

AI based portable Hospital

Team Members:

1. Abdul Etibaroghlu---Zhejiang Normal University
2. Sidik Mohamed Hassan---Zhejiang university of science and technology
3. Angel---Tsinghua University

Team Members

- Abdul Etibaroghlu
Azerbaijan,
Zhejiang Normal University,
Master of Biology

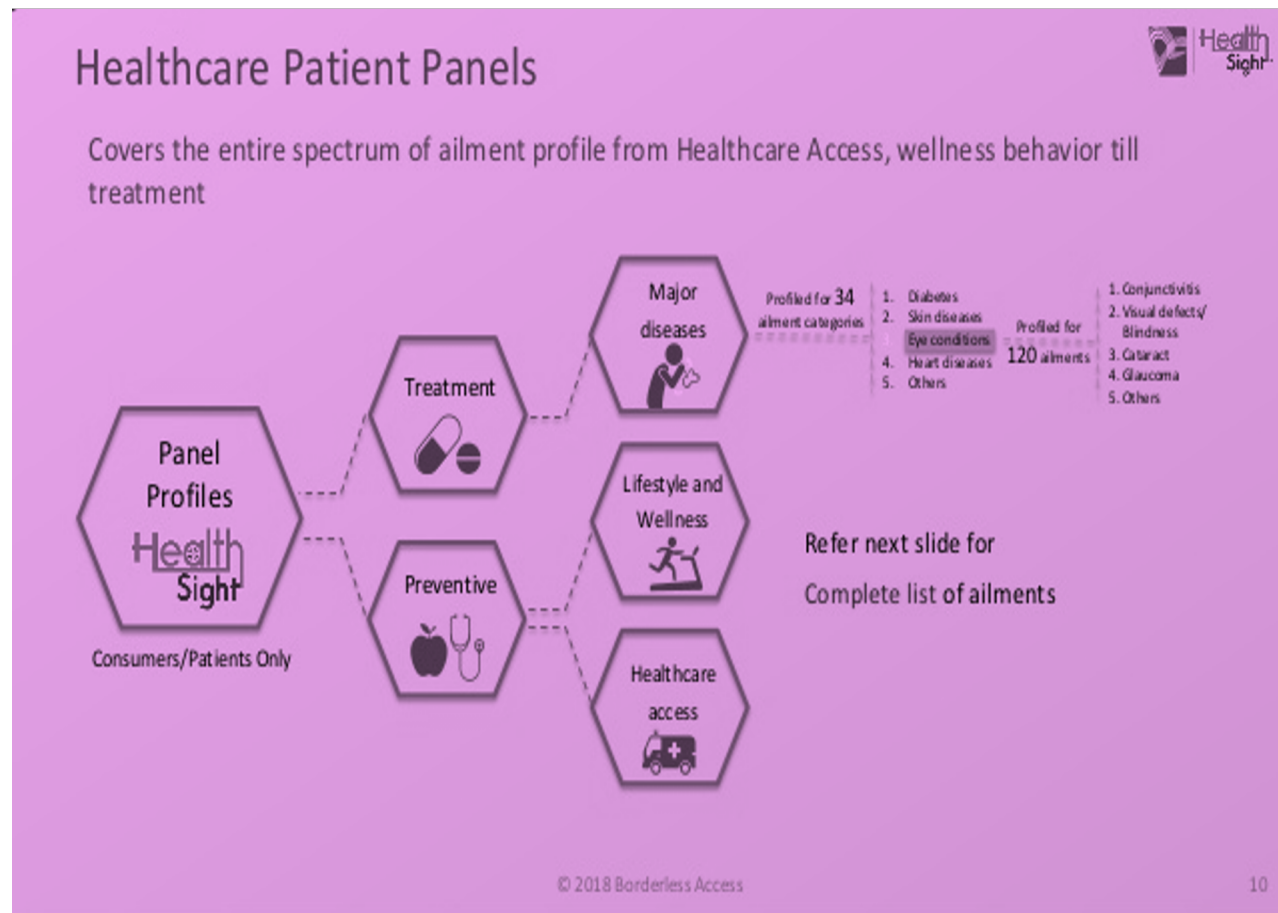


What we designed?

