

Supplemental Materials

Towards a Typology of Risk Preference: Four Risk Profiles Describe Two Thirds of
Individuals in a Large Sample of the U.S. Population

Renato Frey^{1,2}, Shannon M. Duncan^{3,4}, and Elke U. Weber²

¹University of Zurich

²Princeton University

³Columbia University

⁴University of Pennsylvania

February 3, 2022

Contents

Complementary Analyses and Robustness Checks	3
Criteria for model fits in the confirmatory factor analyses (CFA)	3
Variance decomposition	3
Relationship of factor values with risk–return framework	3
Relationship of factor values with SOEP risk items	4
Illustration of overfitting in traditional implementations of GMMs	22

List of Figures

S1 Psychometric structure of model 1.	10
S2 Psychometric structure of model 2.	11
S3 Psychometric structure of model 3.	12
S4 Psychometric structure of model 4.	13
S5 Psychometric structure of model 5.	14
S6 Psychometric structure of model 6.	15
S7 Comparison of model fits.	16
S8 Correlation of factor values across models and with additional indicators.	17
S9 Distribution of factor scores extracted from M5.	18
S10 Screening analysis to identify the optimal number of clusters.	19
S11 Cluster enrichment analysis.	20
S12 Classification probabilities.	21

List of Tables

S1 Demographic information	6
S2 Models of risk preference tested using confirmatory factor analyses (CFA)	7
S3 Model fits (full sample)	8
S4 List of items used in the DOSPERT scale	9

Complementary Analyses and Robustness Checks

Criteria for model fits in the confirmatory factor analyses (CFA)

Conventional criteria for good model fits include a root mean square error of approximation (RMSEA) of smaller than .05, a standardized root mean square residual (SRMR) of smaller than .08, and both a comparative fit index (CFI) and a Tucker-Lewis index (TLI) of greater than .9. We also report the Bayesian Information Criterion (BIC) and the Akaike information criterion (AIC).¹ Finally, for each model we report the proportion of explained variance (R^2) that is accounted for by i) all factors, ii) the general factor only (if the model involves a general factor), and iii) the specific factors only (if the model involves specific factors).

Variance decomposition

In the selected model (i.e., M5), the general factor accounted for 15% of the total variance and the domain-specific factors cumulatively accounted for 27% of the total variance. That is, of the explained variance (43%), R accounted for somewhat more than one third (i.e., 38%) whereas the six specific factors accounted for approximately the other two thirds (i.e., 62%). Very similar shares of the explained variance between the general and the specific factors were observed in the other two bifactor models M4 and M6.

In comparison to the observations made in past research on risk preference (Frey et al., 2017) and regarding other major personality traits (e.g., intelligence or psychopathology; Caspi et al., 2014; Deary, 2012; Spearman, 1904), the share of variance accounted for by the general factor was somewhat smaller in the present setting. Yet, considering that the DOSPERT scale has been intentionally designed as a domain-specific instrument, the proportion of variance accounted for by the general factor still turned out to be quite high.

Relationship of factor values with risk–return framework

We adopted a complementary perspective to model general risk-taking propensity by implementing a risk–return framework (Weber et al., 2002; Weber & Milliman, 1997). The central assumption of this framework is that inter- and intra-individual differences in risk-taking propensity result because people may vary in terms of either the risks they perceive or the benefits they expect from engaging in risky activities (e.g., domain-specific risk-taking propensity may result from domain-specific differences in risk perception). Yet, the influence of risk perception and expected

¹Note that these two likelihood-based information criteria cannot be compared across different samples, as they directly depend on sample size. However, they are useful for comparing models within the same sample.

benefits on risk-taking propensity might be relatively homogeneous across people and domains (i.e., a negative and a positive relationship, respectively). Therefore, “controlling” for any differences in the proximal processes potentially underlying risk-taking propensity may give rise to people’s general risk-taking propensity, over and beyond specific risk perceptions and expected benefits.

We tested this assumption with a mixed-effects model using the predictors “perceived risk” and “expected benefits” and the dependent variable “risk-taking propensity”. A model with random intercepts across participants as well as random slopes for the two predictors turned out to have parameters that might not be perfectly identifiable, and we thus used a mixed-effects model with random intercepts across participants only and fixed effects for the two predictors.

We used Bayesian estimation techniques as implemented in the R-package *rstanarm* (Stan Development Team, 2016), with the prior $\mathcal{N}(0,10)$ for the intercept and $\mathcal{N}(0,2.5)$ for the predictors. Weakly informative priors provide some statistical regularization and thus guard against overfitting the data. Three chains with 2000 iterations were run per model.

The results indicated a posterior mean of 4.54 (95% highest-density interval [HDI] = 4.5 to 4.59) for the intercept (i.e., the average level of risk-taking propensity were perceived risks and expected benefits equal to zero). As expected, the posterior mean for the fixed effect of risk perception was negative (-.55, HDI = -.56 to -.54) and positive for the fixed effect expected benefit (.34, HDI = .33 to .35).

The random intercepts (i.e., the person-specific risk-taking propensity were risk perception and expected benefits equal to zero) ranged from 2.3 to 7.1 with a standard deviation of .51. These person-specific intercepts correlated highly with the general factor of M5, namely, at .7 (see Fig. S8 for the exact correlations with the other factors of M5). Figure 3 of the main text further illustrates that these intercepts (denoted with “I”) clustered closely to the center of the network of the various measures of risk-taking propensity.

Relationship of factor values with SOEP risk items

We also explored to what extent the extracted factors converged with conceptually related assessments of risk-taking propensity, such as the general and the five domain-specific single-item questions implemented in the German Socio-Economic Panel (SOEP).

As Figure 3 of the main text shows, the general risk item of the SOEP clustered close to the center of the network. It was correlated at .43 with the general factor extracted in M5 (see Fig. S8 for all correlations between the SOEP items and the other factors of M5). Moreover, as can be further

seen in Figure 3, the domain-specific SOEP items were correlated more strongly with each other (and were thus in fact less domain-specific), as compared to the domain-specific factors extracted from the bifactor models (and in particular, M5; see also Fig. [S8](#)).

Table S1: Demographic information			
Variables		N	%
Gender	Female	1,569	50.20%
	Male	1,554	49.80%
Age	18-29	1,107	35.40%
	30-39	1,079	34.50%
	40-49	513	16.40%
	50-59	291	9.30%
	60-69	117	3.70%
	70-79	16	0.50%
Ethnicity	Black or African American	259	8.30%
	American Indian or Alaskan	36	1.20%
	White	2,637	84.40%
	Asian	185	5.90%
	Hawaiian or Pacific Islander	6	0.20%
Education	No degree	30	1.00%
	High school diploma or GED	962	30.80%
	Associates Degree, occupational	251	8.00%
	Associates Degree, academic	315	10.10%
	Bachelor's Degree	1,177	37.70%
	Master's Degree	308	9.90%
	Professional Degree	41	1.30%
	Doctoral Degree	39	1.20%
Political	Democrat	1,482	47.50%
	Independent	1,004	32.10%
	Republican	637	20.40%
Income	Less than \$25,000	663	21.20%
	25,000–34,999	479	15.30%
	35,000–49,999	548	17.50%
	50,000–74,999	723	23.20%
	75,000–99,999	371	11.90%
	100,000–149,999	262	8.40%
	\$150,000 or more	77	2.50%

Table S2: Models of risk preference tested using confirmatory factor analyses (CFA)

Model	Number of factors	Items	Factors
M1	One factor	30	-
M2	Five factors (original)	30	oblique
M3	Six factors (revised)	26	oblique
M4	One general factor + five factors (original)	30	orthogonal
M5	One general factor + six factors (revised)	30	orthogonal
M6	One general factor + four factors (revised)	28	orthogonal

Table S3: Model fits (full sample)

Model	DF	RMSEA	SRMR	CFI	TLI	BIC	AIC	R ²	R ² (R)	R ² (Fs)
M1	405	0.12	0.11	0.43	0.39	350597	350234	0.19	0.19	–
M2	395	0.07	0.07	0.80	0.78	338741	338317	0.38	–	0.38
M3	284	0.05	0.05	0.90	0.88	291072	290667	0.41	–	0.41
M4	375	0.06	0.08	0.86	0.84	337080	336536	0.43	0.15	0.27
M5	378	0.05	0.06	0.89	0.88	335967	335441	0.44	0.17	0.28
M6	322	0.06	0.05	0.89	0.87	315267	314759	0.41	0.17	0.24

Note. DF = degrees of freedom, RMSEA = root mean square error of approximation, SRMR = standardized root mean square residual, CFI = comparative fix index, TLI = Tucker-Lewis index, BIC = Bayesian Information Criterion, AIC = Akaike information criterion, R² = Explained variance (in total, R = by the general factor, Fs = by the specific factors).

