

387 Utilizing alternative data: In cases where some data is missing, alternative data can
388 be sought to fill the gap. For example, sample survey data are used to estimate the
389 overall picture.

390 Model refinement and validation

391 Model Hypothesis testing: The hypotheses in the model are rigorously tested to
392 ensure that they are supported in the data. If some assumptions are found not to be
393 true, the model structure should be adjusted in time.

394 Cross-validation: Cross-validation was used to validate the model and evaluate its
395 performance on different datasets to ensure the stability and generalization ability of
396 the model.

397 Advantages of the Bayesian approach

398 Utilization of prior information: The Bayesian approach allows prior information to be
399 incorporated into the model, which can alleviate the problem of insufficient data to
400 some extent. Additional information support for the model is provided by combining
401 expert knowledge and historical data.

Uncertainty quantification: Bayesian methods are able to provide uncertainty quantification of parameter estimates, such as confidence intervals or probability distributions. This helps to assess the impact of data quality on model predictions.

Through these measures, the dependence of BAPC model on data can be effectively reduced, and the accuracy and reliability of model prediction can be improved.

The autoregressive integrated moving average (ARIMA) model is a statistical model widely used in time series analysis to predict future data points. Although the ARIMA model has been successful in many areas, it is based on a series of assumptions and suffers from some limitations. The following are the main assumptions and limitations of the ARIMA model:

The main assumptions of the ARIMA model

Stationarity:

Assumptions: Time series data are stationary, that is, their statistical properties (e.g., mean, variance) remain constant in time.

Limitations: Many real time series data are non-stationary, such as trends and seasonal variations in economic data. If the stationarity assumption is not met by the data, the predictive performance of the model may decrease significantly. Although it is possible to smooth the data through methods such as difference, excessive difference may lead to loss of information.

Linearity:

Assumption: The relationship in time series data is linear, that is, future values can be represented as linear combinations of past values.

Limitations: Many real time series data have nonlinear relationships, such as volatility clustering in financial markets. The inability of ARIMA models to capture these nonlinear features may lead to prediction bias.

Independent identically distributed (IID) error term:

Assumption: The error terms of the model are independently and identically distributed with zero mean and constant variance.

Limitations: In practice, the error term may not satisfy the assumption of independent identically distributed, such as the presence of autocorrelation or heteroscedasticity. This may lead to inaccurate estimation results of the model and affect the reliability of the prediction.

435 Model Order Selection:

436 Assumptions: The order of the model (p , d , q) is known or can be determined in
437 some way.

438 Limitations: Selecting the appropriate model order is a complex process that usually
439 relies on information criteria such as AIC, BIC, or methods such as cross validation.
440 Inappropriate order selection may lead to overfitting or underfitting of the model,
441 affecting the prediction performance.

442 Limitations of the ARIMA model

443 Limited capacity for non-stationary data:

444 The ARIMA model requires the data to be stationary, but many actual data are non-
445 stationary. Although it is possible to smooth the data through methods such as
446 difference, excessive difference may lead to loss of information and affect the
447 predictive ability of the model.

448 Failure to capture nonlinear relationships:

449 Based on the assumption of linearity, ARIMA model cannot effectively capture the
450 nonlinear relationship in time series. For data with nonlinear features, ARIMA models
451 may not provide accurate predictions.

452 Sensitive to outliers:

453 The ARIMA model is sensitive to outliers in the data, which may lead to inaccurate
454 model parameter estimation and thus affect the prediction results.

455 Model assumptions are strict:

456 The assumptions of the ARIMA model are strict, and it is often difficult for the actual
457 data to fully meet these assumptions. For example, the assumption of independent
458 identically distributed error terms may not hold in many cases, leading to inaccurate
459 estimation results of the model.

460 The forecast is limited:

461 ARIMA model is mainly used for short-term prediction, but may not be effective for
462 long-term prediction. This is because the model assumes that the statistical
463 properties of the data remain constant over the forecast period, whereas the
464 statistical properties of the actual data may change over time.

465 Model selection and parameter estimation are complex:

466 Selecting the appropriate ARIMA model order is a complex process that usually relies

on methods such as information criteria or cross-validation. Inappropriate order selection may lead to overfitting or underfitting of the model, affecting the prediction performance.

Summary

The ARIMA model is a powerful tool for time series analysis, but its assumptions and limitations need to be carefully considered when applied. For data that are non-stationary, nonlinear, or have outliers, it may be necessary to combine other models, such as seasonal ARIMA models, GARCH models, or machine learning methods, to improve the accuracy and reliability of prediction.

Compared with other prediction models, ARIMA model has the following advantages and disadvantages:

Advantages of the ARIMA model

The model is simple, easy to understand and apply: The structure of ARIMA model is relatively simple, and it is mainly modeled by three parts: autoregression (AR), difference (I) and moving average (MA), without the help of other exogenous variables.

High prediction accuracy: For short-term prediction, ARIMA model has high prediction accuracy and reliability, and can better capture the trend and seasonality of time series.

High maturity: ARIMA model has a solid theoretical foundation, and after long-term development and application, it has formed a set of relatively complete theoretical system.

