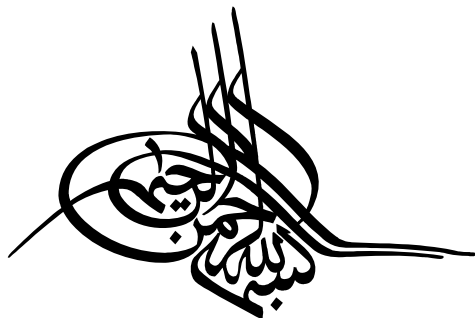




شیمی تجزیه پیشرفته

مجموعه شیمی تجزیه

گردآورنده: دکتر مجید حاجی مسینی

عضو هیأت علمی پژوهشکده علوم و فنون هسته‌ای

آمادگی آزمون دکتری

حاج حسینی، مجید
شیمی تجزیه پیشرفته رشته شیمی تجزیه / دکتر مجید حاجی حسینی (عضو هیأت علمی
پژوهشکده علوم و فنون هسته‌ای)
مشاوران صعود ماهان: ۱۴۰۴
۱۹۴ ص: جدول، نمودار (آمادگی آزمون دکتری مجموعه شیمی تجزیه)
ISBN/N: 978-600-457- 662-7

فارسی - چاپ اول
شیمی تجزیه پیشرفته
دکتر مجید حاجی حسینی
ج - عنوان
رده بندی کنگره:
رده بندی دیویی
کتابخانه ملی ایران

LB ۲۳۵۳ / ح ۲ / ۱۳۹۲

۳۷۸/۱۶۶۴

۳۳۳۴۸۸۴



انتشارات مشاوران صعود ماهان



- ☐ نام کتاب: شیمی تجزیه پیشرفته
- ☐ مدیران مسئول: هادی و مجید سیاری
- ☐ گردآورنده: دکتر مجید حاجی حسینی
- ☐ مدیر برنامه ریزی و تولید محتوا: سمیه بیگی
- ☐ ناشر: مشاوران صعود ماهان
- ☐ نوبت و تاریخ چاپ: چاپ اول / ۱۴۰۴
- ☐ تیراژ: ۱۰۰۰ نسخه
- ☐ قیمت: ۴۳۰۰/۰۰۰ ریال
- ☐ شابک: ISBN ۹۷۸-۶۰۰-۴۵۸-۶۶۲-۷

انتشارات مشاوران صعود ماهان: تهران - خیابان ولیعصر، بالاتر از تقاطع ولیعصر مطهری، پلاک ۲۰۵۰

تلفن: ۸۸۱۰۱۱۳ و ۸۸۴۰۱۳۱۳

کلیه حقوق مادی و معنوی این اثر متعلق به موسسه آموزش عالی آزاد ماهان می‌باشد. و هرگونه اقتباس و
کپی برداری از این اثر بدون اخذ مجوز پیگرد قانونی دارد.

بنام خدا

ایمان داریم که هر تغییر و تحول بزرگی در مسیر زندگی بدون تحول معرفت و نگرش میسر نخواهد بود. پس بیایید با اندیشه توکل، تفکر، تلاش و تحمل در توسعه دنیای فکریمان برای نیل به آرامش و آسایش توأمان اولین گام را برداریم. چون همگی یقین داریم دانایی، توانایی می‌آورد.

شاد باشید و دلی را شاد کنید

برادران سیاری

مقدمه گردآورنده

شیمی تجزیه کلاسیک مقدمه و اساس شیمی تجزیه می باشد. بنابراین ایجاب می نماید که متخصصین شیمی با گرایش تجزیه تسلط کامل بر شیمی تجزیه پایه داشته باشند. متأسفانه توجه بسیار زیادی اساتید در دانشگاهها به بعد پژوهشی، کیفیت بعد آموزشی را به شدت پایین آورده، امید است کتاب مذکور که شامل مباحث شیمی تجزیه پایه برای آزمون ورودی به دکتری شیمی تجزیه گردآوری شده است جهت ارتقا بعد آموزشی مفید واقع گردد. در مطالب ارایه شده وجود اشتباهات اجتناب ناپذیر است در نتیجه از تمامی دانشجویان تقاضا می شود که با ارائه پیشنهادات به تکمیل این مطالب در نوبت چاپ بعدی یاری نمایند.

با تشکر

دکتر مجید حاجی حسینی

هیئت علمی پژوهشگاه علوم و فنون هسته ای

فهرست مطالب

فصل اول - آمار در شیمی تجزیه	۷
فصل دوم: تعادلات شیمیایی.....	۲۹
فصل سوم: تعادلات اسید و باز در محیط آبی	۴۰
فصل چهارم: تعادلات اسید و باز در حلال‌های ناآبی.....	۶۶
فصل پنجم: استانداردهای شیمیایی	۹۲
فصل ششم: حلالیت رسوب‌ها.....	۱۰۵
فصل هفتم: خواص رسوب	۱۱۵
فصل هشتم: کاربردهای واکنش‌های رسوبی.....	۱۳۰
فصل نهم: کمپلکسها.....	۱۴۰
فصل دهم: خواص روش‌های تجربی برای مطالعات تعادلات کمپلکس.....	۱۵۷
فصل یازدهم: پتانسیل‌های الکترودی	۱۷۵
فصل دوازدهم: اندازه‌گیری غلظت با استفاده از روش‌های سینتیکی.....	۱۸۲

فصل اول

آمار در شیمی تجزیه

آنچه که در این فصل می‌خوانیم

- ☒ مقدمه
- ☒ روشهای بیان دقت
- ☒ روشهای بیان صحت
- ☒ انواع خطاهای سیستماتیک
- ☒ نتیجه‌گیری از منحنی خطای نرمال

فصل اول

آمار در شیمی تجزیه

مقدمه

شرط استفاده از آمار این است که حداقل هر آنالیز را سه بار تکرار کنیم.

۱-۱- اهداف استفاده از آمار:

۱-۱-۱ گزارش بهترین نتیجه:

برای گزارش کردن بهترین نتیجه از چهار روش استفاده می شود:

- استفاده از میانگین حسابی (متداول ترین روش)

$$\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$$

$$\bar{X}_G = \sqrt[n]{X_1 \cdot X_2 \cdot \dots \cdot X_n}$$

- متوسط هندسی:

- میانه: باید اعداد را به ترتیب مرتب کنیم. اگر تعداد اعداد فرد باشد، عدد وسط و اگر تعداد اعداد زوج باشد میانگین دو عدد

وسط میانه نام دارد.

- مُد (شیوه): عددی که بیشترین تکرار را داشته باشد.

۱-۲- بررسی داده ها از نظر صحت و دقت:

دقت: نزدیکی و تکرارپذیری داده ها به هم

صحت: نزدیکی داده ها به مقدار حقیقی

روش های بیان دقت:

۱. رنج (گستره): تفاوت بزرگ ترین عدد از کوچک ترین عدد. هرچه رنج کم تر، دقت بیش تر

۲. متوسط انحراف از میانگین ($\bar{d}$): هرچه $\bar{d}$ کوچک تر، دقت بیش تر

$$\bar{d} = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})$$

۳. انحراف معیار ← انحراف معیار نمونه، هرچه انحراف معیار کوچکتر، دقت بیشتر
 ← انحراف معیار جمعیت

جمعیت: آنالیزهایی که بی‌نهایت بار تکرار شده باشد جمعیت نام دارد (حالت ایده‌آل)

متوسط آنالیزهایی که بی‌نهایت بار انجام شده باشند برابر با مقدار واقعی آنالیت (μ) است. $\bar{x} = \mu$

نمونه: به آنالیزهایی که به تعداد محدود تکرار می‌شوند نمونه یا Sample می‌گوییم (حالت واقعی)

$$S = \sqrt{\frac{\sum (x_i - \bar{x})^2}{n-1}} \quad \text{انحراف معیار نمونه} \quad \delta = \sqrt{\frac{\sum (x_i - \mu)^2}{n}} \quad \text{انحراف معیار جمعیت}$$

تعداد آنالیزهای تکراری انجام شده

نکته: جمعیت و انحراف معیار جمعیت حالت ایده‌آل هستند.

نکته: S و δ جمع‌پذیر نیستند.

اما در تمام مراحل تجزیه از جمله Analysis, Sample Preparation, Sampling و... باید دقت را بررسی کنیم اما چون S و δ جمع‌پذیر نیستند برای جمع کردن خطا در تمام قسمت‌ها باید از یک روش جمع‌پذیر استفاده کنیم.

۴. واریانس، که جمع‌پذیر است، هرچه واریانس V کوچکتر، دقت بیشتر.

$$V = S^2 \quad \text{or} \quad \delta^2$$

$$S_T = \sqrt{v_1 + v_2 + \dots + v_n}$$

نکته: چون واریانس، واحد اندازه‌گیری را به توان ۲ می‌رساند، خیلی متداول نیست.

۵. ضریب تغییر (C_v): هرچه C_v کوچکتر، دقت بیشتر

$$C_v = \frac{S}{\bar{x}} \quad \text{or} \quad \frac{\delta}{\mu}$$

واقعی ایده‌آل

۶. RSD (Relative Standard Deviation)، هرچه RSD کوچکتر، دقت بیشتر

$$RSD = \frac{S}{\bar{x}} \times 100 \quad \text{or} \quad \frac{\delta}{\mu} \times 100$$

نکته: متداول‌ترین روش بیان دقت RSD و بعد از آن S است.

نکته: عدد گزارش شده به‌عنوان نتیجه آزمایش فاکتور دقت $\bar{x} \pm$ است.

روش‌های بیان صحت:

۱. خطای مطلق:

$$\text{خطای مطلق} = \bar{x} - \mu$$

واحد خطای مطلق همان واحد داده مورد اندازه‌گیری است. مثبت یا منفی بودن نشان‌دهنده جهت خطاست.

۲. خطای نسبی:

$$\begin{aligned} \text{خطای نسبی} &= \frac{\bar{x} - \mu}{\mu} \times 100 \Rightarrow \% \\ &\times 1000 \Rightarrow \text{ppt} \\ &\times 10^6 \Rightarrow \text{ppm} \end{aligned}$$



خطای نسبی واحد ندارد، پس می‌توانیم خطای نسبی چند آنالیز را با هم مقایسه کنیم.