Table S4: List of items used in the DOSPERT scale

Item	Activity
eth1	Taking some questionable deductions on your income tax return.
eth2	Having an affair with a married man/woman.
eth3	Passing off somebody else's work as your own.
eth4	Revealing a friend's secret to someone else.
eth5	Leaving your young children alone at home while running an errand.
eth6	Not returning a wallet you found that contains \$200.
fin1	Betting a day's income at the horse races.
fin2	Investing 10% of your annual income in a moderate growth diversified fund.
fin3	Betting a day's income at a high-stakes poker game.
fin4	Investing 5% of your annual income in a very speculative stock.
fin5	Betting a day's income on the outcome of a sporting event.
fin6	Investing 10% of your annual income in a new business venture.
hea1	Drinking heavily at a social function.
hea2	Engaging in unprotected sex.
hea3	Driving a car without wearing a seat belt.
hea4	Riding a motorcycle without a helmet.
hea5	Sunbathing without sunscreen.
hea6	Walking home alone at night in an unsafe area of town.
rec1	Going camping in the wilderness.
rec2	Going down a ski run that is beyond your ability.
rec3	Going whitewater rafting at high water in the spring.
rec4	Taking a skydiving class.
rec5	Bungee jumping off a tall bridge.
rec6	Piloting a small plane.
soc1	Admitting that your tastes are different from those of a friend.
soc2	Disagreeing with an authority figure on a major issue.
soc3	Choosing a career that you truly enjoy over a more prestigious one.
soc4	Speaking your mind about an unpopular issue in a meeting at work.
soc5	Moving to a city far away from your extended family.
soc6	Starting a new career in your mid-thirties.

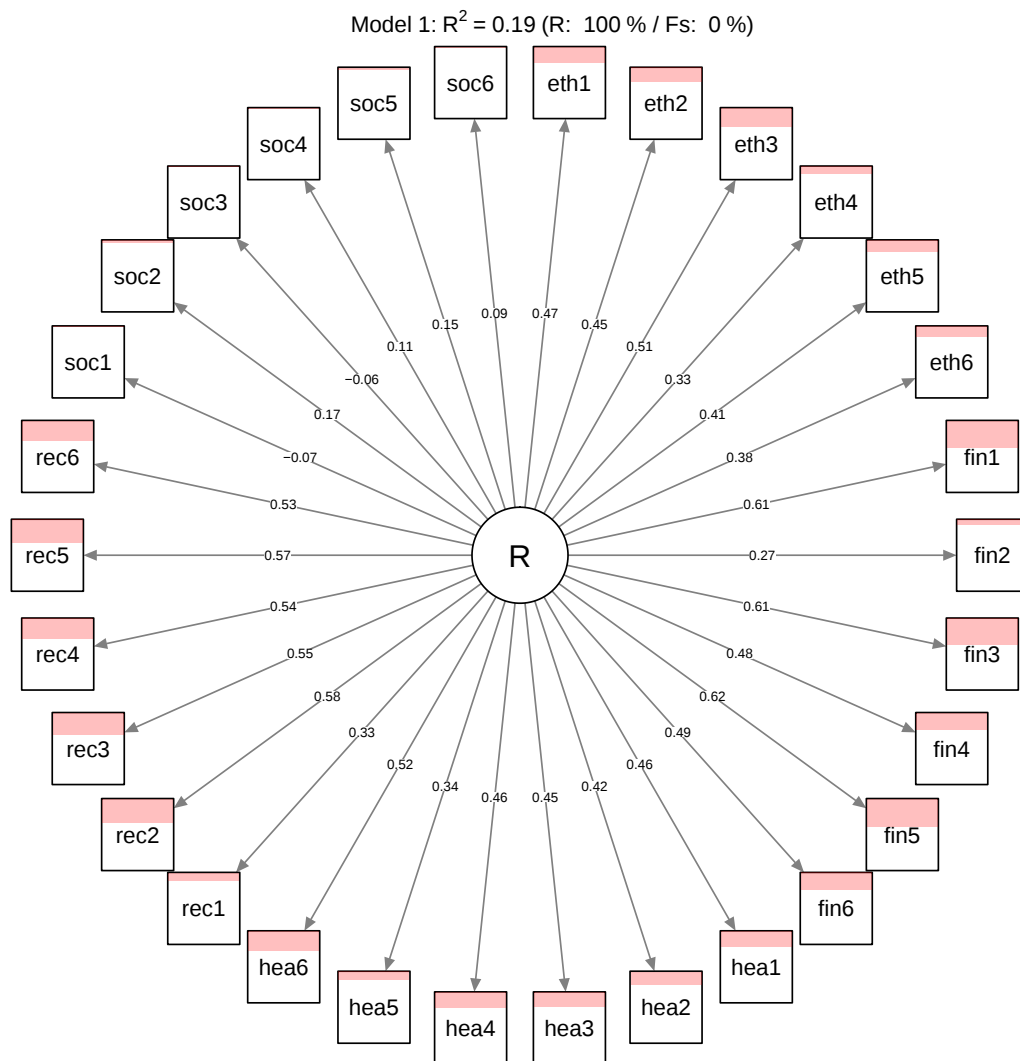


Figure S1: Depiction of the factor structure implemented in M1. The red shaded areas depict the proportion of variance that is accounted for by the general factor.

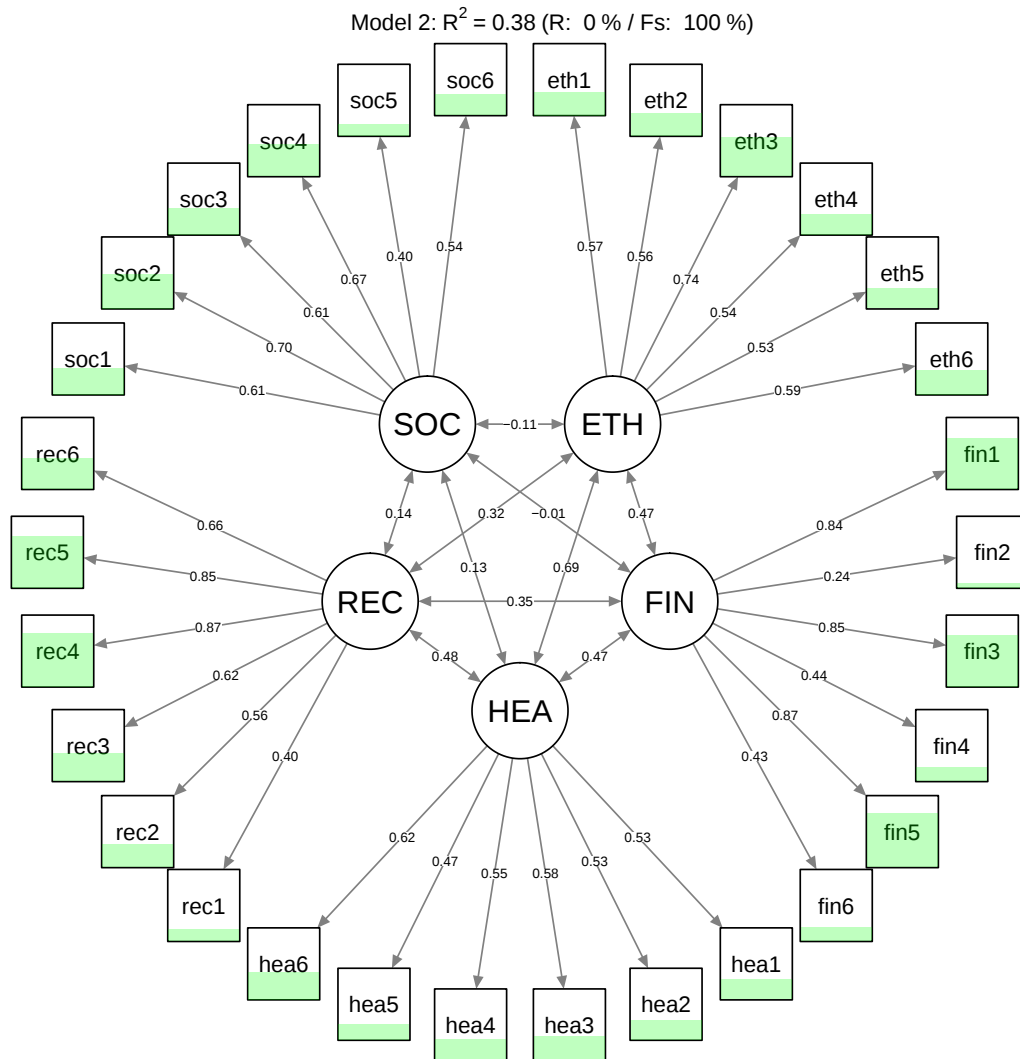


Figure S2: Depiction of the factor structure implemented in M2 (i.e., the factor structure originally proposed for the DOSPERT scale). The green shaded areas depict the proportion of variance accounted for by the domain-specific factors. Factors were permitted to be correlated.