Portable Hospital, let AI decides:

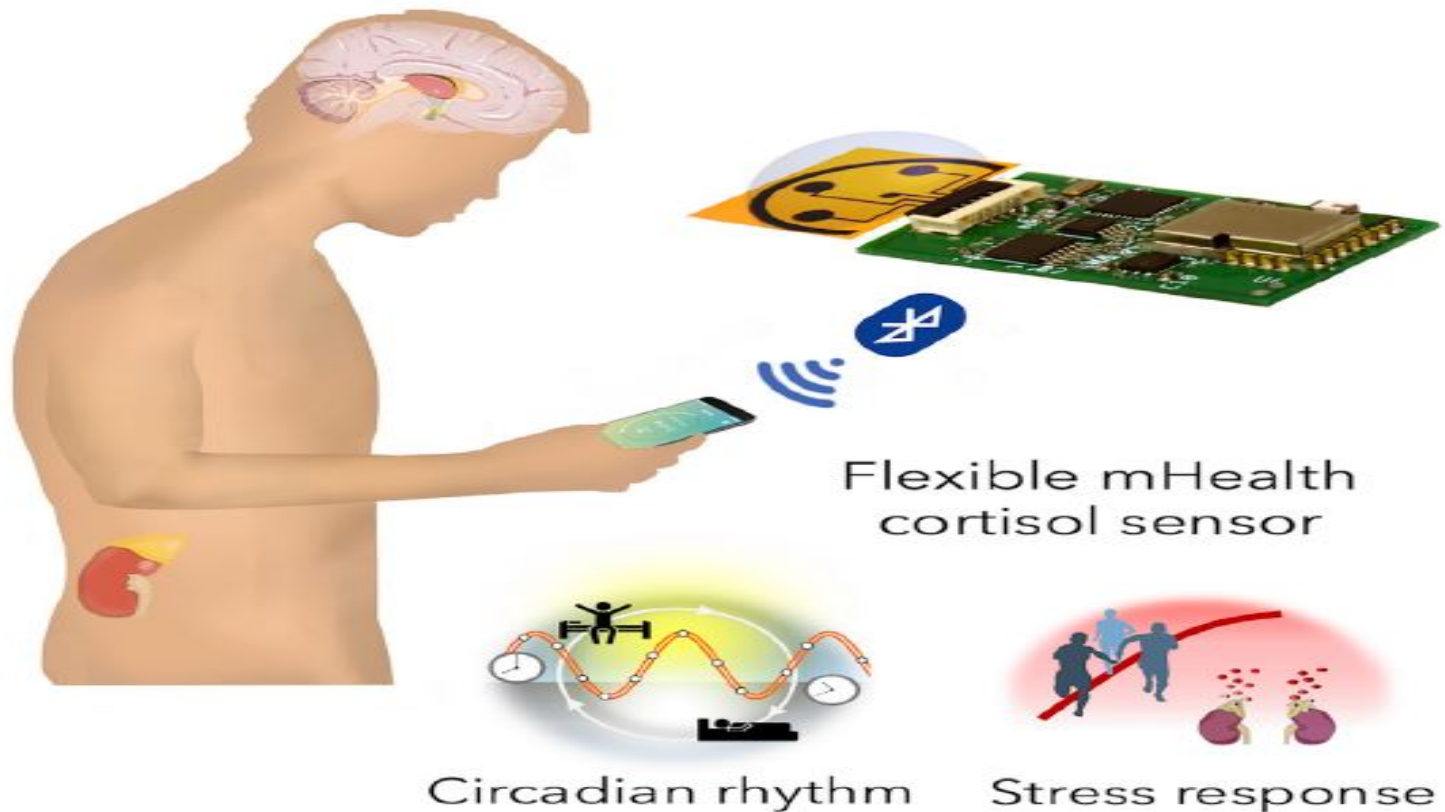
Decides based on
your daily routine
and health status
history:

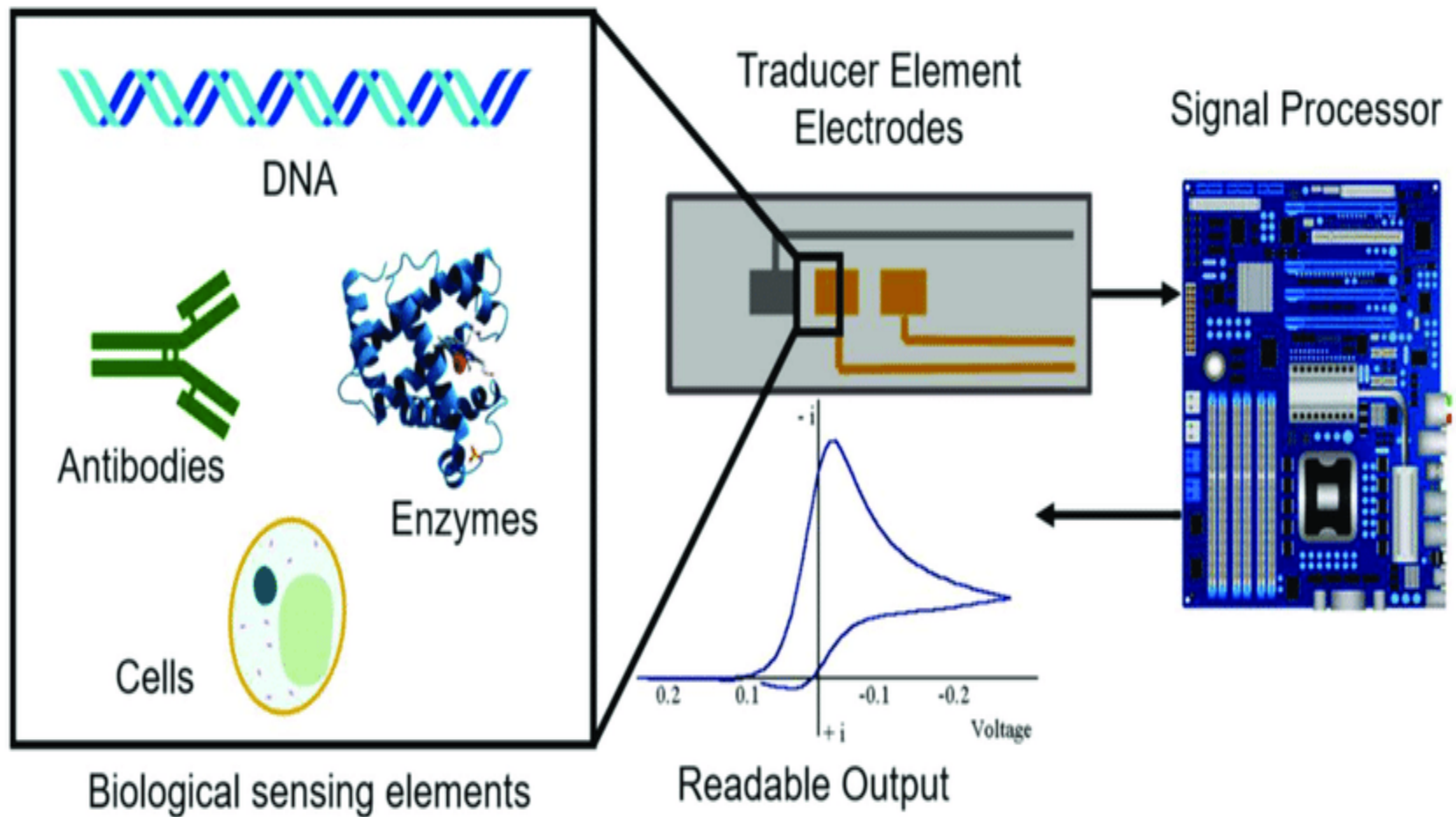
Tracks: blood
pressure,
melatonin level in
blood, testosterone
level, sweat
analysis,



How it works?