۱-۲- انواع خطا در تجزیه:

۱- **خطای فاحش (بزرگ):** اصلاً قابل قبول نیست، به راحتی قابل تغییر است و باید حذف شود. مقدار خطا بزرگ بوده و مشخص می‌باشد که مربوط به تغییرات و یا اشتباهات بزرگ است.

۲- **خطای سیستماتیک:** این خطا هم اصلاً قابل قبول نیست. این خطا جهت مثبت و منفی دارد. سه منبع خطای سیستماتیک را ایجاد می‌کند (منبع روش، منبع دستگاهی، منبع شخصی)

۳- خطای راندوم یا تصادفی

انواع خطاهای سیستماتیک

ثابت: با تغییر میزان آنالیت و نمونه تغییر نمی‌کند، مانند خطای مصرف تیترانت توسط معرف در تیتراسیون.

متغیر: مانند خطای جذب رطوبت در هنگام توزین یک نمک جاذب رطوبت.

متناسب: با تغییر میزان آنالیت و نمونه، تغییر می‌کند. مانند خطای ۲ درصدی حضور هافونیم در سنگ معدن زیرکونیوم.

خطای سیستماتیک را باید تشخیص داد، اندازه‌گیری کرد و آن را حذف نمود.

روش‌های اندازه‌گیری خطای سیستماتیک:

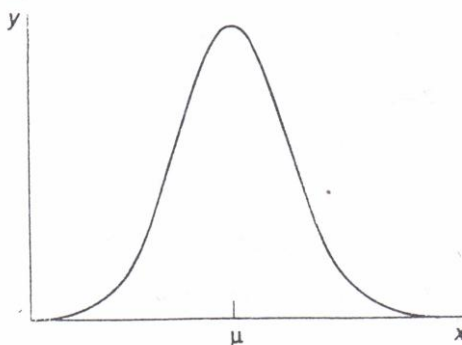
۱. استفاده از نمونه استاندارد.

۲. استفاده از نمونه شاهد (Blank). نمونه‌ای که تمام اجزا به جز آنالیت را دارد.

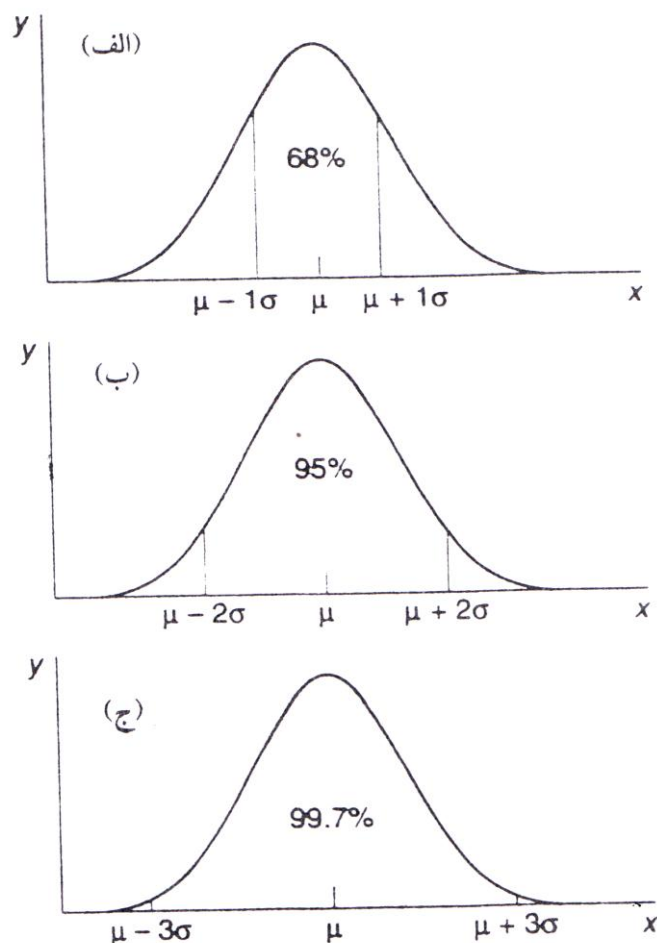
۳. استفاده از یک روش دیگر که روشی استاندارد باشد.

۴. روش افزودن مقدار مشخصی استاندارد به نمونه حقیقی آنالیز شده و محاسبه بازیابی (recovery). به این روش spike می‌گویند.

خطای تصادفی (راندوم / نامعین): این نوع خطاها منبع و علت مشخصی ندارند و عواملی باعث ایجاد این خطا می‌شوند که در کنترل ما نیستند، پس نه قابل اندازه‌گیری و نه قابل حذف شدن هستند پس تنها می‌توانیم مقدار خطای راندوم را تخمین بزنیم. برای این کار از منحنی خطای نرمال استفاده می‌کنیم.



نکته: خطای سیستماتیک روی صحت تأثیر می‌گذارد اما خطای رندوم روی دقت تأثیر می‌گذارد. چون خطای سیستماتیک جهت دارد و یا مثبت است یا منفی اما خطای رندوم در هردو جهت ایجاد می‌شود.



نتیجه‌گیری از منحنی خطای نرمال:

احتمال وقوع خطای صفر، از تمام خطاهای مثبت یا منفی بیش‌تر است.

احتمال وقوع خطای راندوم + و خطای راندوم - با هم برابرند.

احتمال وقوع خطاهای راندوم خیلی بزرگ، خیلی کم است.

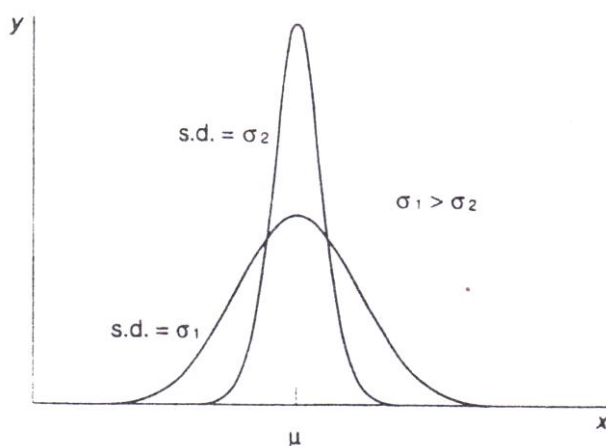
منحنی خطای نرمال را می‌توان با یکی از روش‌های بیان دقت درجه‌بندی کرد:

۶۸/۳٪ از منحنی بین +۱S و -۱S است یعنی ۶۸/۳٪ از آنالیزها تنها خطای بین +۱S و -۱S را دارند.

۹۵/۵٪ از منحنی بین +۲S و -۲S است.

۹۹/۷٪ از منحنی بین +۳S و -۳S است.

برطبق فرمول انحراف استاندارد، هرچه n (تعداد تکرار آزمایش‌ها) بیش‌تر شود، S (انحراف استاندارد) کم‌تر و در نتیجه خطای راندوم کم‌تر می‌شود. با افزایش تعداد تکرار آنالیزها S کوچک شده و منحنی خطای نرمال فشرده‌تر خواهد شد.



۳-۱- کاربردهای تستهای آمار در شیمی تجزیه:

۱. Student t – test:

هدف از این تست، تعیین فاصله و حدود اطمینان است که با یک درصد اطمینان مشخص، μ در آن محدوده باشد.

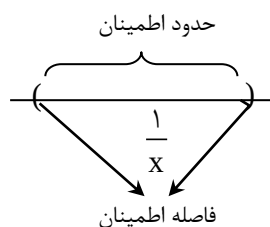
$N - 1$ = درجه آزادی

$$\mu = \bar{x} \pm \frac{tS}{\sqrt{N}}$$

N = تعداد آنالیزها

S = انحراف استاندارد

t = یک عدد آماری که از جداول با سطح اطمینان و درجه آزادی بدست می‌آید.



هر جا S داشته باشیم می‌توانیم آن را با δ هم جابه‌جا کنیم و فرمول بر حسب S بصورت زیر می‌گردد:

برای به‌دست آوردن Z از جداول به درجه آزادی نیاز نداریم و فقط با داشتن سطح اطمینان به‌دست می‌آید.

$$\Rightarrow \mu = \bar{x} \pm \frac{z\delta}{\sqrt{N}}$$

۲. تعیین دقت انحراف استاندارد:

$$q = \sqrt{\frac{S^2(N-1)}{N}} = \sqrt{\frac{\sum (x_i - \bar{x})^2}{N}}$$

X^x : پارامتر آماری بزرگ

$$S \pm \sqrt{\frac{X^x \times q}{X^y \times q}}$$

S = حدود اطمینان برای

X^y : پارامتر آماری کوچک

مثال: ۱۰ بار اندازه گیری تکراری سرب در نمونه خاک مقدار متوسط ۰/۱۴۶۲ با انحراف استاندارد ۰/۰۰۷۴ را به دست می دهد. با ۹۵٪ اطمینان، فاصله اطمینان را برای میانگین و انحراف استاندارد محاسبه کنید.

با درصد اطمینان ۹۵٪ و درجه آزادی ۹ از روی جدول به دست می آید

$$\bar{x} \pm \frac{tS}{\sqrt{N}} = 0.1462 \pm \frac{2.262 \times 0.0074}{\sqrt{10}}$$

$$q = \sqrt{\frac{S^2(N-1)}{N}} = \sqrt{\frac{(0.0074)^2 \times 9}{10}} = 0.007$$

$$\text{فاصله اطمینان برای انحراف استاندارد} = 0.0074 \pm \begin{cases} 1.77 \times 0.007 \\ 0.66 \times 0.007 \end{cases}$$

۳. حذف نتایج مشکوک:

می خواهیم بدانیم خطایی که باعث اختلاف نتایج شده ناشی از خطای سیستماتیک (و یا فاحش) است (و باید آن را حذف کنیم) و یا ناشی از خطای راندوم است (و باید آن را نگه داریم)

روش اول: عدد مشکوک را کنار می گذاریم و برای سایر عددها $\bar{x}$ و $\bar{d}$ را حساب می کنیم. اگر اختلاف عدد مشکوک از میانگین $\bar{d} \geq 4\bar{d}$ باشد آن عدد مشکوک نیست و تفاوت آن عدد ناشی از خطای راندوم است و باید عدد را نگه داشت، اما اگر اختلاف عدد مشکوک از میانگین $\bar{d} < 4\bar{d}$ باشد تفاوت ناشی از خطای سیستماتیک است و باید آن عدد را حذف کرد:

if $|x_i - \bar{x}| \leq 4\bar{d}$ باید عدد را نگه داریم

if $|x_i - \bar{x}| > 4\bar{d}$ باید عدد را حذف کنیم.