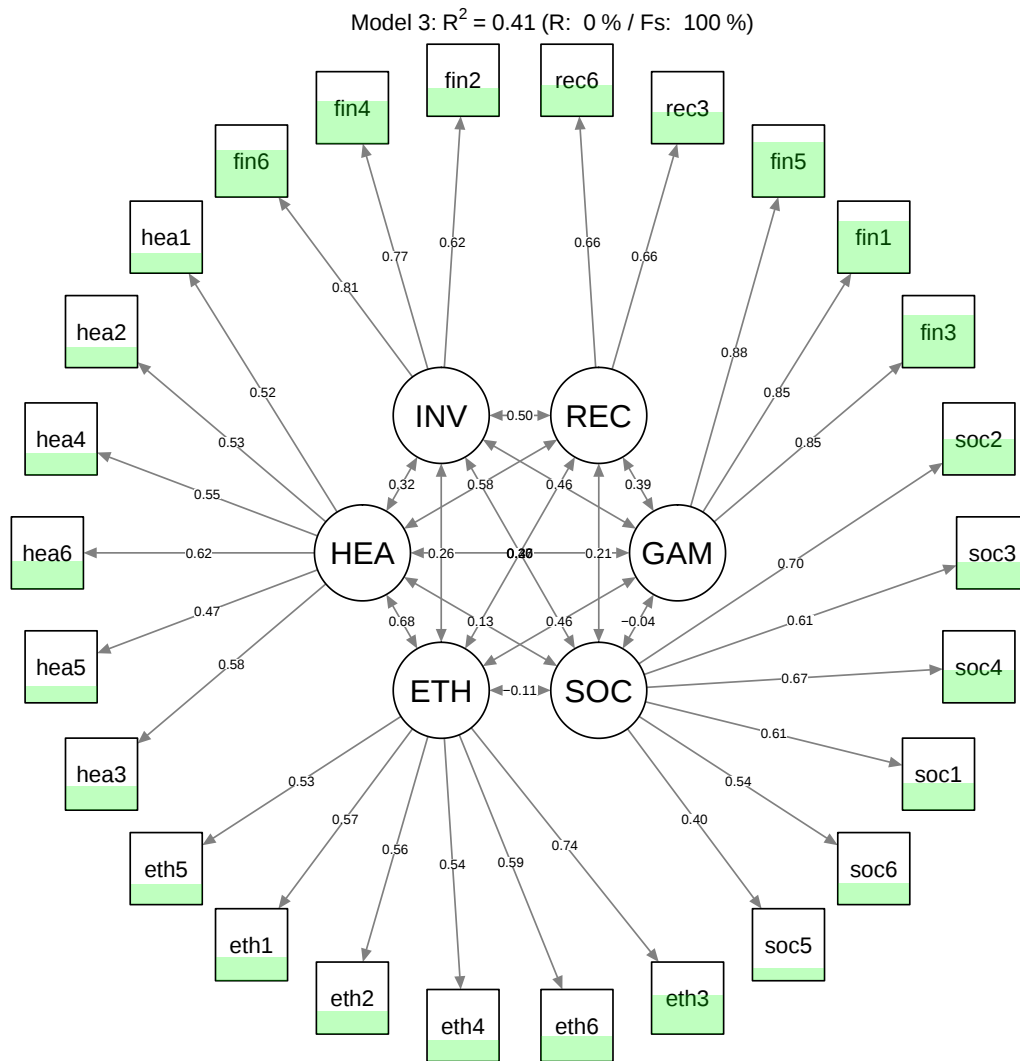


Figure S3: Depiction of the factor structure implemented in M3. The green shaded areas depict the proportion of variance accounted for by the domain-specific factors. Factors were permitted to be correlated.

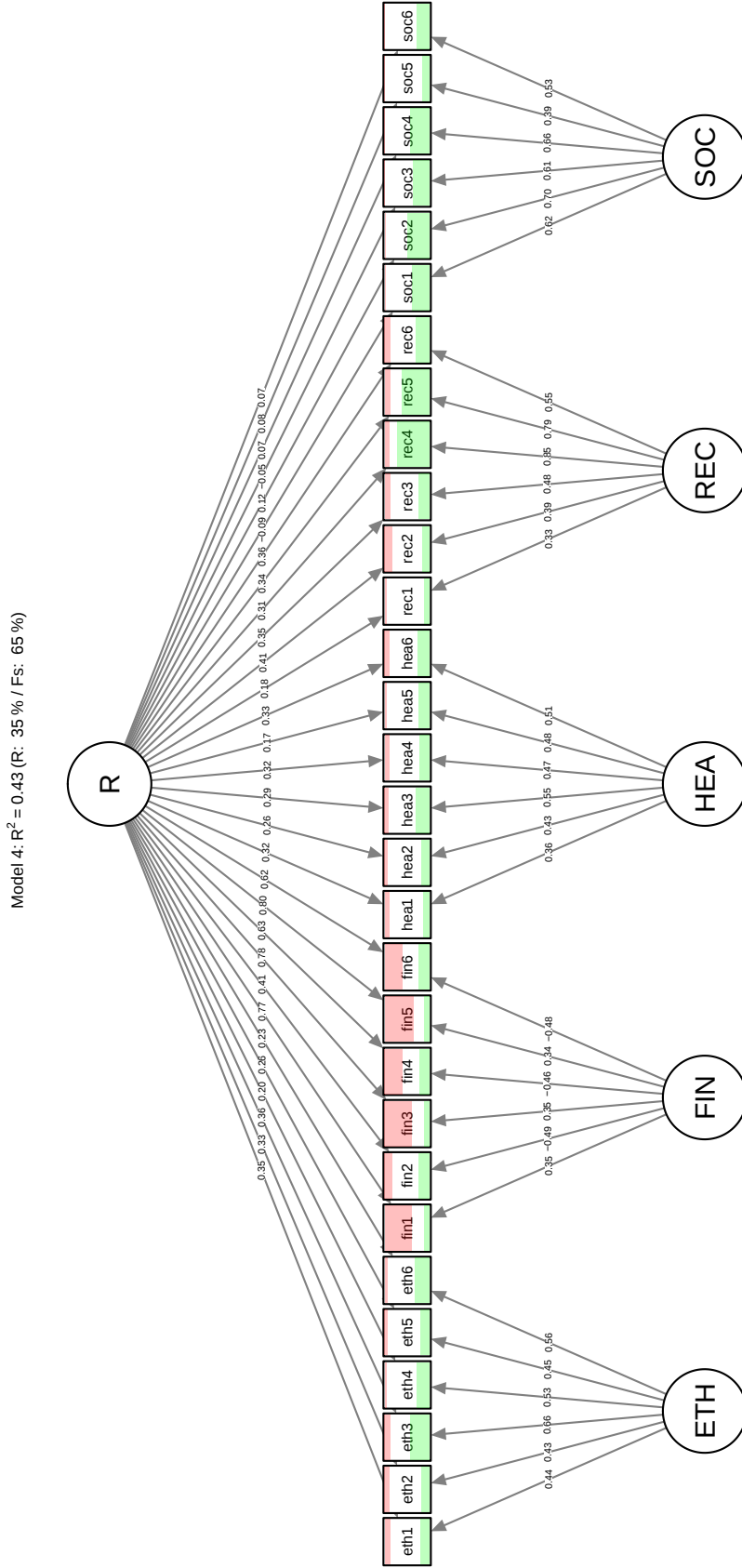


Figure S4: Depiction of the factor structure implemented in M4. The red (green) shaded areas depict the proportion of variance accounted for by the general (domain-specific) factors. Factors were constrained to be orthogonal.

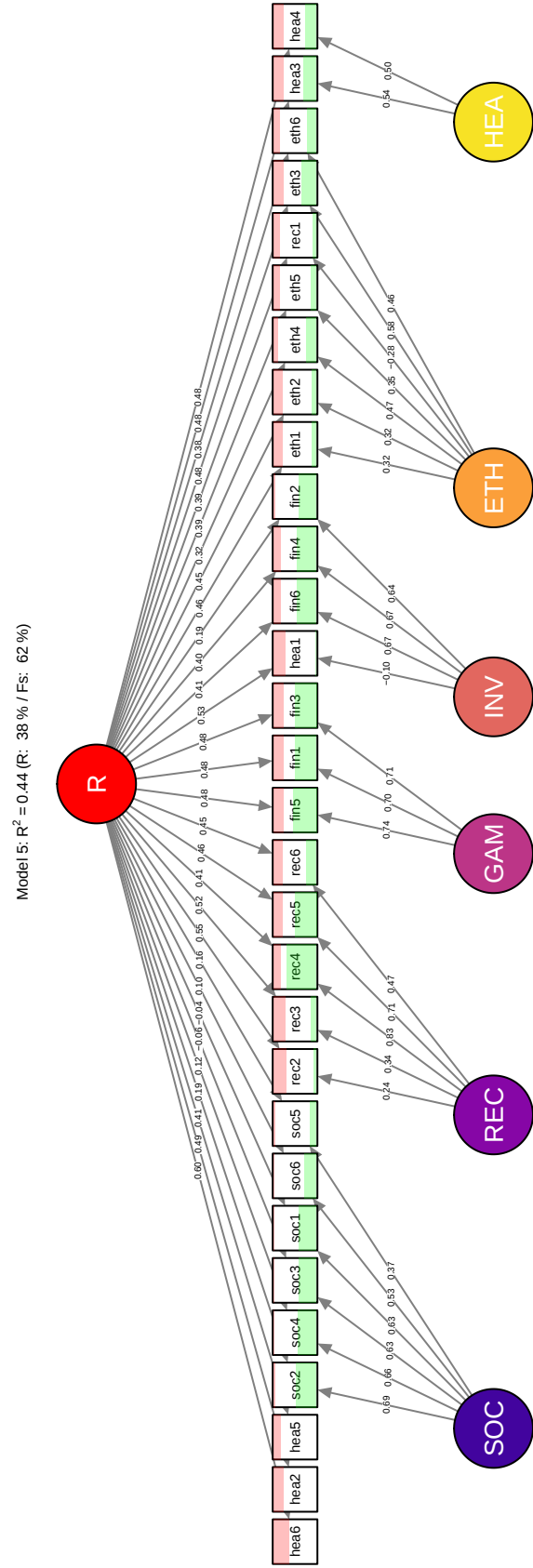


Figure S5: Depiction of the factor structure implemented in M5. The red (green) shaded areas depict the proportion of variance accounted for by the general (domain-specific) factors. Factors were constrained to be orthogonal.

