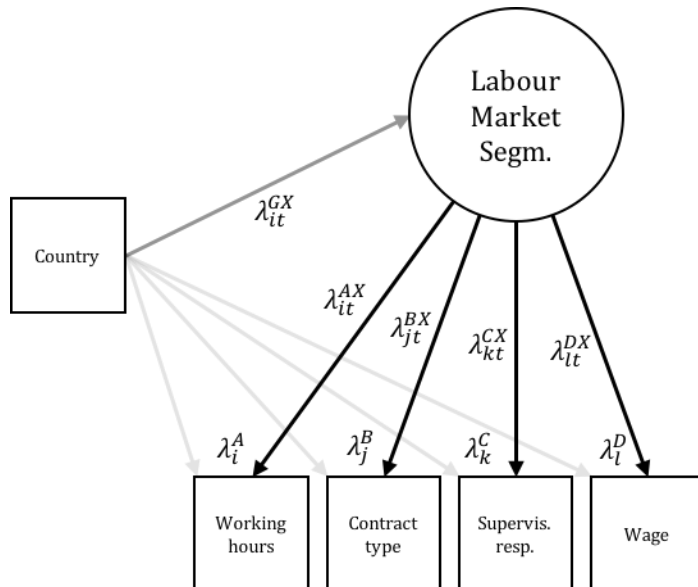


1: Additional information regarding the quantitative research approach

To create a model of labour market segmentation, we use LC analysis (Hagenaars and McCutcheon, 2002) — a model-based clustering method that creates groups of observations that are similar on a number of characteristics. Specifically, LC analysis allows us to identify discrete latent classes (i.e. labour market groups or segments) from a set of observed indicator variables (McCutcheon, 1987). The obtained labour market segments can then be used in subsequent analysis as independent variables to further investigate their population. The use of LC analysis is based on theoretical deliberation that observed indicator variables are statistically associated due to an unobserved common factor rather than being causally related (McCutcheon, 2002). LC analysis thus offers scope to investigate patterns and characteristics of labour market segmentation that are not directly observable. The applied model can be found in Figure 1, presenting the conceptual model used in the LC analysis. It shows the variables of working hours, contract type, supervisory responsibilities, and wage which serve as indicator variables for the unobserved variable of labour market segment. Although each of these indicator variables has a country-specific intercept, they all share the same links with the latent variable (i.e. λ_{it}^{AX} , etc). The arrow between the country variable and the unobserved segment variable allows us to model various proportions of segments occurring per country.

Figure A1: Measurement LC model of labour market segmentation



The LC model fit is conventionally evaluated by several model fit criteria: the likelihood-ratio (L^2), Pearson's chi-squared (χ^2), and information criteria (IC) indices. As the former two options have a number of limitations (see McCutcheon, 2002), we rely on the information criteria. According to a simulation study by Nylund et al. (2007), the Bayesian information criterion (BIC) and sample-size adjusted Bayesian information criterion (SABIC) were the most successful in identifying the correct number of classes with larger sample sizes and thus will also guide our model selection. Information criteria have limitations too, as they tend towards overly parsimonious models and there are numerous other studies that investigate the efficiency of other model fit indices for latent variable models (see for example Tofighi and Enders 2007; Celeux and Soromenho, 1996; Soromenho, 1993). Nevertheless, working with large samples, model fit criteria would always justify the inclusion of additional latent classes up to a point where the model starts extracting irrelevantly small classes with extreme observations. For this reason, the final model will be selected also based on its practical ability to extract large and robust segments. Despite various suggestions, there is no commonly accepted best criterion to determine the number of classes in latent class models. To address these issues and to provide some sensitivity analysis, we also evaluate information criteria with the five-fold

cross-validation procedure, where repeatedly four parts of the sample are used to fit the model and the remainder is used to evaluate its model fit (Donovan and Chung, 2015; Collins et al., 1994).