نکته: $4\bar{d}$ به سطح اطمینان بستگی دارد. برای $4\bar{d}$ ، سطح اطمینان ۹۹٪ است.

روش دوم (تست P): عدد مشکوک را کنار گذاشته و $\bar{x}$ و S یا δ را حساب می کنیم: (S را بدون احتساب x_i به دست می آوریم)

اختلاف در حد خطای راندوم و باید عدد را نگه داریم if $|x_i - \bar{x}| \leq 3S$ or 3δ

خطای سیستماتیک و باید عدد را حذف کنیم. if $|x_i - \bar{x}| > 3S$ or 3δ

اصل صفر: هر تست آماری یک جمله منفی دارد، مثلاً این عدد مشکوک نیست جمله منفی است. هر تست آماری یا اصل صفر را تأیید می کند یا اصل صفر را رد می کند. اصل صفر را با H^0 نمایش می دهند. در تست های آماری ممکن است دو نوع خطا ایجاد شود: این که تست آماری اصل صفری را که درست است را رد کند (خطای نوع اول) یا این که تست آماری اصل صفری را که درست نیست تأیید کند (خطای نوع دوم) هرچه تعداد تکرار انجام آزمایش ها بیش تر باشد، احتمال وقوع این دو نوع خطا کم تر می شود.

روش سوم (تست Q یا دیکسون): ابتدا باید داده ها را از کوچک به بزرگ مرتب کنیم. و دو حالت پیش می آید، عدد مشکوک یا از همه بزرگ تر است و یا از همه کوچک تر است.

$$Q = \frac{|X_i - X_n|}{|X_i - X_1|} \quad \text{اگر عدد مشکوک از همه بزرگتر باشد.} \quad Q = \frac{|X_1 - X_i|}{|X_n - X_i|} \quad \text{اگر عدد مشکوک از همه کوچکتر باشد.}$$

$$Q_{\text{test}} = \frac{|\text{تفریق نزدیکترین عدد به عدد مشکوک}|}{|\text{تفریق دورترین عدد به عدد مشکوک}|}$$

فرمول کلی

Q_{table} را می‌توانیم با درصد اطمینان و درجه آزادی به‌دست آوریم.

نکته: درجه اطمینان $(N-1)$ شامل عدد مشکوک هم هست.

$$\text{if } Q_{\text{test}} \leq Q_{\text{table}}$$

عدد مشکوک نیست / اصل صفر تأیید می‌شود / عدد نگه داشته می‌شود.

$$\text{if } Q_{\text{test}} > Q_{\text{table}}$$

عدد مشکوک است / اصل صفر رد می‌شود / عدد حذف می‌شود.

روش چهارم (تست t_n): عدد مشکوک را کنار می‌گذاریم و $\bar{X}$ و S را حساب می‌کنیم. S را بدون در نظر گرفتن $X_?$ حساب می‌کنیم.

$$T_{n_{\text{test}}} = \frac{|X_? - \bar{X}|}{S}$$

$$\text{if } T_{n_{\text{exp}}} \leq T_{n_{\text{table}}}$$

خطا راندوم است / عدد مشکوک نیست / اصل صفر تأیید می‌شود / عدد باید نگه داشته شود.

$$\text{if } T_{n_{\text{exp}}} > T_{n_{\text{table}}}$$

خطا سیستماتیک است / عدد مشکوک است / اصل صفر رد می‌شود / عدد باید حذف شود.

نکته جالب: تمام پارامترهای جداول آماری حداکثر خطای راندوم را در نظر گرفته‌اند. پس اگر پارامترهای $\exp$ یا test کوچک‌تر و یا برابر آن‌ها باشند، در نتیجه خطا در حد خطای راندوم خواهد بود، اما اگر پارامترهای $\exp$ یا test بزرگ‌تر از آن باشند یعنی خطا بیش‌تر از میزان خطای راندوم و احتمالاً خطای سیستماتیک رخ داده است و خطای سیستماتیک قابل قبول نیست و باید حذف شود.

۴. تست t :

این تست دو کاربرد دارد:

۱. بررسی این‌که آیا اختلاف معنی‌داری بین μ و $\bar{X}$ وجود دارد یا نه؟

$$t_{\text{test}} = \frac{|\mu - \bar{X}|}{\frac{S}{\sqrt{N}}} \quad \mu = \bar{X} \pm \frac{ts}{\sqrt{N}} \quad \text{از این فرمول به‌دست آمده است} \Leftarrow$$

$$\text{if } t_{\text{test}} \leq t_{\text{table}}$$

اختلاف بین μ و $\bar{X}$ معنی‌دار نیست / اصل صفر تأیید می‌شود / اختلاف $\bar{X}$ و μ ناشی از خطای راندوم است.

$$\text{if } t_{\text{test}} > t_{\text{table}}$$

اختلاف μ و $\bar{X}$ معنی‌دار است / اصل صفر رد می‌شود / اختلاف $\bar{X}$ و μ ناشی از خطای سیستماتیک است.

۲. مقایسه بین دو $\bar{X}$ و بررسی این‌که اختلاف بین دو میانگین معنی‌دار است یا نه؟

$$t_{\text{test}} = \frac{|\bar{X}_r - \bar{X}_1|}{S_p} \sqrt{\frac{N_1 N_r}{N_1 + N_r}}$$

t_{table} با استفاده از درصد اطمینان و درجه آزادی $(n_1 + n_r - 2)$ نیاز داریم:

$$S_p = \sqrt{\frac{(n_1 - 1)S_1^2 + (n_r - 1)S_r^2}{n_1 + n_r - 2}}$$

if $t_{\text{test}} \leq t_{\text{table}}$

اختلاف در $\bar{X}$ معنی دار نیست / اختلاف ناشی از خطای راندوم است

if $t_{\text{test}} > t_{\text{table}}$

اختلاف دو $\bar{X}$ معنی دار است / اصل صفر رد می شود.

فرمول اصلی t_{test} این است:

$$t_{\text{test}} = \frac{|\bar{X}_r - \bar{X}_1|}{S} \text{ و } S = SP \sqrt{\frac{n_1 + n_r}{n_1 n_r}}$$

در فرمول قبل فرض بر این است که اختلاف بین S_1 و S_r معنی دار نیست و می توان آن را برابر گرفت.

$$S^2 = \frac{S_1^2}{n_1} + \frac{S_r^2}{n_r} \Rightarrow S^2 = S_p^2 \left(\frac{1}{n_1} + \frac{1}{n_r} \right) \Rightarrow S^2 = S_p^2 \left(\frac{n_1 + n_r}{n_1 n_r} \right)$$

اما اگر S ها خیلی با هم تفاوت داشته باشند و اختلاف آنها معنی دار باشد، باید t_{test} را از این رابطه به دست آوریم.

$$t_{\text{test}} = \frac{|\bar{X}_r - \bar{X}_1|}{\sqrt{\frac{S_1^2}{n_1} + \frac{S_r^2}{n_r}}}$$

و برای به دست آوردن درجه آزادی مربوط به t_{table} از این رابطه به دست می آید:

$$\text{تعداد درجات آزادی که باید روند شود} = \frac{\left(\frac{S_1^2}{n_1} + \frac{S_r^2}{n_r} \right)^2}{\left(\frac{S_1^2}{n_1(n_1 - 1)} + \frac{S_r^2}{n_r(n_r - 1)} \right)}$$

اگر تعداد آنالیزهای حالت اول و دوم برابر باشند و جفت-جفت با هم متناظر باشند مثلاً یک قرص را دو نیم کرده باشیم و هر نیم را با روش ۱ و نیم دیگر را با روش ۲ آنالیز کرده باشیم جهت مقایسه نتایج روش ۱ و ۲، با فرمول زیر به دست می آید. ستون d اختلاف نتایج ۱ و ۲ از یکدیگر می باشد.

$$\begin{array}{ccc} 1 & 2 & 3 \\ x_1 & x'_1 & d_1 = x_1 - x'_1 \\ x_r & x'_r & d_r = x_r - x'_r \\ x_r & x'_r & d_r = x_r - x'_r \end{array} \Rightarrow t_{\text{test}} = \frac{\bar{d} \sqrt{N}}{S_d}$$

مثال: مقادیر زیر برای وزن اتمی کادمیم به دست آمده اند.

۱۱۲/۳۲ ۱۱۲/۲۵ ۱۱۲/۳۶ ۱۱۲/۲۱ ۱۱۲/۳۰ ۱۱۲/۳۶

آیا متوسط مقادیر فوق با مقدار پذیرفته شده برای کادمیم با ۹۹٪ اطمینان اختلاف معنی دار دارد یا نه؟ (جرم اتمی کادمیم: ۱۱۲.۴)

$$\bar{x} = 112/3, S_x = 0.603$$

$$t_{\text{test}} = |\mu - \bar{x}| \frac{\sqrt{N}}{S} = |112/4 - 112/3| \frac{\sqrt{6}}{0.603} = 4/0.6$$

$$t_{\text{table}} = 4/0.32 \Rightarrow t_{\text{test}} > t_{\text{table}} \quad \text{اختلاف معنی دار است}$$

خطای روش سیستماتیک است و قابل قبول نیست.