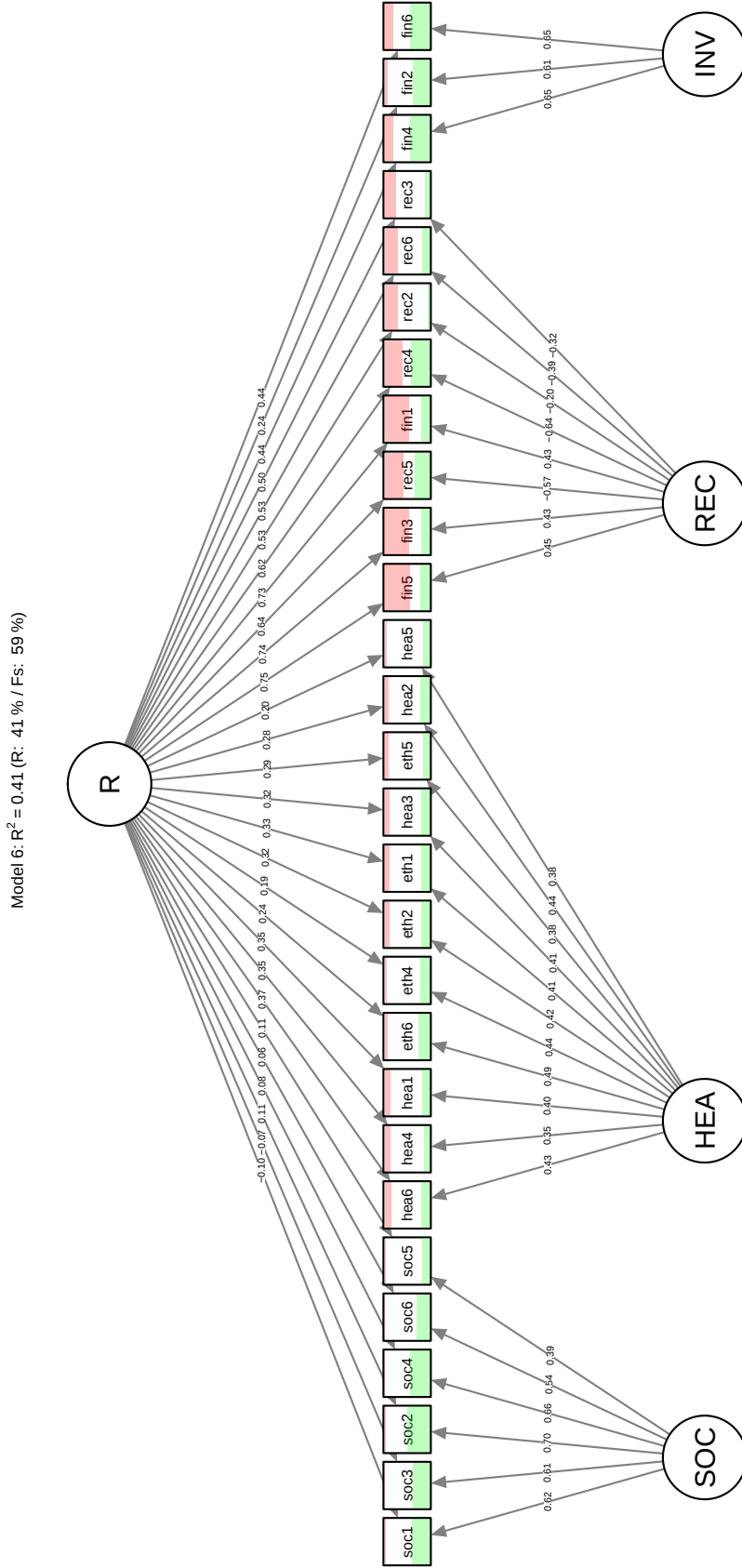


Figure S6: Depiction of the factor structure implemented in M5. The red (green) shaded areas depict the proportion of variance accounted for by the general (domain-specific) factors. Factors were constrained to be orthogonal.

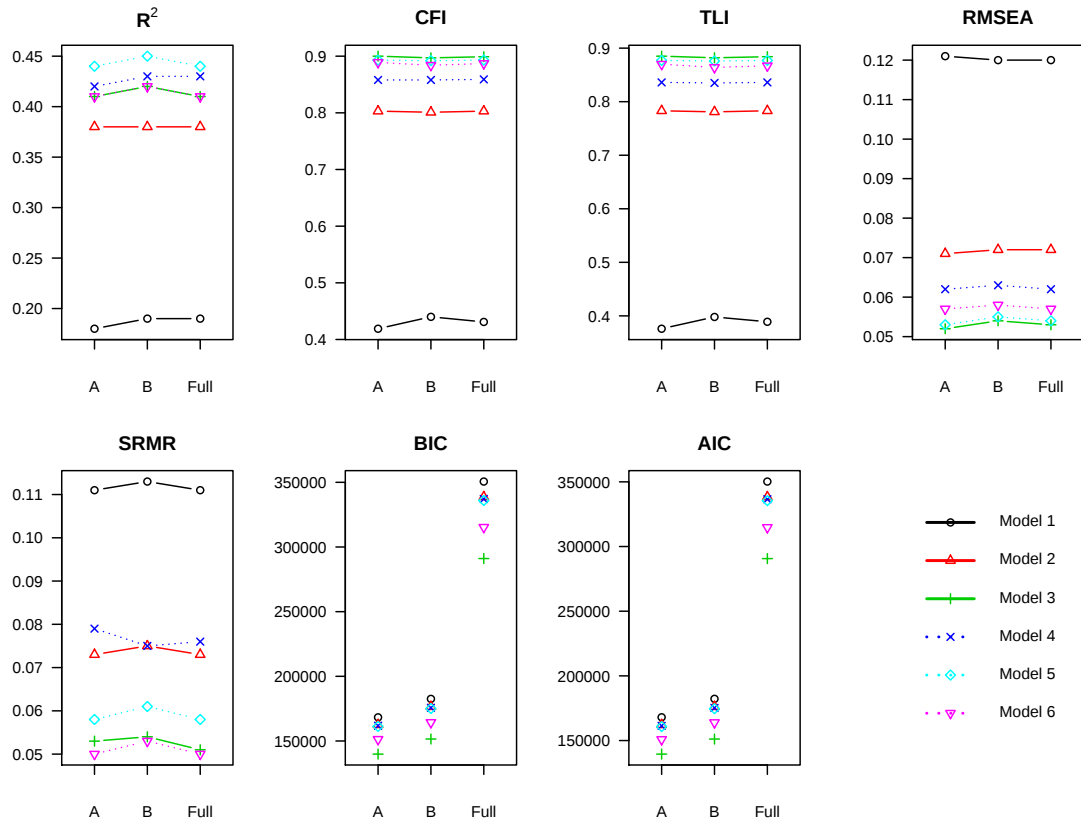


Figure S7: Model fits obtained from the confirmatory factor analyses. Fit indices are displayed separately for subsample A, B, and the full dataset. R^2 = proportion of explained variance, RMSEA = root mean square error of approximation, SRMR = standardized root mean square residual, CFI = comparative fit index, TLI = Tucker-Lewis index, BIC = Bayesian Information Criterion, AIC = Akaike information criterion. Note that the BIC and AIC can only be used for comparisons of models within but not across the different samples, because these likelihood-based information criteria directly depend on the specific sample size.