Electrochemical Biosensors





Team member

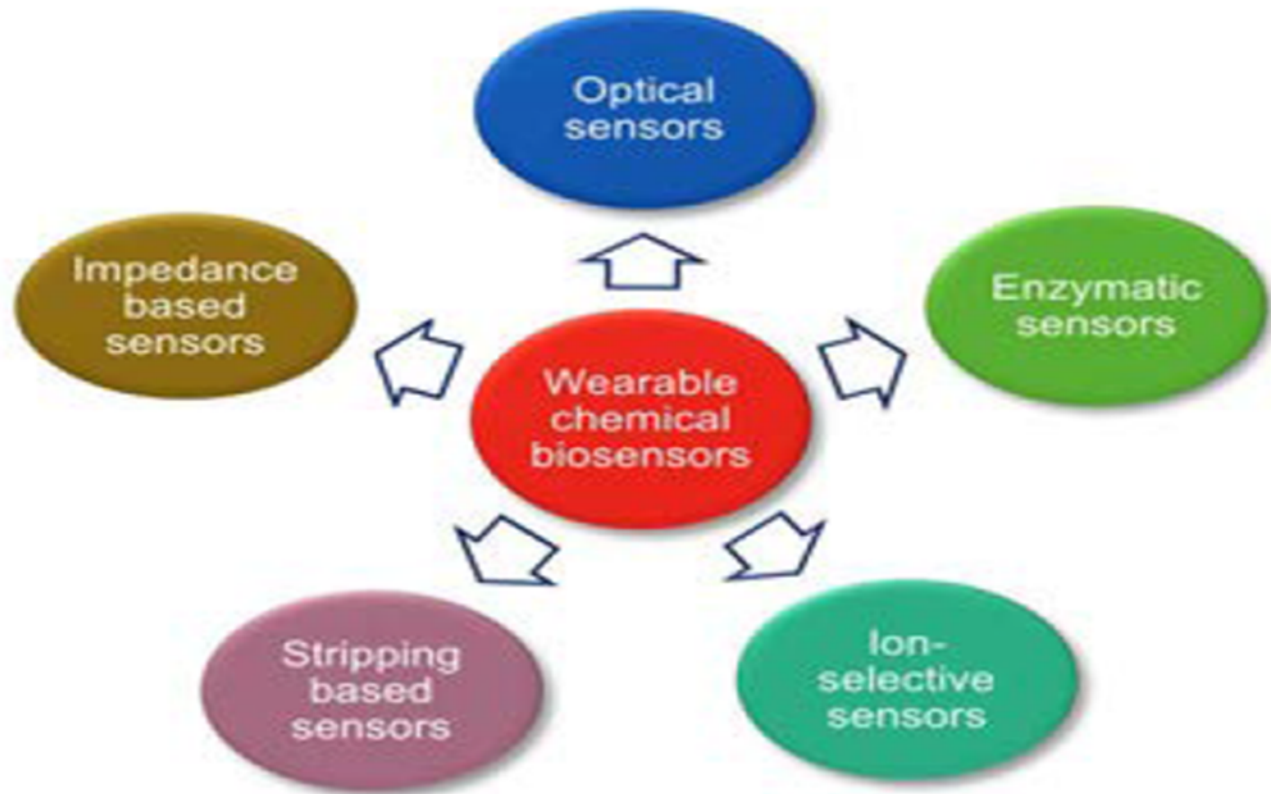
- Name: Sidik Mohamed Hassan
- School : Zhejiang university of science and technology
- Master of chemical engineering



Sweat biosensor

- Human sweat, a very important bio fluid, consists of water (99%), ions (Na^+ , K^+ , Ca^{2+} , etc.), metabolites (glucose, lactate, ethanol, etc.), hormones, small proteins and peptides, providing a wealth of chemical information about physiological and metabolic status.