مثال: دو روش آنالیز برای اندازه‌گیری جزء X در نمونه اعمال و درصد X به صورت زیر به دست آمده است. آیا اختلاف معنی‌داری بین نتایج وجود دارد؟

	test1	test2	test3	test4	
Method A	4/68	4/63	4/69	4/55	$\Rightarrow \bar{x}_A = 4/64$ $S_A = 0.637$

Method B	4/81	4/7	4/74	$\Rightarrow \bar{x}_B = 4/75$ $S_B = 0.556$
----------	------	-----	------	---

$$S_p = \sqrt{\frac{3 \times (0.637)^2 + 2 \times (0.556)^2}{4 + 3 - 2}} = \sqrt{0.36} = 0.6$$

$$t_{\text{test}} = \frac{|4/75 - 4/64|}{0.6} \times \sqrt{\frac{4 \times 3}{4 + 3}} = 2/4 \quad t_{\text{table}} = 2/52 \text{ (درجه آزادی ۵ و ۹۵\%)}$$

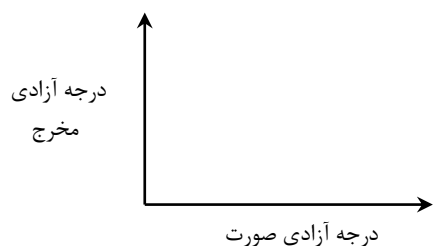
$$t_{\text{test}} < t_{\text{table}} \quad \text{اصل صفر تأیید می‌شود / خطا راندوم است:}$$

۵. تست F: مقایسه دقت دو روش با هم

نکته: F همیشه باید بزرگ‌تر از یک باشد پس همیشه V بزرگ‌تر را در صورت می‌گذاریم.

$$F_{\text{test}} = \frac{V_1}{V_2}$$

F_{table} را بر حسب درجه آزادی صورت و درجه آزادی مخرج به دست می‌آوریم.



$$F_{\text{test}} \leq F_{\text{table}} \quad \text{اختلاف معنادار نیست / اصل صفر تأیید می‌شود.}$$

$$F_{\text{test}} > F_{\text{table}} \quad \text{اختلاف معنادار است / اصل صفر رد می‌شود. اختلاف دقت دو روش معنی‌دار می‌باشد.}$$

۶. تست chi:

۱. کاربرد اول این تست: داده‌ها را بررسی می‌کنیم تا ببینیم آیا در ثبت داده‌ها تعصب خاصی بوده یا نه؟

فرکانس موردانتظار فرکانس مشاهده شده

$$\chi^2_{\text{test}} = \frac{\sum (f_i - F_i)^2}{F_i}$$

نکته: اگر F_i برای همه یکسان نباشد باید از این فرمول استفاده کنیم.

$$\chi^2_{\text{test}} = \sum \frac{(f_i - F_i)^2}{F_i}$$

$\chi^2_{\text{test}} \leq \chi^2_{\text{table}}$ در خواندن اطلاعات تعصب نبوده / اصل صفر تأیید شد / خطا راندوم است

$\chi^2_{\text{test}} > \chi^2_{\text{table}}$ در خواندن اطلاعات تعصب بوده / اصل صفر رد شد / خطا سیستماتیک است

توضیح: فرکانس مشاهده شده تعداد دفعاتی است که یک نتیجه به دست آمده است و فرکانس موردانتظار معمولاً تعداد کل اندازه‌گیری‌ها تقسیم بر تعداد نتایج مختلف به دست آمده است.

مثال: نتایج حاصل از خواندن حجم بورت و خواندن وزن در ترازو به صورت زیر است:

عدد آخر	خواندن حجم بورت	F_i	خواندن ترازو	F_i
.	۲۱۲	$\frac{۱۵۰۰}{۱۰}$	۱۲۶	۱۰۰۰/۱۰
۱	۲۱۲	۱۵۰	۹۲	۱۰۰
۲	۲۲۹	۱۵۰	۹۵	۱۰۰
۳	۱۶۶	۱۵۰	۸۵	۱۰۰
۴	۱۲۴	۱۵۰	۱۱۲	۱۰۰
۵	۱۰۷	۱۵۰	۱۳۲	۱۰۰
۶	۱۱۰	۱۵۰	۷۷	۱۰۰
۷	۸۱	۱۵۰	۹۵	۱۰۰
۸	۱۸۴	۱۵۰	۱۰۷	۱۰۰
۹	$\frac{۱۲۵}{۱۵۰۰}$	۱۵۰	$\frac{۸۵}{۱۰۰۰}$	۱۰۰

مثلاً ۸۱ بار عدد ۷ خوانده شده‌اند.

اگر هیچ خطایی نبود با توجه به ۱۵۰۰ بار تکرار آزمایش

و وجود امکان ۱۰ عدد (از صفر تا نه) عدد ۷ را باید ۱۵۰ بار مشاهده می‌کردیم.

$$\chi^2_{x, \text{st}} = ۲۷/۱۶/۹$$

$$\chi^2_{\text{test}} = \frac{\sum (f_i - F_i)^2}{F_i} = ۱۵۹/۶۸ > \chi^2_{\text{table}} = ۱۶/۹ \Rightarrow \text{تعصب وجود داشته}$$

(۹۵٪ و ۹ درجه آزادی)

بنابراین در خواندن ترازو نیز تعصب وجود داشته است $\Rightarrow \chi^2_{\text{test}} = ۲۷/۱ > ۱۶/۹$

۷. آنالیز واریانس: برای مقایسه چندین $\bar{x}$ و این که آیا اختلاف بین آن‌ها معنی‌دار است یا نه به عنوان مثال یک تیتراسیون را

به k دانشجو می‌دهیم و هر دانشجو این تیتراسیون را n بار انجام می‌دهد.

$$\begin{array}{c|c}
 \begin{array}{cccc}
 1 & x_1 & x_2 & \dots & x_n \\
 2 & x_1 & x_2 & \dots & x_n \\
 \vdots & \vdots & \vdots & & \vdots \\
 k & x_1 & x_2 & \dots & x_n
 \end{array} &
 \begin{array}{c}
 \bar{x}_1 \\
 \bar{x}_2 \\
 \vdots \\
 \bar{x}_k \\
 \hline
 \bar{x}_{total}
 \end{array}
 \end{array}$$

x_{ij} یعنی مثلاً نتیجه x_j دانشجو i ام

$$\sum_{i=1}^k \sum_{j=1}^n (x_{ij} - \bar{x})^2 = \underbrace{\sum_{i=1}^k \sum_{j=1}^n (x_{ij} - \bar{x}_i)^2}_{S_w} + \underbrace{n \sum_{i=1}^k (\bar{x}_i - \bar{x}_{total})^2}_{S_b}$$

جمع مربعات تغییرات درون گروهی جمع مربعات تغییرات بین گروهی

S_w within

S_b between

نکته: $n \times k$ را در بعضی کتاب‌ها N^* هم نشان داده‌اند.

$$F_{test} = \frac{S_b / (k - 1)}{S_w / (n \times k) - k}$$

نکته: S_b را در صورت کسر می‌گذاریم چون انحراف داده‌های بین گروهی بیش‌تر است و F باید بزرگ‌تر از یک باشد.

نکته: F_{table} را با استفاده از درجه آزادی صورت $(k - 1)$ و درجه آزادی مخرج $(n \times k) - k$ از جدول به‌دست می‌آوریم.

نکته: اگر تعداد آزمایش‌ها در هر گروه‌ها با هم برابر نباشند (تعداد کل آزمایش‌های همه گروه‌ها - تعداد گروه‌ها) درجه آزادی

مخرج است و در ضمن در این حالت فرمول S_b کمی تغییر می‌کند: $\sum_{i=1}^k n_i (\bar{x}_i - \bar{x}_{total})^2$

$F_{test} \leq F_{table}$ اختلاف معنی‌دار وجود ندارد.

$F_{test} > F_{table}$ اختلاف معنی‌دار وجود دارد.

مثال: فرض کنید ۵ آنالیز مختلف به نام‌های A, B, C, D, E مقدار آهن را در آب اندازه‌گیری و مقادیر آن را با واحد ppb

گزارش کرده‌اند. با آنالیز واریانس مشخص نمایید که آیا مقادیر میانگین مشاهده شده توسط آن‌ها به‌صورت معنی‌دار با هم

اختلاف دارند یا نه؟

A	B	C	D	E
۱۰.۳	۹.۵	۱۲.۱	۷.۶	۱۳.۶
۹.۸	۸.۶	۱۳	۸.۳	۱۴.۵
۱۱.۴	۸.۹	۱۲.۴	۸.۲	۱۵.۱

$$\bar{x}_A = 10.5 \quad \bar{x}_B = 9 \quad \bar{x}_C = 12.5 \quad \bar{x}_D = 8.033 \quad \bar{x}_E = 14.4 \quad \bar{x}_{total} = 10.81$$

$$S_b = n \sum (\bar{x}_i - \bar{x}_{total})^2 = 80.396$$

$$S_w = 3.61$$

$$\frac{80.396}{4}$$

$$F_{test} = \frac{4}{3.61} = 55.68 > F_{table} = 5.99$$

$$(3 \times 5) - 5$$

اختلاف معنادار وجود دارد



۴-۱ یادآوری چند مفهوم و فرمول ابتدایی:

$$M = \frac{\text{تعداد مولها}}{\text{حجم (lit)}} = \text{مولاریته}$$

$$\text{تعداد مولها} = \frac{m(\text{gr})}{f_w \leftarrow \text{وزن فرمولی}}$$

$$N = \frac{\text{تعداد اکی والانها}}{\text{حجم (lit)}} = \text{نرمالیه}$$

$$n \times \text{تعداد مولها} = \frac{m}{f_w / n} = \text{تعداد اکی والانها}$$

n در اسیدها: تعداد H^+ اسیدی

n در بازها: تعداد OH^- بازی

n در نمکها: ظرفیت فلز $\times$ تعداد فلز

n در واکنش اکسایش و کاهش: تعداد الکترونهای جابه‌جا شده

$$\text{نرمالیه} = \frac{m(\text{gr})}{f_w} = \frac{m}{f_w \cdot V(\text{lit})} \times n = M \times n$$

$$\text{نرمالیه} = \frac{n \cdot V(\text{lit})}{f_w} \quad \text{دانسیتة درصد خلوص}$$

$$\text{نرمالیه} = \frac{10 \times a \times d}{f_w} \quad \text{دانسیتة درصد خلوص}$$

$$\frac{V}{V} = \frac{\text{حجم خالص}}{\text{حجم ناخالص}} \times 100$$

$$\frac{W}{W} = \frac{\text{وزن خالص}}{\text{وزن ناخالص}} \times 100$$

سه نوع درصد خلوص داریم:

$$\frac{W}{W} = \frac{\text{وزن خالص}}{\text{وزن ناخالص}} \times 100$$

تمرین ۱: در صورتی که انحراف معیار مراحل مختلف در یک آزمایش برابر ۰.۳ و ۰.۴ باشد، انحراف معیار کل چند است؟

$$S_{\text{total}}^2 = S_1^2 + S_2^2 = 0.3^2 + 0.4^2 = 0.25$$

$$S_{\text{Total}} = \sqrt{V_T} = \sqrt{0.25} = 0.5$$

تمرین ۲: یک روش تجزیه‌ای خاص، خطای ثابتی به اندازه ۲.۵mg در اندازه‌گیری مقدار آهن از خود نشان می‌دهد. اگر بخواهیم نمونه‌ای از یک سنگ معدن حاوی ۱۲.۵٪ آهن را توسط این روش مورد آنالیز قرار دهیم و در این کار، خطای بیش از ۱٪ نداشته باشیم، حداقل نمونه مورد نیاز جهت انجام این آزمایش چقدر است؟

$$\bar{x} - \mu = 2.5 \text{mg}$$

$$\frac{\bar{x} - \mu}{\mu} \times 100 = 1\% \Rightarrow \frac{2.5 \text{mg}}{\mu} = \frac{1}{100} \Rightarrow \mu = 250 \text{mg}$$