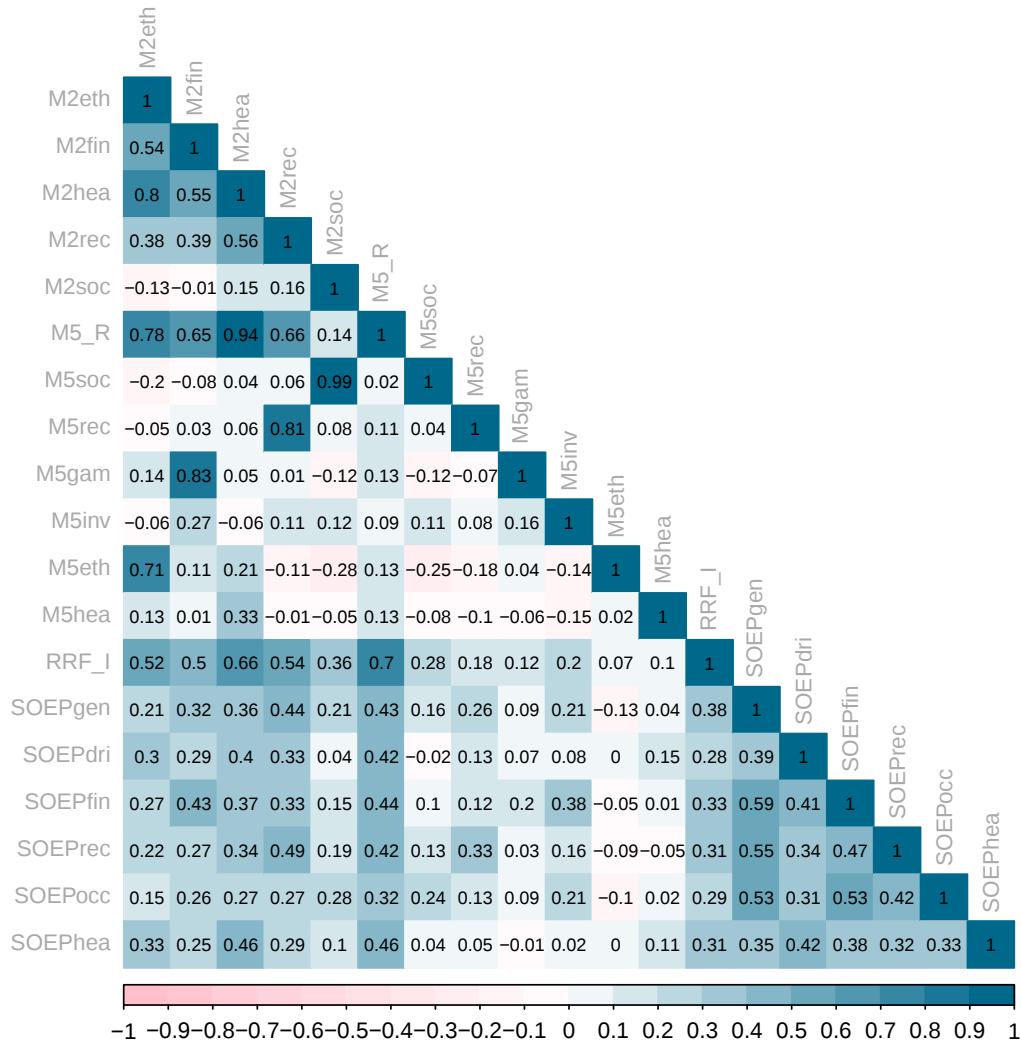


Figure S8: Correlation of factor values across models and with additional indicators. The plot shows correlations between i) the domain-specific factors extracted from M2, ii) the general and domain-specific factors extracted from M5, iii) the person-specific intercepts extracted from the risk-return framework (RRF_I), and iv) the general and domain-specific one-question items assessing risk-taking propensity as implemented in the German Socio-Economic Panel (SOEP).

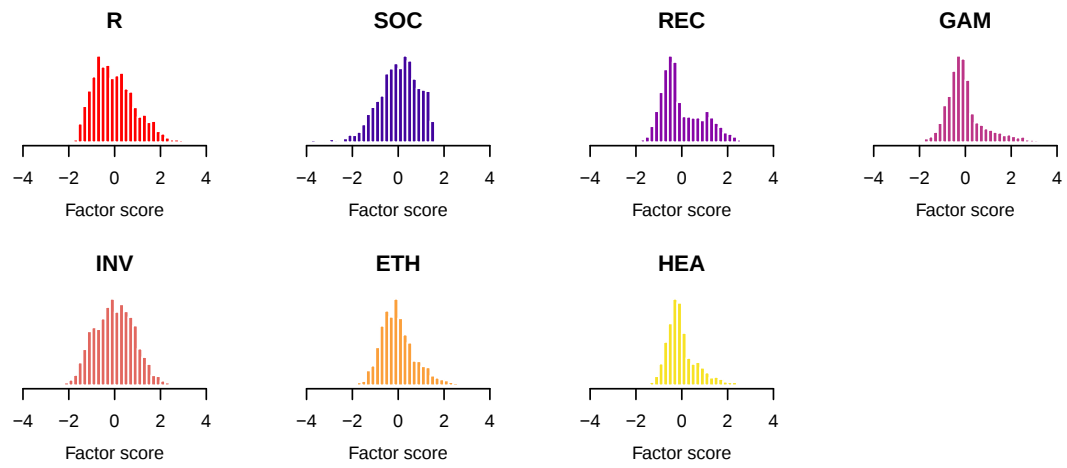


Figure S9: Distribution of factor scores extracted from M5. Each panel shows the distribution of participants' factor scores for one factor.

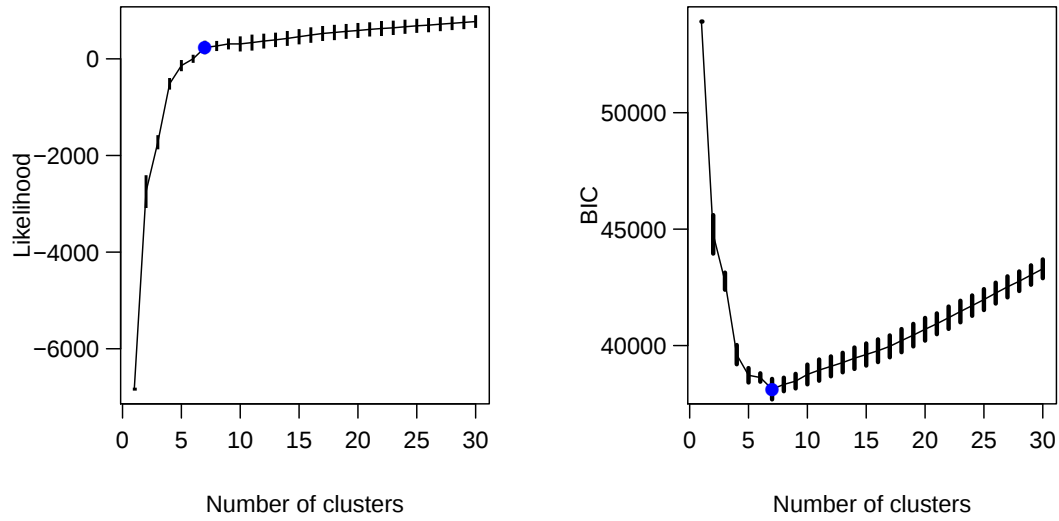


Figure S10: Screening analysis to identify the optimal number of clusters. This analysis consisted of estimating from 1 to 50 clusters in the factor space of M5, across 100 iterations with different, randomly selected starting parameters. The panels show the absolute fits (left panel) and Bayesian information criteria (BIC; right panel) for different numbers of clusters. The lines depict the average fit indices across the 100 iterations, and the confidence intervals depict $\pm$ one standard deviation. According to this analysis, the best number of clusters was 7 (marked with blue circles).

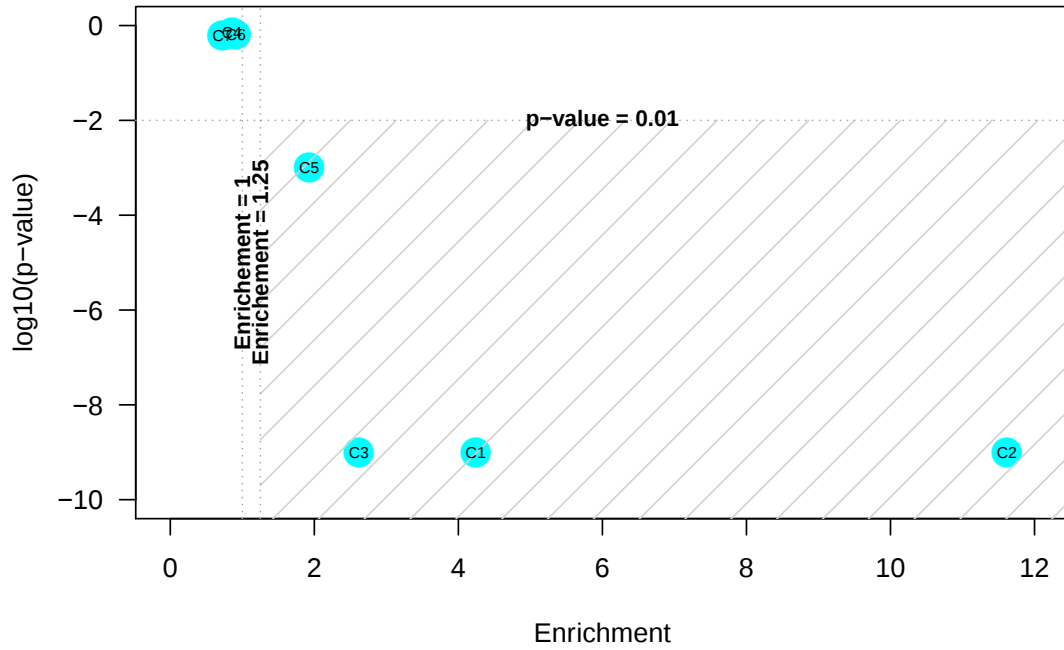


Figure S11: Cluster enrichment analysis (see Gerlach et al., 2018). The plot shows the enrichment and p-values for the seven clusters obtained from the initial BIC model selection (see Fig. S10). Only clusters 1, 2, 3, and 5 fulfilled the two imposed criteria (hence representing the four identified profiles), namely, an enrichment of at least 1.25 and a p-value of smaller than 0.01. To obtain the enrichment and p-value of each cluster, we conducted a permutation test with 1000 shuffled datasets, each implying that there exist no clusters while keeping the marginal distributions of the variables constant (i.e., the null-model). The enrichment refers to the ratio between the actual density at the cluster center and the average density across the shuffled datasets at the same point. The p-value refers to the proportion of simulation runs in which the density of the shuffled dataset was higher than the density at the cluster center (i.e., where the cluster center reflects a “false positive” and thus a spurious cluster). The clusters’ p-values (enrichments) were correlated at -0.46 (0.15) with the weights they were assigned to from the Bayesian GMMs, with the four non-spurious clusters having markedly higher weights as compared to the spurious clusters.