DETECTION TECHNIQUES

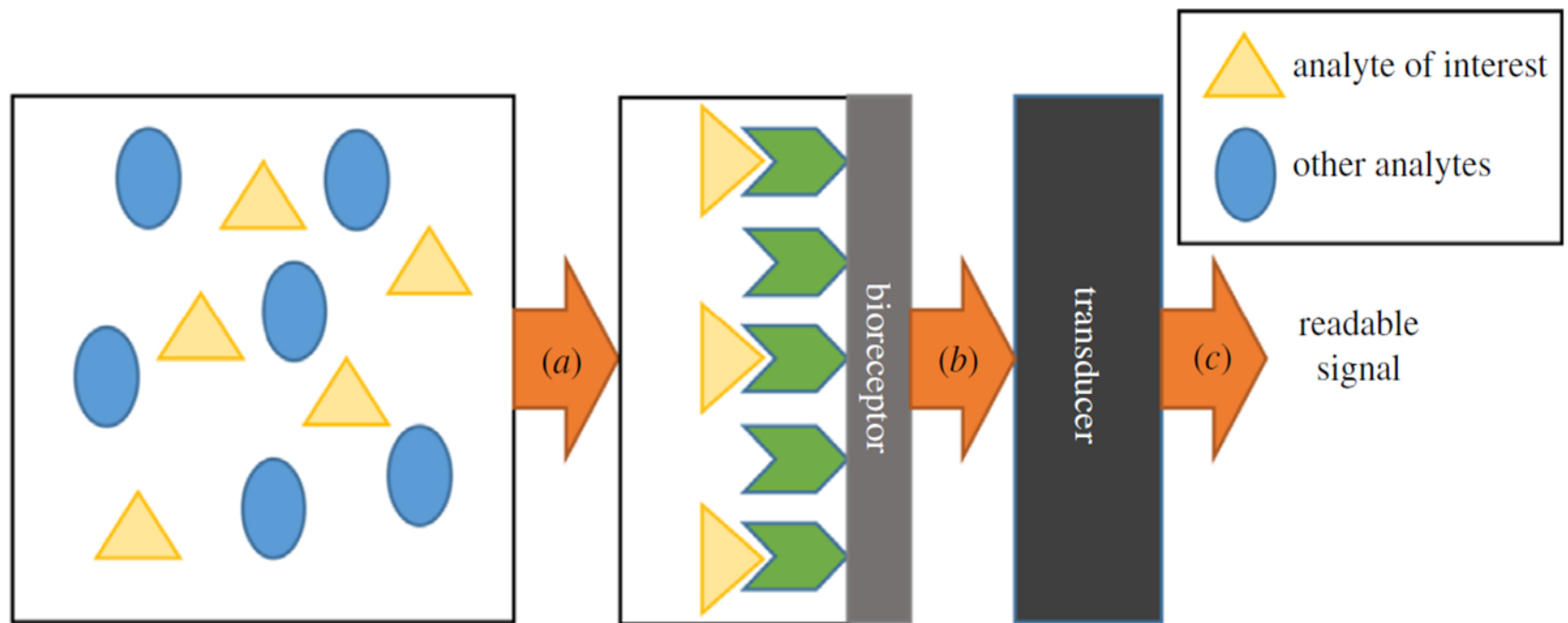


Types of biosensor and the analytes they can detect in sweat

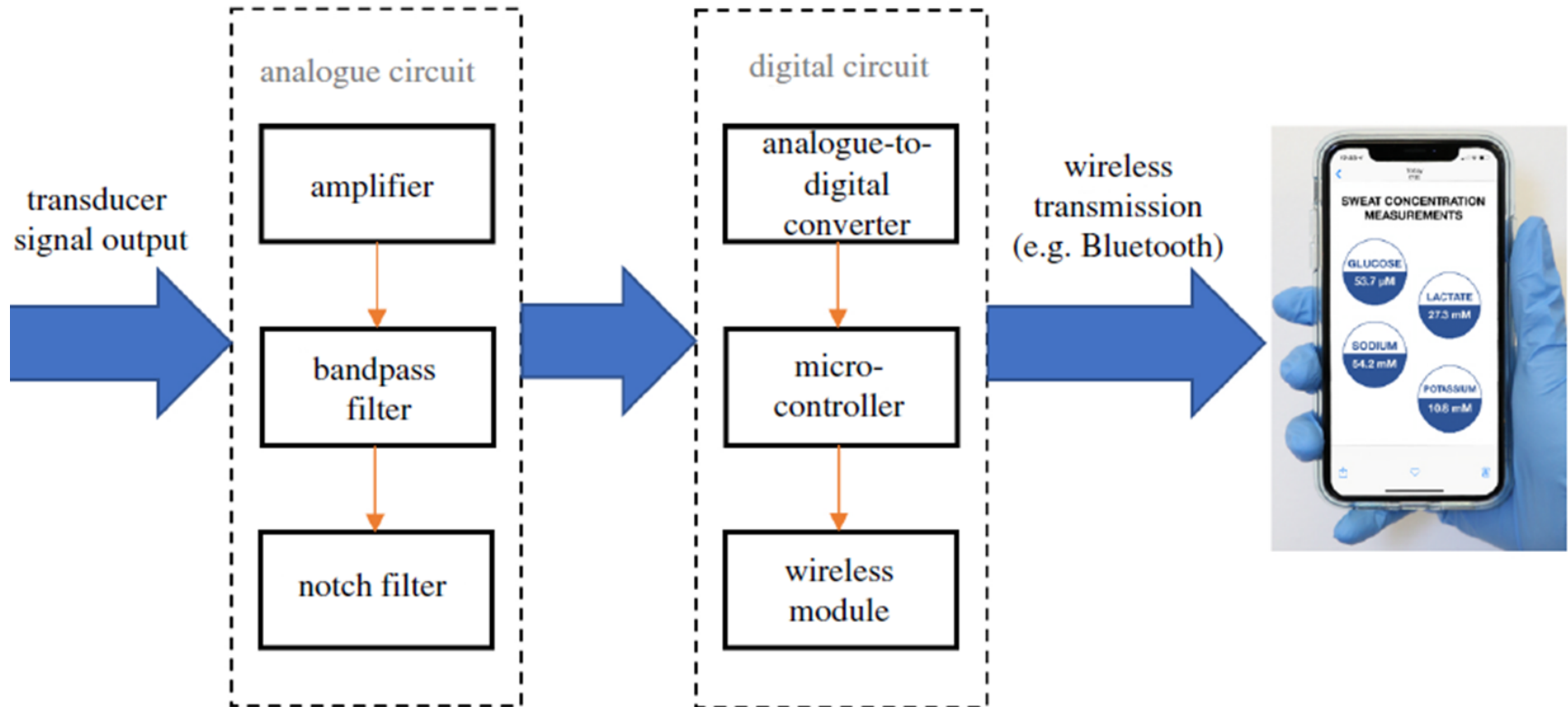
sensor	analyte/data of interest
optical	sweat rate pH level OH^- , H^+ , Cu^+ , Fe^{2+}
impedance-based	sweat rate sweat conductivity
ion-selective electrodes (ISEs)	pH level Na^+ , H^+ , K^+ , NH_4^+ , Cl^- , Mg^{2+} , Zn^{2+} , Ca^{2+}
enzymatic amperometric	metabolites (glucose, lactate, ethanol, uric acid)
stripping-based	heavy metals (Cu, Zn, Pb, Cd, Hg)