The variables displayed in figure 1 are used as indicators in the following multi-group LC model:

$$\pi_{ijklts}^{ABCDX|G} = \pi_{ts}^{X|G} \pi_{its}^{A|XG} \pi_{jts}^{B|XG} \pi_{kts}^{C|XG} \pi_{lts}^{D|XG}$$

where A - D stand for working hours, contract type, supervisory responsibilities, and wage respectively, and G stands for country. Overall, $\pi_{ts}^{X|G}$ stands for the probability of being in a labour market segment t , given that the observation comes from country s . Conditional indicator probabilities ($\pi_{its}^{A|XG} \dots \pi_{lts}^{D|XG}$) can be expressed as log-linear equations:

$$\pi_{its}^{A|XG} = \frac{\exp(\lambda_i^A + \lambda_{it}^{AX} + \lambda_{is}^{AG})}{\sum_i \exp(\lambda_i^A + \lambda_{it}^{AX} + \lambda_{is}^{AG})}$$

where λ_{it}^{AX} is of highest interest, as it shows the link between a respective indicator (here A) and the latent class (X). The log-linear parameter λ_{is}^{AG} stands for different intercept probabilities per indicator by country. The measurement model used to model segmentation is visualized in figure 1. For data manipulation and cleaning, we used the software R (i.e. R Core Team, 2018), RStudio (i.e. RStudio Team, 2016), and dplyr package (Wickham and Francois, 2016). Data visualizations were made by ggplot2 package (Wickham, 2009). Models were estimated in LatentGOLD 5.1 (Vermunt and Magidson, 2005), using the ‘Cluster’ option for estimating the measurement model and the ‘Step3’ option for estimating associations with covariates.

2: Model fit criteria

Table A1 and Figure A2 show model fit indices, indicating that the best fit is achieved by the 5- and 6-class models. The relative improvement in BIC and SABIC decreases significantly after the 5- and 6-class models for both 5-fold cross-validation and pooled model fit. BIC from cross-validation even indicates that the 6-class model fits the best (lowest value across classes). Yet, having inspected the extracted labour market segment profiles, we select to continue our analyses with five classes, based on the interpretability of the solution. The 6-class model adds a labour market group that splits one of the segments into conceptually similar groups with slightly different distributions and adds little to our interpretation. Furthermore, Figure 2 illustrates that moving from the 5- to 6-class model adds relatively little to the model fit. Hence, we argue that there are five labour market segments, which demonstrates that there is substantively more fragmentation than assumed by binary approaches.

Table A1: Model fit indices

Model	L^2	Df	5-fold cross-validation		Pooled model fit	
			BIC (LL)	SABIC (LL)	BIC(L^2)	SABIC(L^2)
3-class	11280.67	38	1615546	1615447	10796.45	10917.22
4-class	1095.28	30	1605483	1605359	713.01	808.35
5-class	479.58	22	1604997	1604848	199.24	269.16
6-class	293.59	14	1604758	1604583	115.20	159.69
7-class	29.84	6	1604782	1604581	-46.62	-27.55