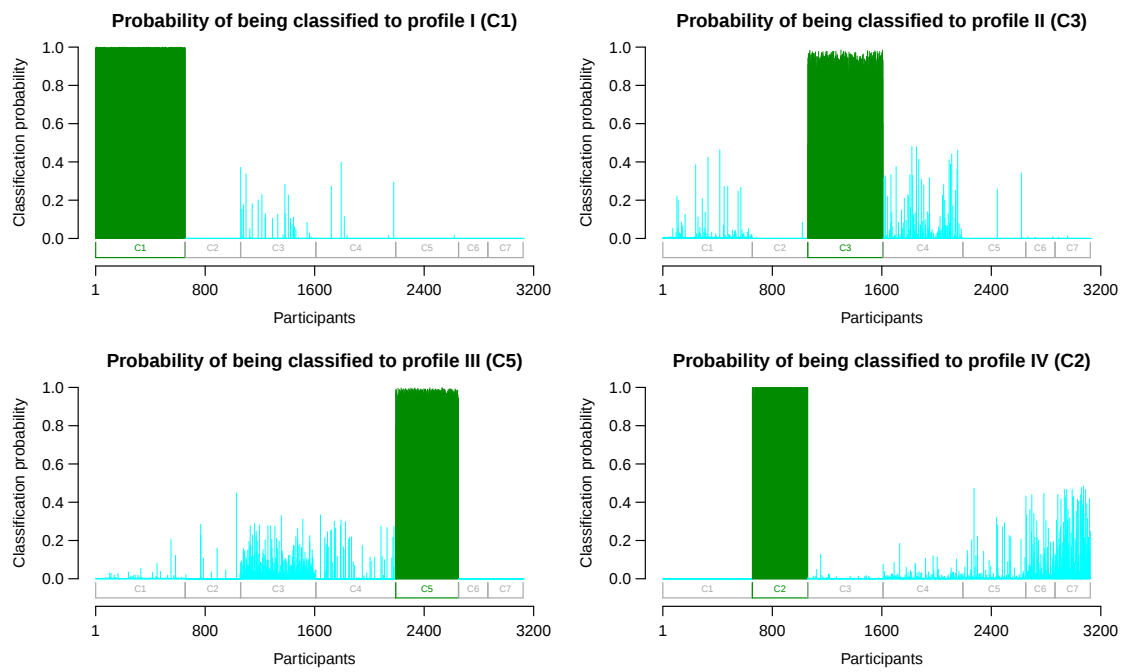


Figure S12: Classification probabilities. The figure depicts each participant's probability of being classified to one of the four profiles. Participants are grouped according to the profile they were actually classified to, and participants who were classified to the respective profile shown in the four panels are depicted in green.

Illustration of overfitting in traditional implementations of GMMs

The analysis below illustrates how a traditional implementation of Gaussian mixture models (GMM) may overfit the empirical structure even in simple two-dimensional cases. In comparison, a Bayesian GMM as implemented with the Python package *scikit-learn* (Pedregosa et al., 2011) proved to recover the correct (number of) clusters.

Step 1: Generate 2d-dataset with two obvious clusters.

```
set.seed(123)

m_x1 <- 20
m_x2 <- 25
m_y1 <- 50
m_y2 <- 80
means <- data.frame(x=c(m_x1, m_x2), y=c(m_y1, m_y2))

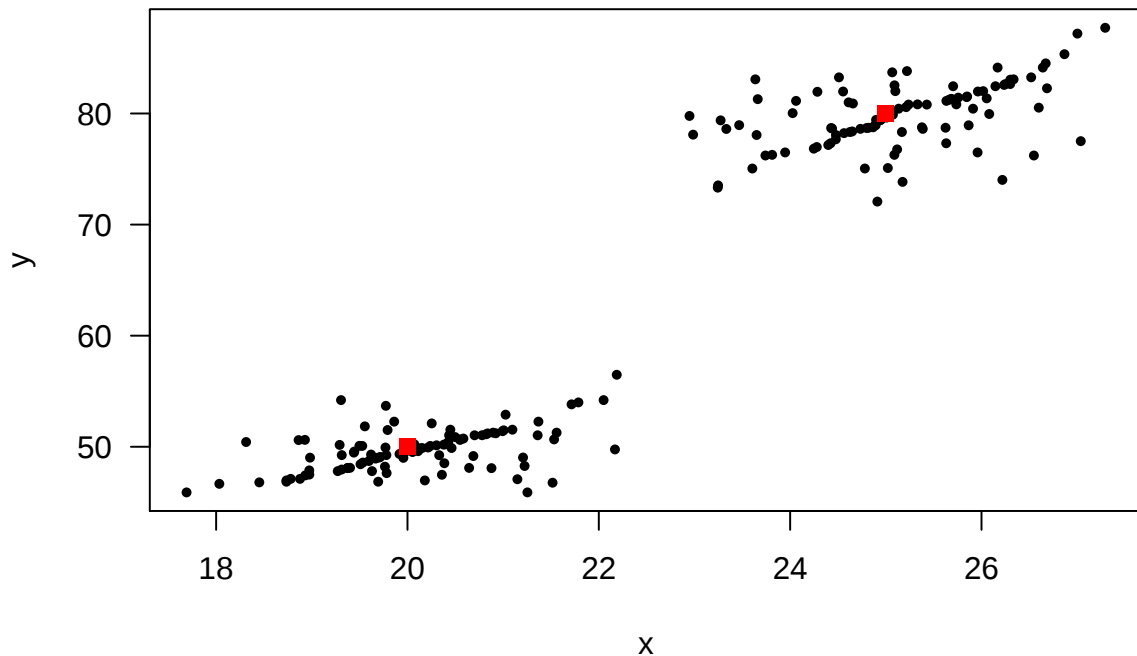
x1 <- sort(rnorm(100, mean=m_x1, sd=1))
y1 <- sort(rnorm(100, mean=m_y1, sd=2))
ind1 <- sample(1:100, 50)
ind2 <- sample(1:100, 50)
y1[ind1] <- y1[ind2]

x2 <- sort(rnorm(100, mean=m_x2, sd=1))
y2 <- sort(rnorm(100, mean=m_y2, sd=3))
ind3 <- sample(1:100, 50)
ind4 <- sample(1:100, 50)
y2[ind3] <- y2[ind4]

x <- c(x1, x2)
y <- c(y1, y2)
d <- cbind(x,y)
write.csv(d, file="simdat.csv")
plot(d, las=1, pch=16, cex=.7, main="Dataset with cluster means")
```

```
points(means, pch=15, cex=1.2, col="red")
```

Dataset with cluster means



Step 2: Use Mclust to fit three different GMMs.

```
library(mclust)
```

```
## Package 'mclust' version 5.4.3
```

```
## Type 'citation("mclust")' for citing this R package in publications.
```

```
m1 <- Mclust(d, G=1:10)
```

```
m2 <- Mclust(d, G=1:10, modelNames=c("EEE"))
```

```
m3 <- Mclust(d, G=1:10, modelNames=c("EEV"))
```

```
# save the cluster means
```

```
cmeans_m1 <- t(summary(m1)$mean)
```

```
cmeans_m2 <- t(summary(m2)$mean)
```

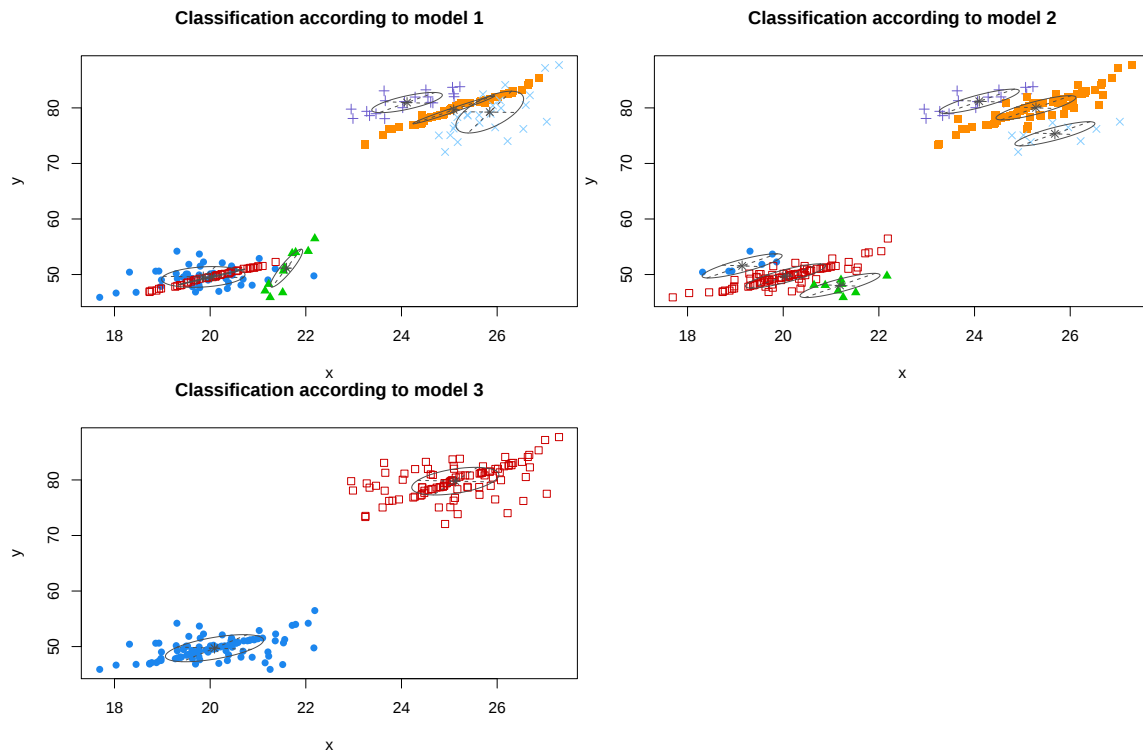
```
cmeans_m3 <- t(summary(m3)$mean)
```

```
for (i in 1:3) {
```

```
  plot(get(paste("m", i, sep="")), what="classification")
```

```
  title(paste("Classification according to model", i))
```

```
}
```