The most common methods of detection are enzymatic amperometric and potentiometric ion-selective electrode (ISE) sensors .

How To Design Sweat Biosensor



How does it work



Applications

cystic fibrosis

Chloride

hypoglycaemia

glucose

Ischaemia

lactate

Hypokalaemia

Potassium

Hyponatremia

sodium

Team member

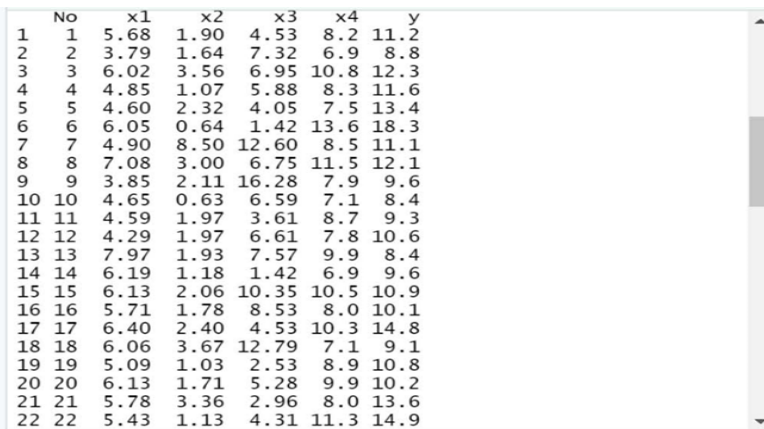
- Name: Angel
- School: Tsinghua university
- Master degree of Mathematics Statistics



Data analysis of blood sugar data

- Brief statement: This data is about 27 diabetes patients' health conditions. Serum total cholesterol(x_1), glycerine(x_2), fasting insulin(x_3), glycated haemoglobin(x_4), fasting blood sugar(y). We need to analysis how x_1, x_2, x_3, x_4 impact y

- Calculation:
- `> mydata<-read.table("clipboard",header=T)`
- `> mydata`



	No	x1	x2	x3	x4	y
1	1	5.68	1.90	4.53	8.2	11.2
2	2	3.79	1.64	7.32	6.9	8.8
3	3	6.02	3.56	6.95	10.8	12.3
4	4	4.85	1.07	5.88	8.3	11.6
5	5	4.60	2.32	4.05	7.5	13.4
6	6	6.05	0.64	1.42	13.6	18.3
7	7	4.90	8.50	12.60	8.5	11.1
8	8	7.08	3.00	6.75	11.5	12.1
9	9	3.85	2.11	16.28	7.9	9.6
10	10	4.65	0.63	6.59	7.1	8.4
11	11	4.59	1.97	3.61	8.7	9.3
12	12	4.29	1.97	6.61	7.8	10.6
13	13	7.97	1.93	7.57	9.9	8.4
14	14	6.19	1.18	1.42	6.9	9.6
15	15	6.13	2.06	10.35	10.5	10.9
16	16	5.71	1.78	8.53	8.0	10.1
17	17	6.40	2.40	4.53	10.3	14.8
18	18	6.06	3.67	12.79	7.1	9.1
19	19	5.09	1.03	2.53	8.9	10.8
20	20	6.13	1.71	5.28	9.9	10.2
21	21	5.78	3.36	2.96	8.0	13.6
22	22	5.43	1.13	4.31	11.3	14.9