Suitable for non-seasonal data: ARIMA model is suitable for time series data with no obvious seasonal characteristics, and can smooth non-stationary time series by difference operation.

Disadvantages of the ARIMA model

The requirement for data stationarity is high: ARIMA model requires time series to be stationary, or to be stationary after difference operation. If there is an obvious trend or seasonality in the data, it needs to be smoothed first, which may affect the prediction effect of the model.

Parameter selection is complex: ARIMA model needs to determine three parameters p (autoregressive order), d (difference order) and q (moving average order). Improper

parameter selection may lead to overfitting or underfitting of the model.

Only linear relationships can be captured: The ARIMA model can intrinsically only capture linear relationships, and the prediction effect may not be good for data with nonlinear relationships.

The long-term prediction effect is poor: the ARIMA model will gradually decrease the prediction accuracy due to the accumulation of errors in the long-term prediction.

The comparison between ARIMA model and BAPC model in time series prediction is as follows:

ARIMA model

Advantages:

The model is simple, easy to understand and apply: The structure of ARIMA model is relatively simple, and it is mainly modeled by three parts: autoregression (AR), difference (I) and moving average (MA), without the help of other exogenous variables.

High prediction accuracy: For short-term prediction, ARIMA model has high prediction accuracy and reliability, and can better capture the trend of time series.

High maturity: ARIMA model has a solid theoretical foundation, and after long-term development and application, it has formed a set of relatively complete theoretical system.

Suitable for non-seasonal data: ARIMA model is suitable for time series data with no obvious seasonal characteristics, and can smooth non-stationary time series by difference operation.

Disadvantages:

The requirement for data stationarity is high: ARIMA model requires time series to be stationary, or to be stationary after difference operation. If there is an obvious trend or seasonality in the data, it needs to be smoothed first, which may affect the prediction effect of the model.

Parameter selection is complex: ARIMA model needs to determine three parameters p (autoregressive order), d (difference order) and q (moving average order). Improper parameter selection may lead to overfitting or underfitting of the model.

Only linear relationships can be captured: The ARIMA model can intrinsically only capture linear relationships, and the prediction effect may not be good for data with nonlinear relationships.

The long-term prediction effect is poor: the ARIMA model will gradually decrease the prediction accuracy due to the accumulation of errors in the long-term prediction.

The BAPC model

Advantages:

Comprehensive consideration of multiple factors: BAPC model took into account the influence of age, period and cohort on the distribution of NCDS, which could more comprehensively analyze the trends of incidence and mortality of NCDS.

Unique advantages of the Bayesian framework: In the Bayesian framework, BAPC model can use both sample information and prior information to obtain unique parameter estimates, which makes the results more robust and reliable.

Improved prediction accuracy: Compared with the traditional APC model, the BAPC model can obtain more stable and reliable results and improve the prediction accuracy.

Solving the difficulty in parameter estimation: There is a linear relationship among the three factors of age, period, and cohort in the classical APC model, which may lead to difficulty in parameter estimation. The BAPC model effectively solves this problem through the Bayesian approach.

Flexible model selection: The BAPC model allows the flexible selection of different model components and prior probability distributions, which improves the applicability and flexibility of the model.

Advanced computational methods: The BAPC model uses the ensemble nested Laplace approximation (INLA) for posterior marginal distribution approximation, which avoids mixing and convergence problems and shows a relatively low error rate.

Automatic and simplified calculation: The use of the BAPC package for parameter estimation not only does not require complex coding procedures and convergence checks, but also allows flexible selection of age-standardized ratio (ASR) parameters and prior probability distributions to simplify the calculation process.

Predicting future changes: BAPC models can not only describe trends in age, period, and birth cohort, but also predict future changes based on these trends.

Disadvantages:

High computational complexity: The computational process of BAPC model is relatively complex and requires high computational resources and time.

High data requirements: BAPC models require sufficient data to estimate model

566 parameters, and for cases with a small amount of data, the accuracy and reliability of
567 the model may be affected.

568 There are many model assumptions: the BAPC model is based on a Bayesian
569 framework and requires assumptions about the prior distribution, which may affect
570 the prediction results of the model.

571 Comparison

572 Applicable scenarios:

573 ARIMA model: Suitable for time series with no or very weak seasonality. The effect
574 was poor when there was obvious seasonality in the data.

575 BAPC model: suitable for disease prediction that requires comprehensive
576 consideration of age, period and cohort factors, especially in the field of public health
577 and epidemiology.

578 Predictive power:

579 ARIMA model: performs well in short-term forecasting, but may not perform well for
580 long-term forecasting and data with nonlinear relationships.

581 BAPC model: It can more accurately predict the change trend of diseases in different
582 ages, periods and cohorts, and is suitable for long-term prediction and complex data
583 structure.

584 Model complexity:

585 ARIMA model: The model structure is relatively simple, with few parameters, which is
586 easy to understand and apply.

587 BAPC model: The model structure is complex and requires Bayesian estimation and
588 complex calculations.

589 Data requirements:

590 ARIMA model: It has high requirements for the stationarity of the data and needs pre-
591 processing such as difference.

592 BAPC model: Sufficient data are needed to estimate the model parameters, and high
593 data volume and data quality requirements are required.