For models 1 and 2, the analysis above shows that several cluster means are very close to each other (small euclidean distance between cluster means).

Step 3: Use a Gaussian kernel to estimate the empirical density in the dataset.

```
library(ks)
dens <- kde(d, binned=F)

# save the points where the density was evaluated (will be needed below)
evalp <- expand.grid(dens$eval.points[[1]], dens$eval.points[[2]])
colnames(evalp) <- c("x", "y")

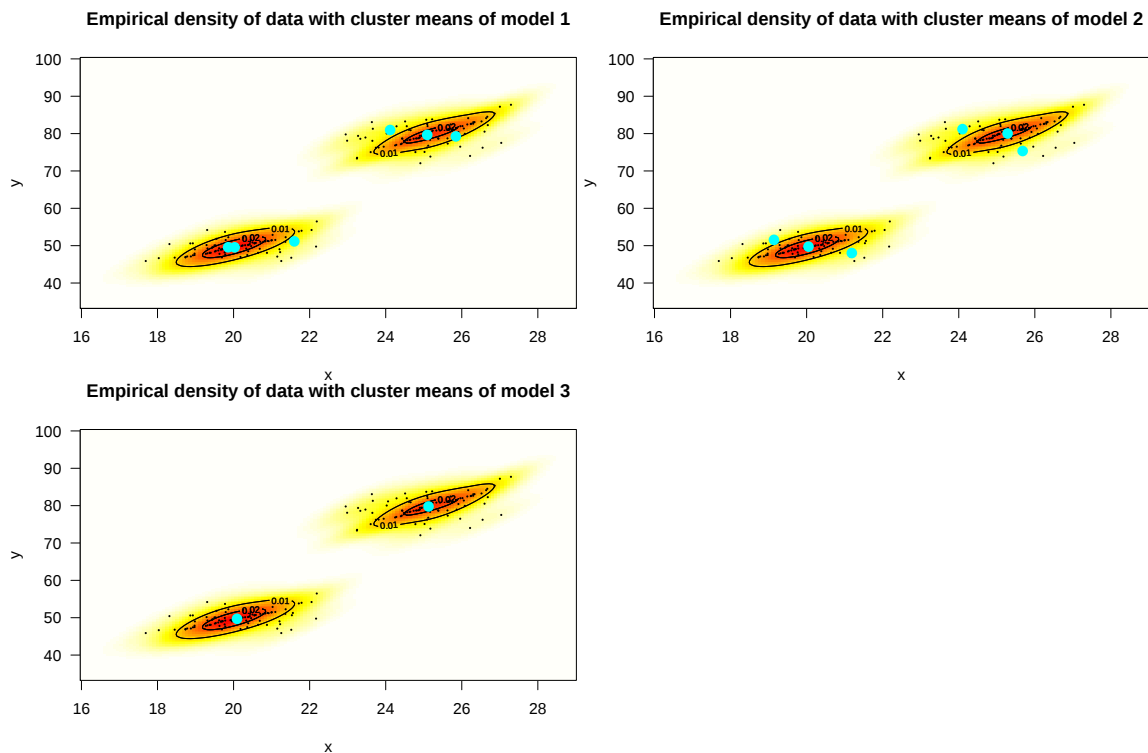
# predict the density at the cluster means
dens_cmeans_m1 <- predict(dens, x=cmeans_m1)
dens_cmeans_m2 <- predict(dens, x=cmeans_m2)
dens_cmeans_m3 <- predict(dens, x=cmeans_m3)

# plot the empirical density and add the cluster means
for (i in 1:3) {
  plot(dens, display="image", las=1,
```

```

    main=paste("Empirical density of data with cluster means of model", i))
  points(d, pch=16, cex=0.3, col="black")
  myconts <- round(seq(0.0001, max(dens$estimate), length.out=8), 2)
  plot(dens, col="black", add=T, abs.cont=myconts)
  points(get(paste("cmeans_m", i, sep="")), col="cyan", pch=16, cex=1.5)
}

```



Step 4: Generate a null-model to estimate the “enrichment” of the identified clusters: specifically, shuffle the dataset 100 times (by column, that is, keeping the marginal distributions of each variable constant, but removing any clusters).

```

null_m1 <- NULL
null_m2 <- NULL
null_m3 <- NULL
null_model_z <- NULL
for (i in 1:100) {

  # shuffle dataset
  d_shf <- apply(d, 2, sample)

```

```

# estimate density
dens_shf <- kde(d_shf, binned=F)

# get (and store) the density at the cluster centers
shf_cmeans_m1 <- predict(dens_shf, x=cmeans_m1)
null_m1 <- rbind(null_m1, shf_cmeans_m1)

shf_cmeans_m2 <- predict(dens_shf, x=cmeans_m2)
null_m2 <- rbind(null_m2, shf_cmeans_m2)

shf_cmeans_m3 <- predict(dens_shf, x=cmeans_m3)
null_m3 <- rbind(null_m3, shf_cmeans_m3)

# get (and store) the density at each actual datapoint
shf_z <- predict(dens_shf, x=evalp)
null_model_z <- rbind(null_model_z, shf_z)
}

```

Step 5: Determine a “non-parametric” p-value for each cluster, that is, the proportion of simulation runs in which the density at the cluster centers were LARGER (under the null-model) than the density of the cluster centers in the actual dataset

```

pvals_m1 <- rowMeans(apply(null_m1, 1, function(x) {x > dens_cmeans_m1}))
pvals_m2 <- rowMeans(apply(null_m2, 1, function(x) {x > dens_cmeans_m2}))
pvals_m3 <- rowMeans(apply(null_m3, 1, function(x) {x > dens_cmeans_m3}))
print(pvals_m1)

```

```
## [1] 0.00 0.00 0.00 0.07 0.00 0.00
```

```
print(pvals_m2)
```

```
## [1] 0.80 0.00 0.75 0.12 0.00 0.96
```

```
print(pvals_m3)
```