23	23	6.50	6.21	3.47	12.3	16.0
24	24	7.98	7.92	3.37	9.8	13.2
25	25	11.54	10.89	1.20	10.5	20.0
26	26	5.84	0.92	8.61	6.4	13.3
27	27	3.84	1.20	6.45	9.6	10.4

- `> lm.health<-lm(y~x1+x2+x3+x4,data=mydata)`
linear regression of y for the variables
x1,x2,x3,x4
- `> summary(lm.health)` ## read the result of
the linear regression

```

Call:
lm(formula = y ~ x1 + x2 + x3 + x4, data = mydata)

Residuals:
    Min       1Q   Median       3Q      Max
-3.6268 -1.2004 -0.2276  1.5389  4.4467

Coefficients:
              Estimate Std. Error t value Pr(>|t|)
(Intercept)   5.9433     2.8286   2.101  0.0473 *
x1             0.1424     0.3657   0.390  0.7006
x2             0.3515     0.2042   1.721  0.0993 .
x3            -0.2706     0.1214  -2.229  0.0363 *
x4             0.6382     0.2433   2.623  0.0155 *
---
Signif. codes:
  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

Residual standard error: 2.01 on 22 degrees of freedom
Multiple R-squared:  0.6008,    Adjusted R-squared:  0.5282

F-statistic: 8.278 on 4 and 22 DF,  p-value: 0.0003121

```

Conclusion

- As the analysis shows, the p-value of x_3 and x_4 is smaller than 0.05, and the p-value of x_1 and x_2 is bigger than 0.05, so fasting insulin(x_3) and glycated haemoglobin(x_4) have the bigger impact on fasting blood sugar(y), serum total cholesterol(x_1) and glycerine(x_2) have the smaller impact. That is to say, fasting insulin(x_3) and glycated haemoglobin(x_4) are two main influence factors of fasting blood sugar(y).

Use the cluster to cluster the patients

- If the patients are in the same category, then they have the similar health conditions.
- Calculation:
- `d<-
dist(mydata,method="euclidean",diag=T,upper=F,p
=2)`
- `KM<-
kmeans(mydata,4,nstart=20,algorithm="Hartigan-
Wong") ## Use K-means to cluster 4 category.`
- `KM`

```
-/ 
K-means clustering with 4 clusters of sizes 10, 3, 7, 7
```

```
Cluster means:
```

	No	x1	x2	x3	x4
1	12.500000	5.434000	2.580000	8.635000	8.240000
2	24.000000	8.673333	8.340000	2.680000	10.866667
3	21.714286	5.501429	1.678571	4.952857	9.200000
4	4.142857	5.438571	2.018571	5.271429	9.542857

	y
1	9.71000
2	16.40000
3	12.57143
4	12.52857

```
Clustering vector:
```

```
[1] 4 4 4 4 4 4 1 4 1 1 1 1 1 1 1 1 3 1 3 3 3 2 2 2 3
[27] 3
```

```
Within cluster sum of squares by cluster:
```

```
[1] 362.48049 56.61493 155.88397 162.13403
(between_SS / total_SS = 70.9 %)
```

```
Available components:
```

```
[1] "cluster" "centers" "totss"
[4] "withinss" "tot.withinss" "betweenss"
[7] "size" "iter" "ifault"
```

```
> |
```

Conclusion of the cluster

- As the result show, if we cluster these 27 diabates patients into 4 categories, then these four categories have 10, 3, 7, 7 patients, respectively.
- First category: Number 7, 9~16, 18 (10 patients)
- Second category: Number 23~25 (3 patients)
- Third category: Number 17, 19~22, 26~27 (7 patients)
- Forth category: Number 1~6, 8 (7 patients)