مقدار واقعی نمونه خالص

$$\frac{m_{\text{خالص}}}{m_{\text{ناخالص}}} \times 100 = 12.5\% \Rightarrow \frac{250 \text{mg}}{m_{\text{ناخالص}}} = \frac{12.5}{100}$$

$$\Rightarrow m_{\text{ناخالص}} = 2000 \text{mg} = 2 \text{gr}$$

سنگ معدن



تمرین ۳: یک کارخانه تولیدکننده لامپ‌های روشنایی مدعی است که محصولات این کارخانه به طور متوسط ۱۲۰۰ ساعت کار می‌کند. در صورتی که انحراف معیار کارکرد این لامپ‌ها ۱۰۰ ساعت باشد از میان ۴۰۰ عدد لامپ تولیدی، انتظار می‌رود چند لامپ بیش از ۱۴۰۰ ساعت کار کند کارکرد بیش از ۱۴۰۰ ساعت یعنی بار انحراف معیار بیش از $2S$ + چون $S=100$ است.

۹۵.۵٪ از این ۴۰۰ لامپ بین $2S$ +، یعنی بین ۱۰۰۰ تا ۱۴۰۰ ساعت کار می‌کنند.
پس ۴.۵٪ از لامپ‌ها کمتر از ۱۰۰۰ ساعت یا بیش‌تر از ۱۴۰۰ ساعت کار می‌کند.

$$\frac{x}{400} = \frac{4.5}{100} \Rightarrow x = 18$$

و تنها نیمی از ۴.۵٪ احتمال دارد که بیش از ۱۴۰۰ ساعت کار کنند یعنی تنها ۹ لامپ از ۴۰۰ لامپ بیش از ۱۴۰۰ ساعت عمر خواهد کرد.

۱-۵- آنالیز رگرسیونی:

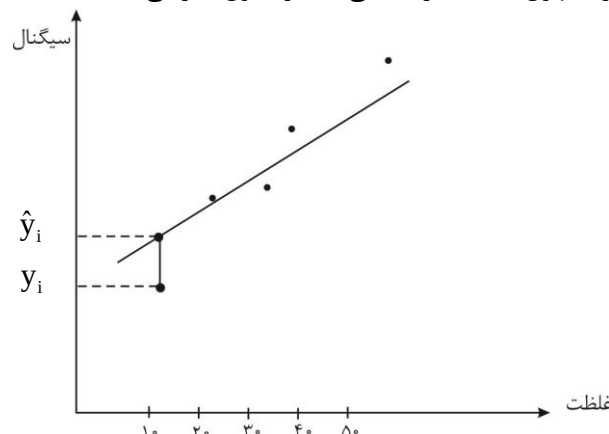
در تجزیه دستگاهی همیشه به دنبال سیگنالی متناسب با غلظت هستیم و ایده‌آل این است که سیگنال و غلظت رابطه خطی داشته باشند.

مثلاً در جذب اتمی:

غلظت $\swarrow$
سیگنال $\searrow$
 $A = \epsilon bc$
یا در پتانسیومتری:

$$E = \text{const} - \frac{0.05915}{n} \log a_x$$

یکی از ساده‌ترین روش‌های آنالیز مجهول استفاده از منحنی کالیبراسیون خارجی است.



امروزه با Excel بهترین خط راست را از بین نقاط عبور می‌دهیم. Excel این کار را با استفاده از روش کم‌ترین مربعات (least Square) انجام می‌دهد. Excel خط را به طریقی رسم می‌کند که مجموع مربعات فاصله نقاط از خط، حداقل مقدار باشد.

نکته: در روش least Square خطا را فقط در محور y در نظر می‌گیریم و می‌گوییم چون x ، یا همان غلظت، مربوط به نمونه استاندارد است پس هیچ خطایی ندارد. روشی که خطای x را هم در نظر می‌گیرد؛ روش وزنی است که در شیمی تجزیه کاربردی ندارد.

اگر سیگنالی را که از دستگاه به دست آمده را با y_i و عرضی را که از روی معادله خط به دست می‌آوریم. با $\hat{y}_i$ نشان دهیم، طبق این معادله بهترین خط باید به طریقی رسم شود که Q کمترین مقدار باشد.

$$Q = \sum (y_i - \hat{y}_i)^2 = \sum (y_i - (a + bx_i))^2$$

با محاسبات ریاضی درمی‌یابیم که اگر

$$b = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sum (x_i - \bar{x})^2} \quad \text{و} \quad a = \bar{y} - b\bar{x}$$

باشند، Q مینیمم است.

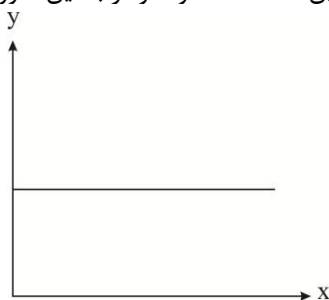
$$Q_{\min} = \sum (y_i - \bar{y})^2 - b^2 \sum (x_i - \bar{x})^2$$

هدف استفاده از آنالیز رگرسیون برای بررسی این است که آیا رابطه بین غلظت و سیگنال دستگاه خطی است یا نه؟

حالت‌های حدی

۱. اگر $\sum (x_i - \bar{x})^2 = 0$ پس $Q_{\min} = 0$ و رابطه کاملاً خطی است.

۲. اگر $b = 0$ پس شیب خط صفر است و بدترین حالت است و نمودار به این صورت می‌شود.



و می‌خواهیم حالت‌های بین دو حالت حدی را بررسی کنیم. برای این کار دو واریانس تعریف می‌کنیم.

$$V_1 = \frac{b^2 \sum (x_i - \bar{x})^2}{1}$$

و

$$V_2 = \frac{\sum (y_i - \bar{y})^2 - b^2 \sum (x_i - \bar{x})^2}{n - 2}$$

تعداد نقاط کالیبراسیون

$$\Rightarrow F_{\text{exp}} = \frac{V_1}{V_2}$$

و برای به دست آوردن F_{table} ، درجه آزادی صورت یک است چون تغییرات در دو محور y ، x بررسی شده و $(2-1) = 1$ درجه آزادی صورت است و درجه آزادی مخرج هم $n - 2$ است.

$$\text{if } F_{\text{exp}} \leq F_{\text{table}}$$

رابطه خطی وجود ندارد / اصل صفر تأیید می‌شود

$$\text{if } F_{\text{exp}} > F_{\text{table}}$$

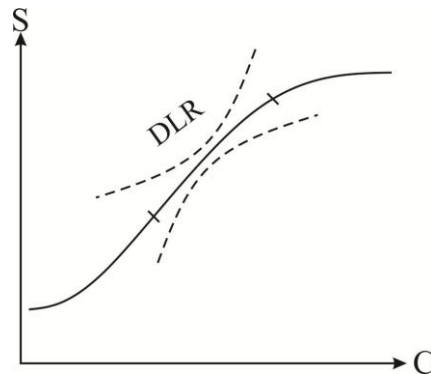
رابطه خطی وجود دارد / اصل صفر رد می‌شود

اما روش راحت‌تر برای بررسی خطی بودن معادله، استفاده از ضریب همبستگی (r) است.



$$r = \frac{\sum(x_i - \bar{x})(y_i - \bar{y})}{\left(\sum(x_i - \bar{x})^2 \sum(y_i - \bar{y})^2\right)^{\frac{1}{2}}}$$

$r = 1$ بهترین حالت است و در این حالت معادله کاملاً خطی است و r باید حداقل 0.99 باشد تا رابطه‌ای را خطی بنامیم. به ناحیه‌ای از منحنی که بین غلظت و سیگنال رابطه خطی وجود دارد (Dynamic Linear Range) DLR می‌گوییم. هرچه گستره DLR بیش‌تر باشد، مزیتی برای آن روش است.



توزیع خطای راندوم روی منحنی را با خط‌چین نشان داده‌ایم. یعنی خطای راندوم در وسط DLR از کناره‌ها کم‌تر است. پس بهترین حالت این است که غلظت (و سیگنال) گونه مجهول در وسط DLR بیفتد. تا مقدار خطای راندوم در اندازه‌گیری غلظت مجهول کمترین مقدار باشد.