Figure A2: Model fit indices

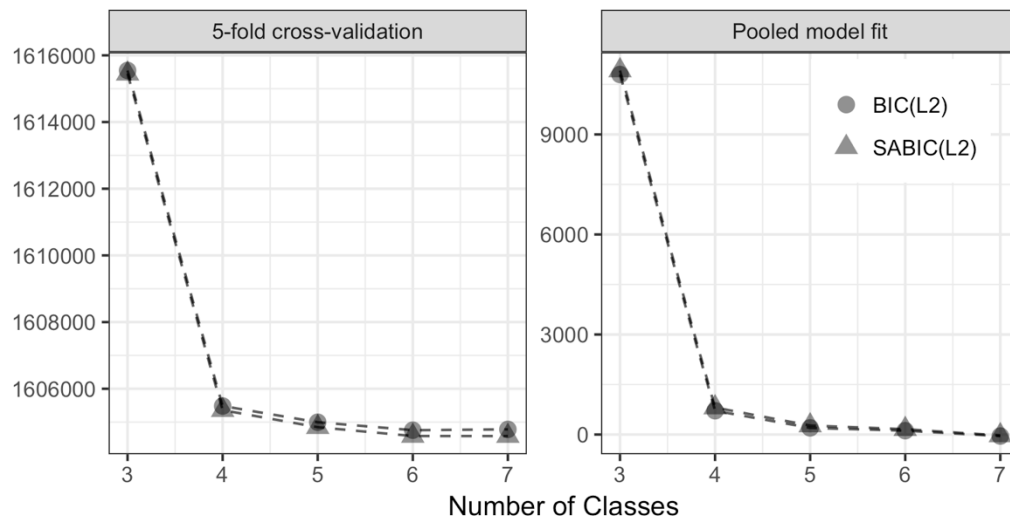
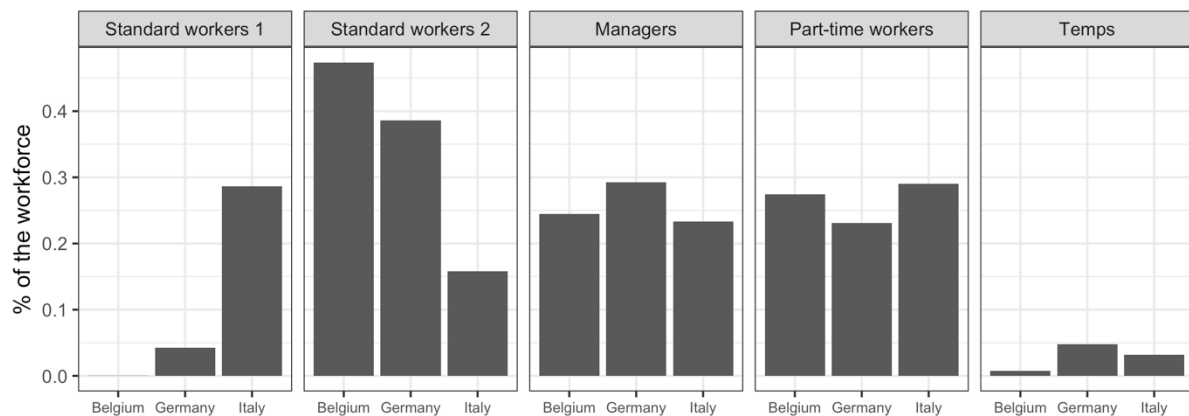


Figure A3

Group sizes per country



References

- Celeux G and Soromenho G. (1996) An entropy criterion for assessing the number of clusters in a mixture model. *Journal of Classification* 13(2): 195-212.
- Collins LM, Graham JW, Long JD and Hansen WB (1994) Crossvalidation of latent class models of early substance use onset. *Multivariate Behavioral Research* 29(2): 165-183.
- Donovan JE and Chung T (2015) Progressive Elaboration and Cross-Validation of a Latent Class Typology of Adolescent Alcohol Involvement in a National Sample. *Journal of Studies on Alcohol and Drugs* 76(3): 419-429.
- Hagenaars JA and McCutcheon A (2002) *Applied latent class analysis*. Cambridge: Cambridge University Press.
- Kankaras M, Moors G and Vermunt, J (2010) Testing for measurement invariance with latent class analysis. In Davidov P, Schmidt P and Billiet J (eds.) *Cross-cultural analysis: methods and applications*. London: Routledge, 361-385.
- McCutcheon A (1987) *Latent class analysis*. London: SAGE.
- McCutcheon A (2002) Basic concepts and procedures in single- and multi-group latent class analysis. In Hagenaars J and McCutcheon A (eds.) *Applied latent class analysis*. Cambridge: Cambridge University Press, 56-88.
- Nylund KL, Asparouhov T and Muthén B (2007) Deciding on the number of classes in latent class analysis and growth mixture modeling: A monte carlo simulation study. *Structural Equation Modeling: A Multidisciplinary Journal* 14(4): 535-569.
- Soromenho G (1994) Comparing approaches for testing the number of components in a finite mixture model. *Computational Statistics* 9(1): 65-78.
- Tofighi D and Enders CK (2007) Identifying the correct number of classes in a growth mixture model. In: Hancock GR (ed.) *Advances in latent variable mixture models*. Greenwich: Information Age, 317-341.
- Vermunt J and Magidson J (2005) Latent Gold 4.0 user's guide. Available at <https://www.statisticalinnovations.com/wp-content/uploads/LGusersguide.pdf>.
- Wickham H (2009) *ggplot2: Elegant Graphics for Data Analysis*. Springer: New York.
- Wickham H and Francois R (2016) *dplyr: A Grammar of Data Manipulation*. R package version 0.5.0. Available at <https://CRAN.R-project.org/package=dplyr>.