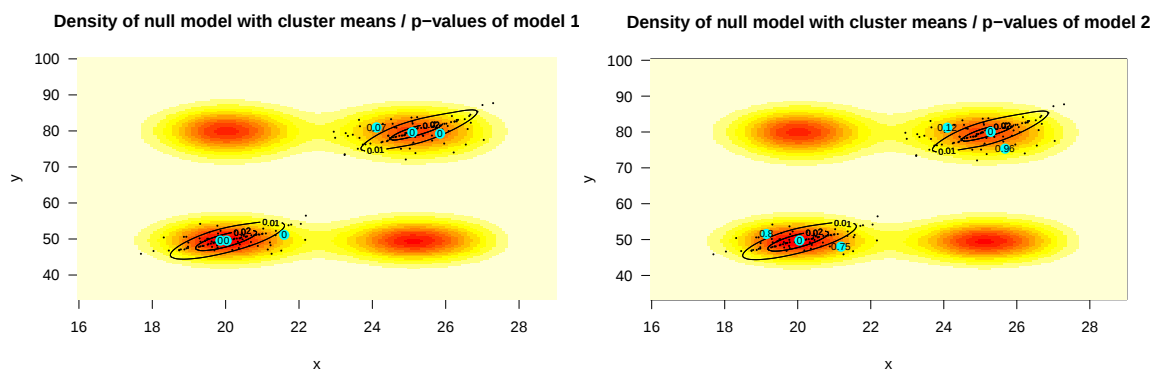
```
## [1] 0 0
```

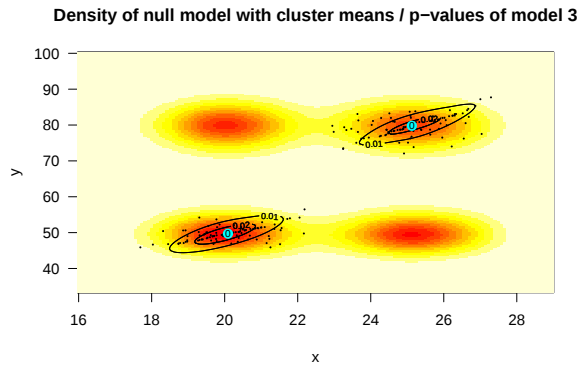
```

# get the average density across all the shuffled datasets (i.e., the null-model)
tmp <- colMeans(null_model_z)
xs <- unique(evalp$x)
ys <- unique(evalp$y)
out <- as.data.frame(matrix(NA, nrow=length(xs), ncol=length(ys)))
rownames(out) <- xs
colnames(out) <- ys
for (k in 1:length(tmp)) {
  out[as.character(evalp$x[k]), as.character(evalp$y[k])] <- tmp[k]
}

# plot the average density add p-values to each of the cluster means
for (i in 1:3) {
  image(x=xs, y=ys, z=1-as.matrix(out), las=1, xlab="x", ylab="y",
        main=paste("Density of null model with cluster means / p-values of model", i))
  points(d, pch=16, cex=0.3, col="black")
  plot(dens, col="black", add=T, abs.cont=myconts)
  points(get(paste("cmeans_m", i, sep="")), col="cyan", pch=16, cex=1.5)
  text(round(get(paste("pvals_m", i, sep="")), 2),
        x=get(paste("cmeans_m", i, sep=""))[,1],
        y=get(paste("cmeans_m", i, sep=""))[,2], cex=.7)
}

```





Interim conclusions:

For model 1 this saturation analysis does not recognize that there exist several spurious clusters, because most of them happen to have means right at the peaks of the empirical density.

For model 2, this analysis works because some of the spurious clusters have means slightly outside the peaks of the empirical density, that is, where the density under the null-model peaks “higher” sufficiently often.

For model 3, the two cluster means are at the peaks of the empirical density, thus correctly resulting in very small p-values.

We next estimate a variational Bayesian gaussian mixture model, as obtained from using the Python-package scikit-learn (the following code estimates the key steps in Python):

```
import sklearn.mixture as sklm
gmm = sklm.BayesianGaussianMixture(n_components = 10,
                                     n_init = 5,
                                     covariance_type = 'full')

c_means = gmm.means_
c_cov = gmm.covariances_
c_weights = gmm.weights_
likelihood = gmm.lower_bound_
```

The following code reads in and visualizes the identified clusters in R:

```
plot(d, las=1, pch=16, cex=.7, main="Dataset with Bayesian GMM.")
points(means, pch=15, cex=1.2, col="red")

# read values from BayesianGaussianMixture (scikit-learn)
means <- read.csv("../python/gmm/out/means.csv", header=F)
```

```

weights <- unlist(read.csv("../python/gmm/out/weights.csv", header=F))
print(round(weights, 4))

##      V11      V12      V13      V14      V15      V16      V17      V18      V19      V110
## 0.5022 0.4973 0.0004 0.0000 0.0000 0.0000 0.0000 0.0000 0.0000 0.0000

# add constant to visualize all clusters (even the ones we really small weights)
weights <- weights + .04
weights <- weights / max(weights) * .2

if (T) {
  means[3:10,1] <- means[3:10,1] + runif(8, min=-.5, max=.5)
  means[3:10,2] <- means[3:10,2] + runif(8, min=-.5, max=.5)
}

for (i in 1:nrow(means)) {
  covars <- read.csv(paste("../python/gmm/out/covars_c", i-1, ".csv", sep=""),
                    header=F)

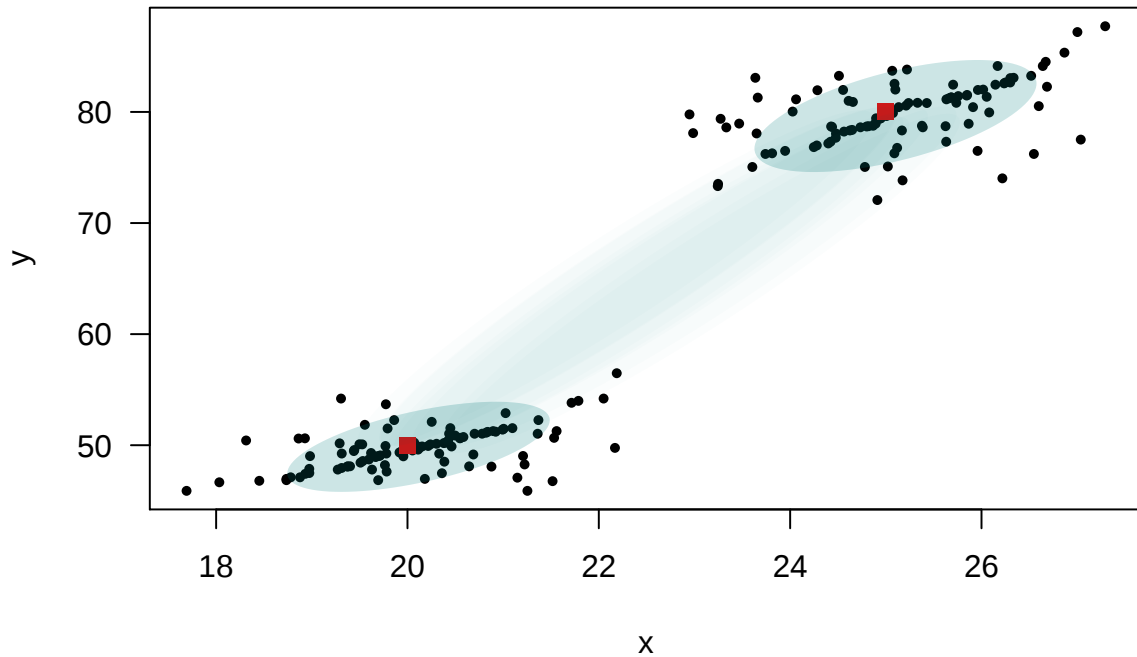
  m <- as.numeric(means[i,])
  c <- matrix(rev(unlist(covars)), nrow=nrow(covars))

  e <- eigen(c)
  eig_vecs <- e$vectors
  unit_eig_vec <- eig_vecs[,2] #!/ norm(eig_vecs)
  angle <- atan2(unit_eig_vec[2], unit_eig_vec[1])
  angle <- (180 * angle / pi) * - 1
  eig_vals <- sqrt(2) * sqrt(e$values)

  library(plotrix)
  draw.ellipse(x=m[1], y=m[2], a=eig_vals[1], b=eig_vals[2], angle=angle,
              deg=T, lwd=0.05, col=rgb(0,.5,.5, alpha = weights[i]), border=0)
}

```

Dataset with Bayesian GMM.



In conclusion, even though the variational Bayesian implementation did fit up to 10 clusters, permitting covariances to fully vary, only the two correct clusters were recovered, each of which being assigned weights close to .5 (and all remaining eight clusters were assigned weights virtually equal to 0).

References

- Caspi, A., Houts, R. M., Belsky, D. W., Goldman-Mellor, S. J., Harrington, H., Israel, S., Meier, M. H., Ramrakha, S., Shalev, I., Poulton, R., & Moffitt, T. E. (2014). The p factor: One general psychopathology factor in the structure of psychiatric disorders? *Clinical Psychological Science*, 2(2), 119–137. <https://doi.org/10.1177/2167702613497473>
- Deary, I. J. (2012). Intelligence. *Annual Review of Psychology*, 63(1), 453–482. <https://doi.org/10.1146/annurev-psych-120710-100353>
- Frey, R., Pedroni, A., Mata, R., Rieskamp, J., & Hertwig, R. (2017). Risk preference shares the psychometric structure of major psychological traits. *Science Advances*, 3, e1701381. <https://doi.org/10.1126/sciadv.1701381>
<http://advances.sciencemag.org/content/advances/3/10/e1701381.full.pdf>
- Gerlach, M., Farb, B., Revelle, W., & Amaral, L. A. N. (2018). A robust data-driven approach identifies four personality types across four large data sets. *Nature Human Behaviour*, 2(10), 735–742. <https://doi.org/10.1038/s41562-018-0419-z>

- Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., Blondel, M., Prettenhofer, P., Weiss, R., Dubourg, V., Vanderplas, J., Passos, A., Cournapeau, D., Brucher, M., Perrot, M., & Duchesnay, É. (2011). Scikit-learn: Machine learning in python. *Journal of Machine Learning Research*, 12, 2825–2830. Retrieved June 12, 2019, from <http://www.jmlr.org/papers/v12/pedregosa11a.html>
- Spearman, C. (1904). "General intelligence," objectively determined and measured. *The American Journal of Psychology*, 15(2), 201–292. <https://doi.org/10.2307/1412107>
- Stan Development Team. (2016). *Rstanarm: Bayesian applied regression modeling via Stan*. <https://mc-stan.org>
- Weber, E. U., Blais, A. R., & Betz, N. E. (2002). A domain-specific risk-attitude scale: Measuring risk perceptions and risk behaviors. *Journal of Behavioral Decision Making*, 15(4), 263–290. <https://doi.org/10.1002/bdm.414>
- Weber, E. U., & Milliman, R. A. (1997). Perceived risk attitudes: Relating risk perception to risky choice. *Management Science*, 43(2), 123–144. <https://doi.org/10.1287/mnsc.43.2.123>