گاهی می‌توان با S_b و یا انحراف استاندارد در اندازه‌گیری شاهد هم جایگزین نمود.

$$S_c = \frac{S_y}{b} \sqrt{\frac{1}{m} + \frac{1}{n} + \frac{(y_c - \bar{y})^2}{b^2 \sum(x_i - \bar{x})^2}}$$

این فرمول برای محاسبه غلظت نمونه مجهول (S_c) بر اساس سیگنال آن (S_y) و سایر پارامترها استفاده می‌شود. توضیحات زیر به اجزای فرمول اشاره می‌کند:

- S_y : سیگنال نمونه مجهول
- b : شیب منحنی کالیبراسیون
- m : تعداد تکرار آنالیز مجهول
- n : تعداد نقاط منحنی کالیبراسیون
- $(y_c - \bar{y})^2$: میانگین سیگنال‌های دستگاه برای نمونه‌های استاندارد
- $\sum(x_i - \bar{x})^2$: انحراف استاندارد در رسم کالیبراسیون
- $\frac{1}{m} + \frac{1}{n}$: غلظت هر نمونه
- $\frac{(y_c - \bar{y})^2}{b^2 \sum(x_i - \bar{x})^2}$: انحراف استاندارد
- $\frac{S_y}{b}$: اندازه‌گیری غلظت
- S_c : شیب منحنی کالیبراسیون نمونه مجهول

برای این که دقت اندازه‌گیری غلظت بیش‌تر باشد و در واقع مقدار S_c کم‌تر باشد، باید تلاش کنیم تا b هرچه بزرگ‌تر، m و n هرچه بزرگ‌تر و $(y_c - \bar{y})$ هرچه کوچک‌تر باشد و DLR وسیع‌تر باشد.

$(y_c - \bar{y})$ کوچک‌تر باشد یعنی این که y_c به $\bar{y}$ نزدیک‌تر باشد یعنی هرچه سیگنال نمونه مجهول نزدیک‌تر به $\bar{y}$ و یا در وسط DLR باشد، دقت روش بیش‌تر است، چون خطای راندوم کم‌تر است.

بهترین کار برای این که غلظت (و سیگنال) مجهول در وسط DLR بیفتد، استفاده از Curve Bracket کردن است. برای این کار باید اول سیگنال نمونه مجهول را بگیریم. بعد یک نمونه استاندارد با غلظت مشخص می‌سازیم. سپس با این یک نمونه استاندارد غلظت حدودی نمونه مجهول را به دست آورده و استانداردها چند تا بالاتر و چند تا پایین‌تر از نمونه مجهول تهیه می‌کنیم.

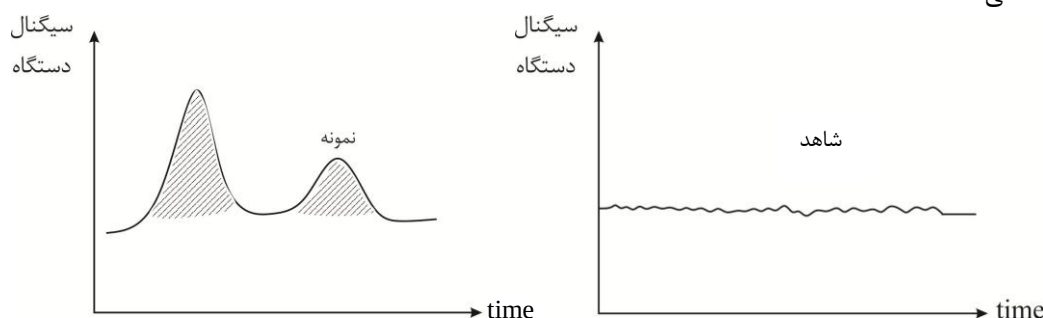
نکته: درست است که روشی که DLR بزرگ‌تری داشته باشد، بهتر است، اما این درست نیست که انواع مجهول با غلظت‌های متفاوت را با یک نمودار سنجش کنیم و بهتر است از Curve Bracket استفاده کنیم.

برای به دست آوردن سیگنال شاهد (S_b)، مثلاً در جذب اتمی که یک روش دیجیتالی است و نتیجه را به صورت یک عدد دیجیتالی به ما می دهد، چندبار نمونه شاهد را آنالیز می کنیم و بعد از روی نتایج به دست آمده برای شاهد، S_b را محاسبه می کنیم (با فرمول مربوط به S که قبلاً داشتیم)

$$S_{\frac{y}{x}} = \sqrt{\frac{\sum (y_i - \hat{y}_i)^2}{n-2}}$$

تعداد نقاط منحنی کالیبراسیون

حالا اگر دستگاه دیجیتالی نباشد یعنی مثلاً نتیجه را مثل GC یا HPLC به صورت یک پیک (غلظت متناسب با سطح زیر پیک است) ارائه می کند.



در مورد این دستگاه ها که نمی توانیم S_b را به صورت دیجیتالی حساب کنیم به جای S_b ، $S_{\frac{y}{x}}$ را استفاده می کنیم.

در فرمول $S_{\frac{y}{x}}$ ، y_i سیگنال خوانده شده از دستگاه برای هر نمونه استاندارد در منحنی کالیبراسیون می باشد و $\hat{y}_i$ سیگنال معادل آن نمونه استاندارد در روی منحنی خطی می باشد.

$$\hat{y}_i = bx_i + a$$

حد تشخیص (Detection limit): حداقل غلظتی است که سیگنال آن از سیگنال زمینه قابل تشخیص است (لزوماً قابل اندازه گیری نیست)

$$DL \text{ or } LOD = \frac{3S_b}{b} \leftarrow \text{اگر دستگاه دیجیتالی نباشد به جای } S_b \text{ از } S_{\frac{y}{x}} \text{ استفاده می کنیم}$$

اگر y_c (سیگنال نمونه) را از y_B (سیگنال شاهد) کم کنیم باید ۳ برابر انحراف استاندارد شاهد باشد تا سیگنال آن نمونه قابل تشخیص از سیگنال زمینه باشد.

$$y_c - y_B = 3S_b \quad (1)$$

$$y = a + bx \Rightarrow \text{if } x = 0 \Rightarrow \text{شاهد} = \text{محلول} \Rightarrow y = a = y_B \Rightarrow \text{سیگنال زمینه} \Rightarrow y_c = y_B + bx$$

$$\Rightarrow y_c - y_B = bx \quad (2)$$

$$LOD \leftarrow x = \frac{3S_b}{b} \text{ است} \Rightarrow \text{اگر این رابطه برقرار باشد } x \text{ همان حداقل غلظتی می شود که قابل تشخیص است}$$

$$(1), (2) \Rightarrow 3S_b = bx$$

LOD حداقل غلظتی که سیگنال آن از سیگنال زمینه قابل تشخیص است.

این که DL لزوماً قابل اندازه گیری نیست یعنی این که ممکن است در قسمت DLR نباشد، پس قابل اندازه گیری نیست.

Determination Limit (LOQ) Limit Of Quantification: حداقل غلظتی است که قابل اندازه گیری است. یا Determination Limit

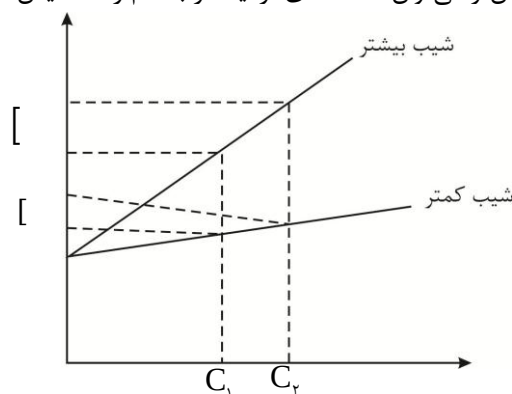
$$LOQ = \frac{1 \cdot S_b}{b}$$

نکته: اگر نمونه مجهول در محدوده DLR نبود، حق اندازه گیری نمونه مجهول را نداریم و برای سنجش باید نمونه مجهول را غلیظتر یا رقیقتر کنیم تا وارد گستره DLR بشود.

حساسیت (Sensitivity) دستگاه

دو نوع حساسیت داریم: (برای مقایسه حساسیت دو روش از مقایسه حساسیت تجزیه‌ای آن دو روش استفاده می‌گردد.)

۱. **حساسیت کالیبراسیونی:** شیب منحنی کالیبراسیون است که هرچه این شیب بیشتر باشد، حساسیت روش بیشتر است. در منحنی با حساسیت کالیبراسیونی بیشتر می‌توان غلظت‌های نزدیک‌تر به هم را تشخیص داد (بهتر تشخیص داد)



۲. **حساسیت تجزیه‌ای:** در یک نقطه تعریف می‌شود و از این رابطه به دست می‌آید.

$$\text{حساسیت تجزیه‌ای} = \frac{b}{S_{\text{sig}}} \leftarrow \begin{array}{l} \text{شیب منحنی کالیبراسیون} \\ \text{انحراف استاندارد سیگنال} \end{array}$$

انحراف استاندارد سیگنال در یک نقطه خاص (مثلاً در نقطه‌ای با غلظت ۲۰ ppm) پس حداقل باید ۳ بار ۲۰ ppm را آنالیز کرده باشیم.

$$S_{\text{sig}} = \sqrt{\frac{\sum (y_i - \bar{y})^2}{n-1}}$$

مثال: در اندازه‌گیری سرب به روش نشر شعله‌ای، معادله زیر به دست آمده است:

$$I = 0.0312 + 1.12 \text{ CPb}$$

در این معادله C غلظت سرب، واحد ppm و I شدت نشر سرب در ICP است.

نتایج زیر از اندازه‌گیری‌ها به دست آمده است.

Conc(Pb)	No. of Repeat	mean Value	S
۱۰	۱۰	۱۱.۶۲	۰.۱۵
۱	۱۰	۱.۱۴	۰.۰۲۵
۰	۲۹	۰.۰۲۹۶	۰.۰۰۸۲

الف. حساسیت کالیبراسیونی را به دست آورید.

ب. حساسیت تجزیه‌ای را در غلظت‌های ۱,۱ به دست آورید.

ج. LOD و LOQ را به دست آورید.

حساسیت کالیبراسیونی

$$I = 0.0312 + 1/12 C_{pb} \quad \text{حساسیت تجزیه‌ای در } 10 \text{ ppm} = \frac{1/12}{0.15}$$

$$1 \text{ ppm} \text{ حساسیت تجزیه‌ای} = \frac{1/12}{0.025}$$

$$LOD = 3 \frac{S_b}{b} = 3 \times \frac{0.0082}{1/12}$$

$$LOQ = 10 \cdot \frac{S_b}{b} = 10 \times \frac{0.0082}{1/12}$$

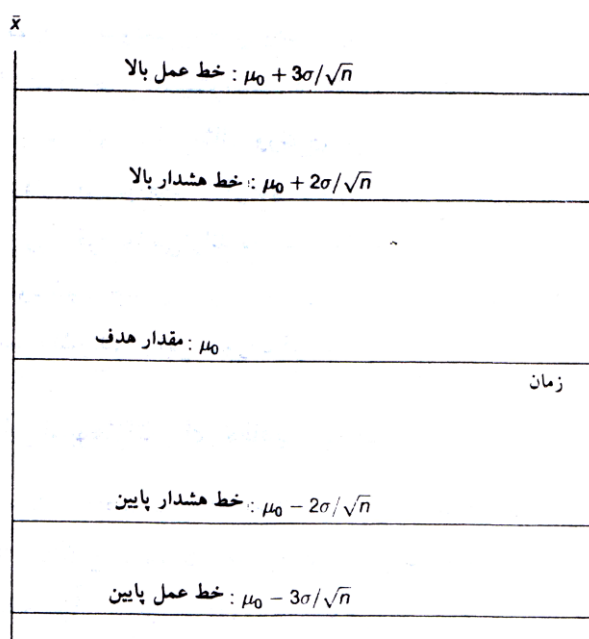
۱-۶- نقشه کنترل کیفی (Control Chart)

این روش (منحنی) در کارخانه‌هایی که یک آنالیز به شکل تکراری انجام می‌شود، برای نگهداری و کنترل خط تولید، خیلی قابل استفاده است.

در این روش میانگین $\bar{X}$ (توضیح این محور فاکتور موردنظر برای اندازه‌گیری $\bar{X}$ بر حسب نمونه‌های مورد اندازه‌گیری است) را روی محور $\bar{X}$ ها تعیین کرده و خطی از آن موازی با محور $\bar{X}$ می‌کشیم.

از داده‌های تکراری S_x را حساب می‌کنیم. بعد دو خط، یکی در $+1.96S_x$ و دیگری در $-1.96S_x$ موازی با خط $\bar{X}$ رسم می‌کنیم. به این دو خط، خطوط هشدار بالا و پایین $upper \text{ warning line}$ و $downer \text{ warning line}$ می‌گوییم. این خطوط هشدار معمولاً با ۹۵٪ درصد اطمینان رسم می‌شود (عدد ۱.۹۶ مربوط به درصد ۹۵٪ است). اگر از ۱۰۰ داده، تا ۵ داده از خطوط هشدار خارج شود مشکلی نیست اما اگر تعداد داده‌های خارج از $Warning \text{ lines}$ بیش‌تر از ۵ داده بود، نشان‌دهنده اشکال در سیستم است. خطوط $+3S$ و $-3S$ هم $Action \text{ lines}$ نام دارند و اگر حتی یک داده خارج از آن‌ها باشد مشکلی جدی در سیستم وجود دارد و باید سیستم را متوقف کنیم.

نکته: اگر در سیستم، خطای سیستماتیک داشته باشیم، اکثر داده‌ها زیر $\bar{X}$ یا بالای $\bar{X}$ می‌افتند چون خطای سیستماتیک یا جهت مثبت و یا جهت منفی دارد.





۷-۱- اندازه‌گیری خطاهای راندوم و سیستماتیک با داشتن خطاهای متغیرها:

اگر R را تابعی از چند پارامتر یا متغیر A و B و C و... در نظر بگیریم.

مثلاً در $A = \varepsilon bc$ ، A تابعی از متغیرهای ε و b و c است.

اگر R تابعی از پارامترهای A و B و C باشد، با داشتن خطاهای راندوم و سیستماتیک A و B و C می‌خواهیم خطای راندوم و سیستماتیک R را حساب کنیم. قوانین جبر معمولی برای این محاسبات قابل اجرا نیست.

$$R = F(A, B, C, \dots)$$

$$\varepsilon_R = \left(\frac{\partial R}{\partial A}\right)\varepsilon_A + \left(\frac{\partial R}{\partial B}\right)\varepsilon_B + \left(\frac{\partial R}{\partial C}\right)\varepsilon_C \leftarrow \text{خطای سیستماتیک یا معین}$$

مثلاً اگر $R = A + B - C$ باشد و خطای سیستماتیک A و B و C را به‌طور جداگانه داشته باشیم:

$$\varepsilon_R = 1 \times \varepsilon_A + 1 \times \varepsilon_B + (-1) \times \varepsilon_C = \varepsilon_A + \varepsilon_B - \varepsilon_C$$

اگر $R = \frac{AB}{C}$ باشد:

$$\varepsilon_R = \left(\frac{\partial R}{\partial A}\right)\varepsilon_A + \left(\frac{\partial R}{\partial B}\right)\varepsilon_B + \left(\frac{\partial R}{\partial C}\right)\varepsilon_C$$

$$\left(\frac{\partial R}{\partial A}\right) = \frac{B}{C} \quad \left(\frac{\partial R}{\partial B}\right) = \frac{A}{C} \quad \left(\frac{\partial R}{\partial C}\right) = \frac{\text{مشتق مخرج در صورت} \times \text{مشتق صورت در مخرج}}{\text{مخرج به توان ۲}} = \frac{-AB}{C^2}$$

$$\varepsilon_R = \frac{B}{C} \varepsilon_A + \frac{A}{C} \varepsilon_B - \frac{AB}{C^2} \varepsilon_C \Rightarrow \frac{\varepsilon_R}{R} = \frac{\varepsilon_A}{A} + \frac{\varepsilon_B}{B} - \frac{\varepsilon_C}{C}$$

$$\delta_R = \left(\frac{\partial R}{\partial A}\right)^r \delta_A^r + \left(\frac{\partial R}{\partial B}\right)^r \delta_B^r + \left(\frac{\partial R}{\partial C}\right)^r \delta_C^r \leftarrow \text{خطای راندوم یا تصادفی}$$

مثلاً اگر $R = A + B - C$ باشد.

$$\delta_R = 1 \times \delta_A^r + 1 \times \delta_B^r + (-1) \times \delta_C^r \Rightarrow \delta_R^r = \delta_A^r + \delta_B^r + \delta_C^r$$

و اگر $R = \frac{AB}{C}$

$$\delta_R^r = \left(\frac{B}{C}\right)^r \delta_A^r + \left(\frac{A}{C}\right)^r \delta_B^r + \left(\frac{-AB}{C^2}\right)^r \delta_C^r \Rightarrow \left(\frac{\delta_R}{R}\right)^r = \left(\frac{\delta_A}{A}\right)^r + \left(\frac{\delta_B}{B}\right)^r + \left(\frac{\delta_C}{C}\right)^r$$

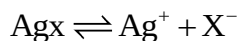
چند رابطه که در شیمی تجزیه کاربرد دارد و به جای مشتق‌گیری، می‌توانیم حفظ کنیم.

$y = a + b$	$\varepsilon_y = \varepsilon_a + \varepsilon_b$	$\delta_y = \sqrt{\delta_a^r + \delta_b^r}$
$y = a - b$	$\varepsilon_y = \varepsilon_a - \varepsilon_b$	$\delta_y = \sqrt{\delta_a^r + \delta_b^r}$
$y = a \times b$	$\frac{\varepsilon_y}{y} = \frac{\varepsilon_a}{a} + \frac{\varepsilon_b}{b}$	$\left(\frac{\delta_y}{y}\right)^r = \left(\frac{\delta_a}{a}\right)^r + \left(\frac{\delta_b}{b}\right)^r$
$y = \frac{a}{b}$	$\frac{\varepsilon_y}{y} = \frac{\varepsilon_a}{a} - \frac{\varepsilon_b}{b}$	$\left(\frac{\delta_y}{y}\right)^r = \left(\frac{\delta_a}{a}\right)^r + \left(\frac{\delta_b}{b}\right)^r$
$y = a^x$	$\frac{\varepsilon_y}{y} = x \frac{\varepsilon_a}{a}$	$\frac{\delta_y}{y} = x \frac{\delta_a}{a}$
$y = \log x$	$\varepsilon_y = 0.434 \frac{\varepsilon_x}{x}$	$\delta_y = 0.434 \frac{\delta_x}{x}$

مثال: حاصل ضرب حلالیت (k_{sp}) برای نمک نقره AgX برابر با $10^{-8} \times (\pm 0.4)$ می‌باشد. خطای تخمینی مربوط به حلالیت محاسبه شده برای نمک AgX در آب چیست؟

نکته: چون در عدد $\pm$ گفته، پس حتماً خطای ران دوم است.

نکته: اگر ± 0.4 در پرانتز باشد، یعنی 10^{-8} هم مربوط به ۴ و هم مربوط به 0.4 است.



$$k_{sp} = [Ag^+][X^-] \Rightarrow [Ag^+]^r = S^r = k_{sp}$$

$$\Rightarrow S = (k_{sp})^{\frac{1}{r}} \Rightarrow S = \sqrt[4]{4 \times 10^{-8}} = 2 \times 10^{-2}$$

حالا می‌خواهیم عدم قطعیت S را هم به دست بیاوریم.

$$\frac{\delta_s}{S} = x \left(\frac{\delta_{k_{sp}}}{k_{sp}} \right) = \frac{1}{2} \left(\frac{0.4 \times 10^{-8}}{4 \times 10^{-8}} \right)$$

$$\delta_s = (2 \times 10^{-2}) \left(\frac{1}{2} \right) \left(\frac{0.4 \times 10^{-8}}{4 \times 10^{-8}} \right) = 1 \times 10^{-5}$$

$$\Rightarrow S = 2 \times 10^{-2} \pm 1 \times 10^{-5} = \underline{2(\pm 0.1) \times 10^{-2}}$$

۸-۱ - طراحی آزمایش:

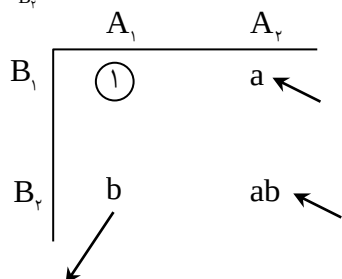
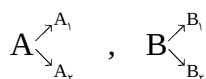
برای بهینه کردن پارامترهای مختلفی که روی سیگنال تأثیر دارند، به روش‌های مختلفی می‌توانیم عمل کنیم:

ساده‌ترین روش $OVAT$ (One Variable At Time) است. اولین فرض در این روش این است که: همه پارامترها مستقل از هم هستند و روی هم تأثیر ندارند. سپس برای هر کدام از پارامترها یک مقدار در نظر گرفته و همه پارامترها، به جز یک پارامتر را در همان مقدار، ثابت نگه داشته و تنها یکی از پارامترها را تغییر می‌دهیم تا عدد بهینه آن را به دست آوریم. سپس برای این پارامتر، عدد بهینه را استفاده کرده و دوباره یک پارامتر دیگر را بررسی می‌کنیم.

۱. چون در این روش تعداد آزمایش‌ها خیلی زیاد می‌شود و ۲. چون نمی‌توان اثر پارامترها را روی هم مشاهده کرد و ۳. چون نمی‌توانیم پارامتر اصلی را تعیین کنیم، این روش خیلی قابل استفاده نیست و برای برطرف کردن این کمبودها، از روش‌های مختلف طراحی آزمایش استفاده می‌کنیم.

ساده‌ترین روش طراحی آزمایش $Full\ Factorial\ Design$ نام دارد. در این روش اگر f تا فاکتور را در L تا level بررسی کنیم باید L^f تا آزمایش طراحی کنیم.

مثلاً اگر دو فاکتور A, B را در دو سطح ۱، ۲ بررسی کنیم $2^2 = 4$ تا آزمایش طراحی می‌کنیم.



چون اختلاف نتیجه آزمایش دوم نسبت به آزمایش ① ناشی از افزایش سطح A است.

اثر افزایش A, B را باهم نشان می‌دهیم.

فقط مربوط به افزایش سطح B است.



در آزمایشی که پاسخ آن را با a نمایش داده‌ایم اختلاف نتیجه نسبت به آزمایش ۱ ناشی از افزایش سطح A می‌باشد. در آزمایش دیگر که پاسخ آن را با b نمایش داده‌ایم اختلاف نتیجه نسبت به آزمایش ۱ ناشی از افزایش سطح A و B می‌باشد و در آزمایش آخر اثر افزایش همزمان A و B را مشاهده می‌کنیم. به آزمایش‌های a, b اثر پارامترهای اصلی می‌گویند. به آزمایش ab اثر متقابل از درجه اول می‌گویند. هرچه درجات آزمایش بیش‌تر شود، اثر آن‌ها روی نتیجه آزمایش کم‌تر می‌شود. مثلاً اگر ۳ فاکتور را هر کدام در دو سطح بررسی کنیم $2^3 = 8$ آزمایش باید طراحی کنیم.

	A_1		A_2	
	B_1	B_2	B_1	B_2
c_1	①	b	a	ab
c_2	②	bc	ac	abc

اثر متقابل از اثر پارامتر اصلی درجه اول اثر متقابل از درجه دوم

در این روش تعداد آزمایش‌ها خیلی زیاد است و وقت و هزینه زیادی صرف می‌شود. ضمن این‌که آزمایش‌هایی که درجه بالاتر دارند، اثر کم‌تری روی نتیجه آزمایش دارند.

طراحی آزمایش Fractional Factorial Design:

در این روش، آزمایش‌های با درجه بالاتر، که در روش Full Factorial Design وجود داشت حذف می‌شود چون هرچه درجات آزمایش بیش‌تر شود، اثر آن‌ها روی نتیجه آزمایش کم‌تر می‌شود. پس نسبت به طراحی آزمایش قبلی تعداد آزمایش‌ها کم‌تر می‌شود و این مزیت این طراحی است.

آزمون‌های یک‌طرفه و دوطرفه:

آزمون‌هایی که در آن احتمال دو رخداد وجود دارد، آزمون دوطرفه نام دارد. مثل این‌که خطای سیستماتیک وجود دارد یا خطای سیستماتیک وجود ندارد. اما آزمون‌های یک‌طرفه آزمون‌هایی است که در آن فقط یک احتمال وجود دارد مثل این‌که استفاده از یک کاتالیزور در یک واکنش فقط سرعت را بیش‌تر خواهد نمود.

برای مقادیر Critical یا Table برای آزمون‌های یک‌طرفه و دوطرفه، تنها یک عدد گزارش شده که براساس آزمون دوطرفه گزارش کرده‌اند و باید بدانیم داده‌های Critical برای آزمون یک‌طرفه را چطور از این جدول به‌دست بیاوریم.

در جداول، داده‌های Critical را برحسب درصد اطمینان بیان کرده‌اند، مثلاً ۹۰٪ (یا همان ۰/۱) یا ۹۵٪ (یا همان ۰/۰۵) و یا ۹۸٪ (که همان ۰/۰۲) و... است. حالا مثلاً اگر بخواهیم مقدار t را برای آزمون یک‌طرفه با سطح اطمینان ۹۵٪ (یا ۰/۰۵) به‌دست بیاوریم. طبق درجه آزادی، مقدار t را برای ۰/۰۵ پیدا می‌کنیم که مربوط به یک آزمون دوطرفه است. حالا برای پیدا کردن t برای آزمون یک‌طرفه ۰/۰۵ (یا هر سطح اطمینانی که داشتیم) را در ۲ ضرب می‌کنیم که مثلاً در مورد ۰/۰۵ $0/05 \times 2 = 0/1$ می‌شود. حالا t مربوط به سطح اطمینان ۰/۱ (یا همان ۹۰٪) را به‌عنوان t برای آزمون یک‌طرفه به کار می‌بریم.

گاهی ممکن است، آزمون یک‌طرفه‌ای داشته باشیم (که برعکس اثر کاتالیزور که اثر افزایش سرعت را بررسی می‌کند) که کاهش یک اتفاق را بررسی می‌کند. مثلاً اگر در یک واکنش یک عامل بازدارنده به کار می‌بریم و می‌خواهیم بدانیم که آیا سرعت را کاهش داده یا نه؟ در این موارد t را که از جدول به روش بالا به‌دست آوردیم را باید با علامت منفی (-) به کار ببریم.



مثال: بررسی کنید آیا اختلاف معنی‌دار بین نتایج به‌دست آمده از دو روشی که در جدول زیر نمایش داده شده، وجود دارد یا نه؟

FTIR	UV-Vis	شماره تولید
83.15	84.63	قرص شماره ۱
83.72	84.38	۲
83.74	84.08	۳
84.20	84.41	۴
83.92	83.82	۵
84.16	83.55	۶
84.02	83.92	۷
83.60	83.69	۸
84.13	83.06	۹
84.24	84.03	۱۰

نکته: چون در هر سری آزمایش هم روش و هم نمونه تغییر یافته نمی‌توانیم همه داده‌ها را با هم بررسی کنیم. و باید تک‌تک هر سری را بررسی کرده و $\bar{d}$ را برای هر کدام به‌دست آورده و سپس t را به‌کار ببریم.

$$t = \frac{\bar{d}\sqrt{n}}{S_d} \quad S_d = 0.5653$$

$$\bar{d} = 0.161$$

$$t_{\text{exp}} = \frac{0.161\sqrt{10}}{0.57} = 0.88 \leq t_{\text{table}} = 2.26 \quad \text{اصل صفر تأیید می‌شود} \quad \Leftarrow 2/26$$

پس اختلاف معنی‌دار وجود ندارد.

اختلاف ناشی از خطای راندوم است.

در این مثال نتایج جفت-جفت با یکدیگر متناظر می‌باشند.

فصل دوم

تصادلات شیمیایی

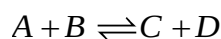
آنچه که در این فصل می‌خوانیم

- ☒ مقدمه
 - ☒ حالت‌های استاندارد
 - ☒ اثرات نمک در ثابت‌های تعادل واکنش‌ها
 - ☒ رفتار گونه‌های بدون بار در محلول حاوی نمک
 - ☒ اندازه‌گیری ضریب فعالیت گونه‌های فرار
- با استفاده از GC

فصل دوم

تعادلات شیمیایی

مقدمه



$$k = \frac{a_C \times a_D}{a_A \times a_B} \text{ ثابت تعادل}$$

زمانی یک واکنش شیمیایی به تعادل می‌رسد که پتانسیل شیمیایی مواد اولیه و پتانسیل شیمیایی فرآورده‌ها، یکسان باشد:

$$\mu_A + \mu_B = \mu_C + \mu_D$$

هدف، به دست آوردن ثابت تعادل است.

پتانسیل شیمیایی برابر است با:

$$\mu_i = k_i + RT \ln a_i \text{ پتانسیل شیمیایی استاندارد}$$

$$\Rightarrow k_A + RT \ln a_A + k_B + RT \ln a_B = k_C + RT \ln a_C + k_D + RT \ln a_D$$

$$\Rightarrow k_A + k_B - k_C - k_D = RT \ln \frac{a_C \cdot a_D}{a_A \cdot a_B} \rightarrow \text{ثابت تعادل}$$

$$\Rightarrow \ln k_{eq} = \frac{k_A + k_B - k_C - k_D}{RT}$$

اگر $a_i = 1$ باشد $\mu_i = k_i$ است که معمولاً آن را با μ_i° نشان می‌دهیم که همان پتانسیل شیمیایی استاندارد می‌باشد. تفاوت غلظت و فعالیت: غلظت مقداری از یک گونه است که ما با محاسبه آن را تهیه کرده‌ایم اما ممکن است تمام این گونه‌ها حلال‌پوش نشده و به صورت Free Ion عمل نکنند و هم‌چنان به صورت Ion Pair در محلول باقی بمانند. مقدار Free Ion موجود فعالیت محلول است.

محاسبه k_i کار راحتی نیست. چون گونه‌ها باید حالت استاندارد داشته باشند.

حالت‌های استاندارد ($a_i = 1$):

- در نمونه محلول، باید مقدار حل‌شونده به سمت صفر میل کند.

- در نمونه گازی، باید فشار جزئی گونه به سمت صفر میل کند.

- در نمونه جامد، گونه باید کاملاً خالص باشد.

به‌خاطر رفع این مشکل، فعالیت را از این رابطه حساب می‌کنیم: